

De Jaegher, Kris

Working Paper

Supplier-induced demand as strategic framing

Discussion Papers Series, No. 10-01

Provided in Cooperation with:

Utrecht University School of Economics (U.S.E.), Utrecht University

Suggested Citation: De Jaegher, Kris (2010) : Supplier-induced demand as strategic framing, Discussion Papers Series, No. 10-01, Utrecht University, Utrecht School of Economics, Tjalling C. Koopmans Research Institute, Utrecht

This Version is available at:

<https://hdl.handle.net/10419/322820>

Standard-Nutzungsbedingungen:

Die Dokumente auf EconStor dürfen zu eigenen wissenschaftlichen Zwecken und zum Privatgebrauch gespeichert und kopiert werden.

Sie dürfen die Dokumente nicht für öffentliche oder kommerzielle Zwecke vervielfältigen, öffentlich ausstellen, öffentlich zugänglich machen, vertreiben oder anderweitig nutzen.

Sofern die Verfasser die Dokumente unter Open-Content-Lizenzen (insbesondere CC-Lizenzen) zur Verfügung gestellt haben sollten, gelten abweichend von diesen Nutzungsbedingungen die in der dort genannten Lizenz gewährten Nutzungsrechte.

Terms of use:

Documents in EconStor may be saved and copied for your personal and scholarly purposes.

You are not to copy documents for public or commercial purposes, to exhibit the documents publicly, to make them publicly available on the internet, or to distribute or otherwise use the documents in public.

If the documents have been made available under an Open Content Licence (especially Creative Commons Licences), you may exercise further usage rights as specified in the indicated licence.

Tjalling C. Koopmans Research Institute

Tjalling C. Koopmans



Universiteit Utrecht

**Utrecht School
of Economics**

**Tjalling C. Koopmans Research Institute
Utrecht School of Economics
Utrecht University**

Janskerkhof 12
3512 BL Utrecht
The Netherlands
telephone +31 30 253 9800
fax +31 30 253 7373
website www.koopmansinstitute.uu.nl

The Tjalling C. Koopmans Institute is the research institute and research school of Utrecht School of Economics. It was founded in 2003, and named after Professor Tjalling C. Koopmans, Dutch-born Nobel Prize laureate in economics of 1975.

In the discussion papers series the Koopmans Institute publishes results of ongoing research for early dissemination of research results, and to enhance discussion with colleagues.

Please send any comments and suggestions on the Koopmans institute, or this series to J.M.vanDort@uu.nl

ontwerp voorblad: WRIK Utrecht

How to reach the authors

Please direct all correspondence to the first author.

Kris De Jaegher
Utrecht University
Utrecht School of Economics
Janskerkhof 12
3512 BL Utrecht
The Netherlands.
E-mail: k.dejaegher@uu.nl

This paper can be downloaded at: [http://
www.uu.nl/rebo/economie/discussionpapers](http://www.uu.nl/rebo/economie/discussionpapers)

Supplier-Induced Demand as Strategic Framing

Kris De Jaegher^a

^aUtrecht School of Economics
Utrecht University

January 2010

Abstract

This paper develops a model of supplier-induced demand as strategic framing where the patient has reference-dependent preferences, and the physician can persuade the patient to buy a treatment by affecting the patient's reference point. In the main result, the patient is assumed to have a constant rate of risk aversion (loss aversion) in the gain (loss) region. Two scenarios are treated. In the cure scenario, the physician wants to frame the patient's decision problem such that he prefers to buy a risky curative treatment rather than no treatment. It is shown that the physician is most persuasive if she sets a high reference point, such that the patient sees all payoffs as losses down from that reference point. In the prevention scenario, the physician wants to frame the patient's decision problem such that he prefers a safe preventive treatment rather than no treatment. In this case, the physician's optimal framing either involves framing all payoffs as gains, thus making the patient risk-averse. Alternatively, loss aversion is exploited by framing only the fact of getting ill (rather than having prevented illness) as a loss.

Keywords: supplier-induced demand, prospect theory, strategic framing.

JEL classification: D82, I11.

1. Introduction

Over the last three decades, one of the most popular themes in health economics is the supplier-induced demand hypothesis, stating that physicians whose income gets under pressure (e.g. because of the entry of new physicians) are able to create demand for their own services (Evans, 1974; for a recent overview, see e.g. Peacock and Richardson, 2007). From a theoretical point of view, the theory of demand inducement is nothing but a modified version of the Dorfman and Steiner (1954) model of advertising. By taking some persuasive effort A , the individual physician is able to shift the demand for her services. The reason that the physician does not always induce demand to the full extent lies in a labour-leisure trade-off (Newhouse, 1970), or in ethical preferences that become less strong as income gets more under pressure (De Jaegher and Jegers, 2000). The theory deviates from the neoclassical model, in that the patient's preferences are not assumed to be fixed. Yet, a weakness of this theory is that it does not model how persuasion actually takes place.

Perhaps partly to fill the theoretical gap, and certainly in response to theoretical developments in microeconomics, some health economic models take a rather different approach to supplier-induced demand (Dranove, 1988; Calcott, 1999; De Jaegher and Jegers, 2001). In these models, the patient's preferences are stable, but there is asymmetric information between the physician and the patient. The physician has incentives to overprescribe treatment. The patient has rational expectations, and knows how often the physician overprescribes. Supplier-induced demand may still exist simply because the patient is better off when putting up with overprescription. The advantage of these theories of supplier-induced demand is that they are more sophisticated than the original, advertising-like, models, and that they are therefore able to produce more testable predictions. Yet at the same time, they are somewhat remote of the original concept of supplier-induced demand as a form of persuasion.

The current paper shows that developments in behavioural economics can contribute to construct a simple model of persuasive demand inducement, where the form that persuasion may take is modelled in more detail, and predictions are made about what form persuasion may take in different medical contexts (for recent explorations of possible application of behavioural economics in health economics, see Frank (2004), and Barigozzi and Levaggi (2008)). Our starting point is the following experiment by Tversky and Kahneman (1981), which happens already to be stated in a medical context. In the two versions of this experiment, subjects were confronted with the following scenario:

“Imagine that the U.S. is preparing for the outbreak of an unusual Asian disease, which is expected to kill 600 people. Two alternative programs to combat the disease have been proposed.”

In version 1 of the experiment, subjects were offered the choice between Programs A and B, which are described as follows:

“If Program A is adopted, 200 people will be saved. If Program B is adopted, there is 1/3 probability that 600 people will be saved and 2/3 probability that no people will be saved.”

In version 2 of the experiment, subjects were offered the choice between Programs C and D:

“If Program C is adopted, 400 people will die. If Program D is adopted, there is 1/3 probability that nobody will die and 2/3 probability that 600 people will die.”

When the choice is between A and B, 72% of the subjects choose A; when the choice is between C and D, 78% choose D. This is in spite of the fact that, from the perspective of expected utility maximization, the two examples are perfectly equivalent. Apparently, the experimenter, by *framing* the example in a different manner, can influence the reference point of the subject, and can cause a preference reversal.

Kahneman and Tversky (1979) use these and other experiments to show that 1) people do not think in absolute terms (as suggested by expected utility theory), but rather think in terms of gains and losses with respect to a reference point. The position of this reference point matters for people’s decision. Presumably, in version 1 of the above experiment, the reference point is “every one dies”, and everything that deviates from that is a gain. In version 2, the reference point is “nobody dies”, and everyone who dies is seen as a loss. 2) Making the concept of losses and gains relevant, people think differently about gains and losses. In the experiment, it seems that people are risk averse when they think in gains, preferring the safe programme where 200 people are saved, and are risk loving when they think in losses, preferring the risky programme where there is a probability of 1/3 that nobody dies. The latter effect seems to be obtained because subjects hate the thought that 400 people would die with certainty, and that in programme D, it is at least possible that all get saved. 3) People care more about losses than about gains. Thus, people feel stronger about losing €100 than they do about gaining €100.

Tversky and Kahneman (1992) summarise these different observations in *prospect theory*. In a steady stream of papers, standard economic models have recently been extended to include reference-dependent preferences. A weakness of these models is that it is not often modelled how an agent’s reference point is actually determined. Yet, the above experiment shows one way in which a reference point may be determined, namely through *framing*. Indeed, in the above experiment, the experimenter is able to affect the subject’s decision by framing the formulation of the experiment in such a way that the subject’s reference point is modified. Our starting point is that the physician may in a similar manner be able to affect the patient’s reference point, and affect the patient’s decision in favour of treatment. We refer to such a phenomenon as *strategic framing*. Strategic framing has attracted interest in, among others, political science (e.g. Levy, 2003) and health policy (e.g. Gerend and Cullen, 2008). Yet, while prospect theory has become one of the most influential theories in economics, strategic framing has received little attention there. Exceptions are Just and Wu (2005), who study the effect of strategic framing of compensations in the principal-agent relationship, and Puppe and Rosenkranz (forthcoming), who study the effect that the retail prices that manufacturers suggest in advertisements have on the price sensitivity of loss averse consumers, and indirectly on the prices that retailers are able to set.

We assume that the patient has reference-dependent preferences, and that the physician can set the patient’s reference point. We study such strategic framing in two different contexts. First, framing in favour of a risky curative treatment, where the alternative of not having any treatment yields a safe but relatively low outcome. Second, framing in favour of a safe preventive treatment, where the risky alternative is not following any treatment, which may yield both a better outcome if the patient remains healthy (as no cost of treatment has then been incurred), or a worse outcome if the patient gets ill.

Section 2 sets out the basic aspects of prospect theory that we consider in our model. Section 3 sets out our model of curative treatment, and of preventive treatment. Section 4

considers strategic framing for our two scenarios under several aspects of prospect theory, where these aspects are treated in isolation. The paper ends with a discussion in Section 5.

2. Patient reference-dependent preferences

We assume that the patient's psychic valuation $f(X)$ function of any outcome X obtained in Section 3, takes the following form, proposed by Tversky and Kahneman (1992):

$$\begin{aligned} X \geq R : f(X) &= v(X - R) \\ X < R : f(X) &= -\lambda v(R - X) \\ \text{with } v' > 0, v'' \leq 0, v(0) &= 0 \text{ and } \lambda \geq 1. \end{aligned} \quad (1)$$

R is the patient's reference point, with respect to which he thinks in gains ($X \geq R$) or in losses ($X < R$). The fact that the function v is used both to measure gains and to measure losses takes into account the *reflection effect*: if the patient is risk averse with respect to gains, then the patient is risk loving with respect to losses ($v'' < 0$). For $\lambda > 1$, the patient is loss averse, and losses have a larger impact on him than equally sized gains. An example of f with $v'' < 0$ and $\lambda > 1$ is given in Figure 1. As can be seen, $-v(R - X)$ is the mirror of $v(X - R)$. Loss aversion shifts the psychic valuation function for losses down from the curve $-v(R - X)$ to the curve $-\lambda v(R - X)$. It is over this psychic valuation function that the patient is assumed to take his expectation.

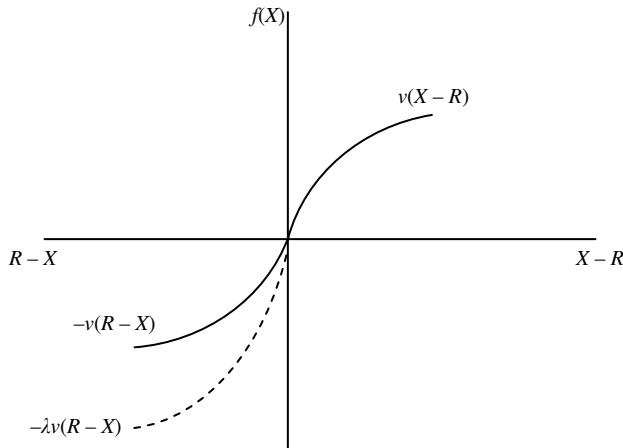


Figure 1

In prospect theory, the psychic valuation function is thus concave for gains, and convex for losses. The degree of concavity in the gain region determines the degree of risk aversion for the losses, and the degree of convexity in the loss region determines the degree of risk lovingness. Concretely, for gains we can apply the well-known Arrow-Pratt measure, where $ARA = |f''(X)|/f'(X)$ is the patient's absolute rate of risk aversion. For losses, we additionally define as $ARL = f''(X)/f'(X)$ the patient's absolute rate of risk lovingness. Unaccounted for by prospect theory, but relevant for strategic framing, is how ARA and ARL depend on X , i.e. how the rate of risk-aversion or of risk-lovingness changes with the outcome. Yet, because the reflection effect, there will be symmetry in the manner in which ARA and ARL change with X . This leads us to distinguish between the following cases for the psychic valuation function. With a *CARA-CARL* psychic valuation function, the patient ARA and ARL does not depend on the level of X . With a *DARA-IARL* psychic valuation function,

the patient becomes less risk averse for high X , and less risk loving for low X . With a *IARA-DARL* psychic valuation function, the patient becomes more risk averse for high X , and more risk loving for low X .

3. Scenarios for the decisions facing the patient

3.1. Prevention

A patient can decide to stay without a preventive treatment, in which case the patient assesses that he will obtain payoff H (healthy) with probability p_G , and payoff S (sick) with probability p_B , where $p_G + p_B = 1$.¹ The preventive treatment yields payoff $H - T$ with certainty, namely the healthy payoff H minus the cost T of treatment. We assume that $H > H - T > S$, so that it is better to incur the cost of treatment than to be ill. The patient has some psychic valuation function $f(\cdot)$ over these outcomes, and prefers to buy the preventive treatment if

$$f(H - T) > p_G f(H) + p_B f(S)$$

Or

$$\frac{p_B}{p_G} > \frac{f(H) - f(H - T)}{f(H - T) - f(S)} = \alpha^P \quad (2)$$

where superscript P refers to the prevention scenario.

3.2. Cure

A patient can decide not to buy any curative treatment, in which case the patient obtains payoff M with certainty. He assesses that the curative treatment will cure him with probability p_G , in which case he obtains payoff $H - T$, where H reflects the healthy outcome and T the cost of treatment, and will with probability p_B obtain outcome $S - T$, where S reflects that the a low payoff from failed treatment, and T again the cost of treatment ($p_G + p_B = 1$). We assume that $H - T > M > S - T$, so that it is better to incur the cost of treatment when treatment is successful, but worse to incur the cost of treatment when treatment fails. The patient has again some psychic valuation function $f(\cdot)$ over these outcomes, and prefers to buy the curative treatment if

$$f(M) < p_G f(H - T) + p_B f(S - T)$$

or

$$\frac{p_B}{p_G} < \frac{f(H - T) - f(M)}{f(M) - f(S - T)} = \alpha^C \quad (3)$$

where superscript C refers to the cure scenario.

¹ This is *secondary prevention*, or medical prevention, in contrast to *primary prevention*, where the latter are actions by the patient that reduce the probability of disease (Kenkel, 2000).

4. Physician strategic framing

We now bring in the physician, who is assumed to be able to influence the patient's treatment/no-treatment decision in both scenarios of Section 3 by framing the decision that the patient faces, in suggesting a reference point with respect to which the patient then thinks in gains and losses. In particular, we assume that the physician always prefers that the patient buys the treatment. To emphasise that we treat a model of persuasion, we do not model any information asymmetry between physician and patient. In as far as information asymmetry is present, this has already been eliminated at the start of our game, as the physician has revealed all information. The patient now faces the decision on whether or not to buy treatment, but is influenced by how the physician frames this decision.

We assume that the physician knows the expression on the right-hand side of equations (2) and (3), but does not know p_B / p_G , namely the odds in favour of the bad outcome assessed by the patient. For this reason, and because the physician always prefers the patient to buy treatment, the physician wants α^P in (2) to be as small as possible, and wants α^C in (3) to be as large as possible. In our model of strategic framing, we thus look for the R that minimises α^P , and that maximises α^C . In order to find these R , we need to know what form the function $\alpha^P(R)$ and $\alpha^C(R)$ take.

A separate analysis for the prevention and cure scenarios is not needed, as $\alpha^P(R)$ and $\alpha^C(R)$ have the same structure. We note that in each of the scenarios, there are only three possible outcomes, namely a low outcome X_l^i , a medium outcome X_m^i , and a high outcome X_h^i , where $i = P, C$ refers to the scenario, and where $X_l^C = S - T$, $X_m^C = M$, $X_h^C = (H - T)$, and $X_l^P = S$, $X_m^P = (H - T)$, $X_h^P = H$. We can then express

$$\alpha^i = \frac{f(X_h^i) - f(X_m^i)}{f(X_m^i) - f(X_l^i)} \text{ for } i = P, C. \quad (4)$$

Because of the asymmetry in the way in which the patient thinks about gains and about losses, rather than formulating one function $\alpha^i(R)$, we in fact need to describe several functions relevant for different levels of R . In particular, depending on the relation between R and the three outcomes, $\alpha^i(R)$ can take four relevant forms, expressed using the valuation function v :

$$\alpha^i(R \leq X_l^i) = \frac{v(X_h^i - R) - v(X_m^i - R)}{v(X_m^i - R) - v(X_l^i - R)} \quad (5)$$

$$\alpha^i(X_l^i < R \leq X_m^i) = \frac{v(X_h^i - R) - v(X_m^i - R)}{v(X_m^i - R) + \lambda v(R - X_l^i)} \quad (6)$$

$$\alpha^i(X_m^i < R < X_h^i) = \frac{v(X_h^i - R) + \lambda v(R - X_m^i)}{\lambda v(R - X_l^i) - \lambda v(R - X_m^i)} \quad (7)$$

$$\alpha^i(R \geq X_h^i) = \frac{\lambda v(R - X_m^i) - \lambda v(R - X_h^i)}{\lambda v(R - X_l^i) - \lambda v(R - X_m^i)} \quad (8)$$

In order to know the shape of $\alpha^i(R)$, we need to know the derivative of $\alpha^i(R)$ for each of the expressions (4)-(8). Rather than doing this for the specification of equation (1) where there

is at the same time a reflection effect and loss aversion, we separately analyze a specification of (1) where there is a reflection effect but no loss aversion (Section 4.1), and a specification of (1) where there is loss aversion but no reflection effect (Section 4.2). In this way, we can see the isolated impact of both the reflection effect, and of loss aversion. This is important as the impacts of these two effects can work in opposite directions. Additionally, inside both the loss and the gain region, $\alpha^i(R)$ may depend on the level of R . As we show below, the sign of $\partial\alpha^i(R)/\partial R$ depends on whether v is *DARA* or *IARA* for the gain region, and on whether v is *DARL* or *IARL* for the loss region. Whether we have a *DARA-IARL* or a *IARA-DARL* psychic valuation function, may again affect whether it is optimal for the physician to frame the payoffs as losses or as gains, so whether f is *DARA-IARL* or *IARA-DARL* is a yet a third effect that may affect the optimal R , on top of the reflection effect and of loss aversion. This is why in Section 4.1, where loss aversion is assumed away, we consider the *CARL-CARA* case, and we thus also assume the effect of the *DARA-IARL* and *IARA-DARL* cases away. In this manner, Section 4.1 purely studies the impact of the reflection effect, as it does not matter to what extent the physician frames the payoffs as gains or losses, but only *whether* she frames the payoffs as gains. Section 4.2 assumes the reflection effect away, so that the patient is everywhere risk neutral, and focuses purely on loss aversion. This is automatically also a *CARL-CARA* case, as the rate of risk aversion and the rate of risk lovingness are both zero. Section 4.3 purely considers the effect of having either the *DARA* or *IARA* cases for the gain region, and the *IARL* or *DARL* cases for the loss region. The focus here thus is on the extent to which the physician should frame the payoffs as gains (losses) once the decision to frame purely as gains (losses) has been made.

4.1 $\lambda = 1$, *CARA-CARL*

As is well-known, the *CARA* utility function takes the form $u(X) = -\exp(-aX)$, where the *ARA* is the constant a . In terms of the psychic valuation function, we can construct in a similar way a *CARA-CARL* psychic valuation function, which has a constant *ARA* = a in the gain region, and a constant *ARL* = a in the loss region ($a > 0$):

$$\begin{aligned} X \geq R: f(X) &= 1 - \exp[-a(X - R)] \\ X < R: f(X) &= -\{1 - \exp[-a(R - X)]\} \end{aligned} \quad (9)$$

where the constant 1 is added to assure that for $R = X$, $f(X) = 0$. Also, in order to exclude loss aversion, we have $\lambda = 1$. We now obtain the following proposition for this case:

Proposition 1. Let the patient's psychic valuation function take the *CARA-CARL* form, and let $\lambda = 1$. Then

- (i) In the prevention scenario, the physician frames all the patient's payoffs as gains ($R \leq S$).
As long as all payoffs are framed as gains, it does not matter how low R is set.
- (ii) In the cure scenario, the physician frames all the patient's payoffs as losses ($R \geq (H - T)$).
As long as all payoffs are framed as losses, it does not matter how high R is set.

Proof:

Step 1. By definition, $\alpha^i(R \leq X_l^i) = \frac{\{1 - \exp[-a(X_h^i - R)]\} - \{1 - \exp[-a(X_m^i - R)]\}}{\{1 - \exp[-a(X_m^i - R)]\} - \{1 - \exp[-a(X_l^i - R)]\}}$. This can be rewritten as $\alpha^i(R \leq X_l^i) = \frac{\exp[-aX_m^i] - \exp[-aX_h^i]}{\exp[-aX_l^i] - \exp[-aX_m^i]}$ and is not a function of R . The same can be checked for $\alpha^i(R \geq X_h^i)$.

Step 2. Given that the patient is risk averse for gains and risk loving for losses, and given Step 1, it follows $\alpha^i(R \leq X_l^i) < \alpha^i(R \geq X_h^i)$.

Step 3. By definition, $\alpha^i(X_l^i < R \leq X_m^i) = \frac{\{1 - \exp[-a(X_h^i - R)]\} - \{1 - \exp[-a(X_m^i - R)]\}}{\{1 - \exp[-a(X_m^i - R)]\} + \{1 - \exp[-a(R - X_l^i)]\}}$.

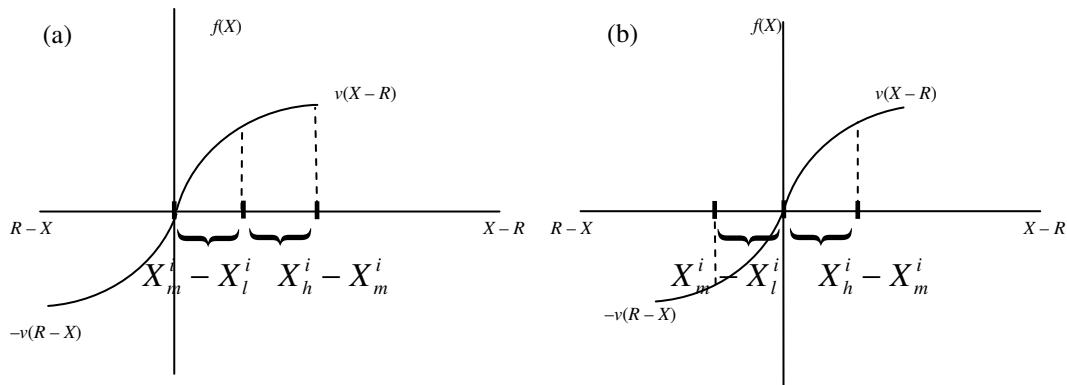
Multiplying both numerator and denominator by $\exp(aR)$, this can be rewritten as $\alpha^i(X_l^i < R \leq X_m^i) = \frac{\exp(aR)[\exp(-aX_m^i) - \exp(-aX_h^i)]}{2 - \exp(-aX_m^i)\exp(aR) - \exp(-aR)\exp(-aX_l^i)}$. The derivative of this

expression with respect to R is larger than zero for $X_l^i < R$. In the same manner, it be checked that $\partial \alpha^i(X_m^i < R < X_h^i) / \partial R > 0$ for $R < X_h^i$.

QED.

The result in Proposition 1 is clear. *CARA-CARL* means that the proportion between any two marginal valuations (which is the form taken by α^i) is constant as long as all payoffs are either framed as gains or as losses. Further, take the case where $X_h^i - X_m^i = X_m^i - X_l^i$. Then for $R = X_m^i$, it is the case that $f(X_h^i) - f(X_m^i) = f(X_m^i) - f(X_l^i)$, meaning that $\alpha^i = 1$ (Figure 2b). For $R \leq X_l^i$, $f(X_h^i) - f(X_m^i) < f(X_m^i) - f(X_l^i)$, so that $\alpha^i < 1$ (Figure 2a). For $R \geq X_l^i$, $f(X_h^i) - f(X_m^i) > f(X_m^i) - f(X_l^i)$, so that $\alpha^i > 1$ (Figure 2c). This is illustrated in Figure 2. As a function of R , $\alpha^i(R)$ takes the form given in Figure 3.

Intuitively, in the prevention scenario, all payoffs are framed as gains in order to make the patient risk averse and prefer a safe, preventive treatment. In the cure scenario, all payoffs are framed as losses in order make the patient risk loving and prefer the risky curative treatment.



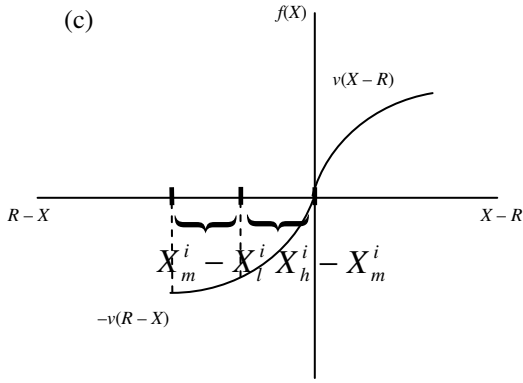


Figure 2

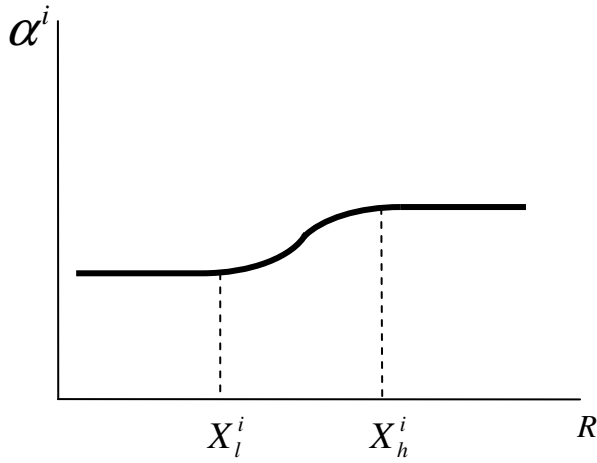


Figure 3

4.2 $\lambda > 1$, $v''=0$

We now consider the case where there is loss aversion ($\lambda > 1$), but where there is no reflection effect, in that the patient is both risk neutral for gains and for losses ($v''=0$). This may be seen as a *CARA-CARL* case with zero risk aversion and risk lovingness. Any strategic framing is then done with the purpose of taking advantage of the patient's loss aversion. In this case, the psychic valuation function is simply linear:

$$\begin{aligned} X \geq R: f(X) &= a(X - R) \\ X < R: f(X) &= -\lambda a(R - X) \end{aligned} \tag{10}$$

where $a > 0$, $\lambda > 1$. This case leads us to Proposition 2.

Proposition 2. Let the patient's psychic valuation function have $v''=0$, and let $\lambda > 1$. Then

(i) In the prevention scenario, the physician puts R exactly at $H - T$.

(ii) In the cure scenario, the physician either puts $R \geq (H - T)$, or $R \leq (S - T)$, where the level of R does not matter as long as it leaves these inequalities valid.

Proof:

Step 1. By definition, $\alpha^i(R \leq X_l^i) = \frac{a(X_h^i - R) - a(X_m^i - R)}{a(X_m^i - R) - a(X_l^i - R)}$. In the same manner,

$\alpha^i(R \geq X_h^i) = \frac{a\lambda(R - X_m^i) - a\lambda(R - X_h^i)}{a\lambda(R - X_l^i) - a\lambda(R - X_m^i)}$. In both these expressions all R 's cancel out, and

they both equal $\frac{X_h^i - X_m^i}{X_m^i - X_l^i}$, as follows simply from the linear form of the psychic valuation function.

Step 2. By definition, $\alpha^i(X_l^i < R < X_m^i) = \frac{a(X_h^i - R) - a(X_m^i - R)}{a(X_m^i - R) + \lambda a(R - X_l^i)}$. This equals

$\frac{X_h^i - X_m^i}{X_m^i - \lambda X_l^i + R(\lambda - 1)}$, so that $\partial \alpha^i(X_l^i < R < X_m^i) / \partial R < 0$. In the same manner,

$\alpha^i(X_m^i < R < X_h^i) = \frac{a(X_h^i - R) + a\lambda(R - X_m^i)}{\lambda a(R - X_l^i) - a\lambda(R - X_m^i)}$. This equals $\frac{X_h^i - \lambda X_m^i + (\lambda - 1)R}{\lambda(X_m^i - X_l^i)}$, so that

$\partial \alpha^i(X_m^i < R < X_h^i) / \partial R > 0$. It follows that the minimal α^i is reached for R such that $R = X_m^i$. QED.

Clearly, framing everything as a loss is equivalent to framing everything as a gain because of the linear form of the psychic valuation form. α^i can be seen as a relation between differential utilities, and for a linear utility function, the relation between any two differential utilities is by definition the same. The reason that α^i is minimised for $R = X_m^i$ can be seen from Figure 4. Consider the horizontal distance between arrows, to which corresponds a differential utility. When there is no loss aversion, this differential utility is the same whatever the reference point, in this case $2d$. However, when as in this case $\lambda = 2$, the part of the differential utility seen as a loss decreases the utility more, so that the differential utility becomes $3d$. Looking at (4), it is clear that the physician can use this principle to increase $f(X_m^i) - f(X_l^i)$. Specific about the linear psychic valuation function is that this can be done without changing $f(X_h^i) - f(X_m^i)$. The physician will set $f(X_m^i) - f(X_l^i)$ at its lowest when $R = X_m^i$. As a function of R , $\alpha^i(R)$ takes the form given in Figure 5.

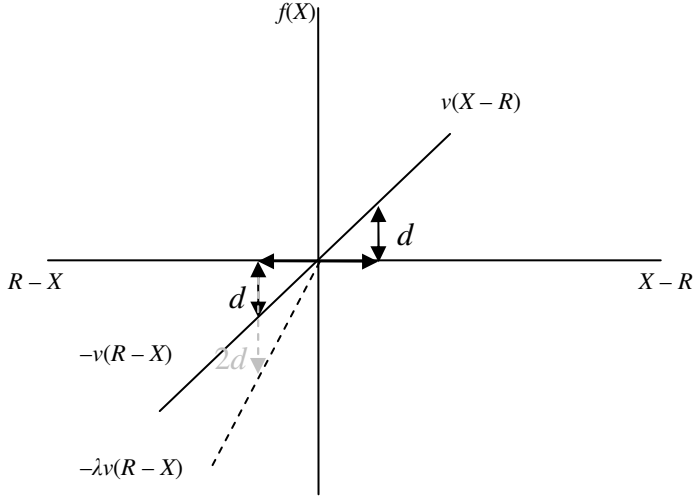


Figure 4

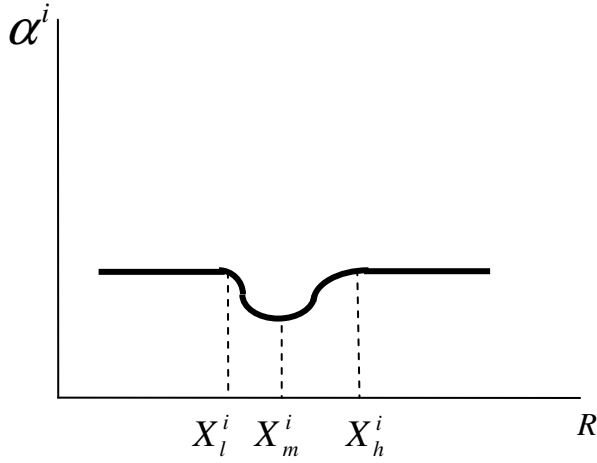


Figure 5

4.3 Physician only able to frame either everything as a gain, or everything as a loss

We now look separately at the case where the physician frames everything as gains (losses). When does changing R have any effect in this case? Let us consider the *DARA-IARL* case. Intuitively, for gain framing, if the patient becomes more risk averse closer to the reference point (increase in R), then $f(X_h^i) - f(X_m^i)$ becomes relatively small compared to $f(X_m^i) - f(X_l^i)$, meaning by equation (4) that α^i decreases. For loss framing, if the patient becomes more risk loving closer to the reference point (decrease in R), then $f(X_m^i) - f(X_l^i)$ becomes smaller compared to $f(X_h^i) - f(X_m^i)$, meaning by equation (4) that α^i increases.

Proposition 3. $\frac{\partial \alpha^i(R \leq X_l^i)}{\partial R} < 0$ (respectively > 0) for a *DARA-IARL* (respectively *IARA-DARL*) psychic valuation function.

Proof:

$$\begin{aligned} \frac{\partial \alpha^i(R \leq X_l^i)}{\partial R} &= \frac{\partial}{\partial R} \left[\frac{v(X_h^i - R) - v(X_m^i - R)}{v(X_m^i - R) - v(X_l^i - R)} \right] = \\ &= - \frac{[v'(X_h^i - R) - v'(X_m^i - R)][v(X_m^i - R) - v(X_l^i - R)]}{[v(X_m^i - R) - v(X_l^i - R)]^2} + \\ &+ \frac{[v(X_h^i - R) - v(X_m^i - R)][v'(X_m^i - R) - v'(X_l^i - R)]}{[v(X_m^i - R) - v(X_l^i - R)]^2} \end{aligned} \quad (11)$$

The sign of the latter expression is equal to the sign of

$$- \frac{[v'(X_h^i - R) - v'(X_m^i - R)]}{[v(X_h^i - R) - v(X_m^i - R)]} + \frac{[v'(X_m^i - R) - v'(X_l^i - R)]}{[v(X_m^i - R) - v(X_l^i - R)]} \quad (12)$$

The two terms in this expression are nothing but the differential versions of two Arrow-Pratt measures. Analogous calculations apply for the case where everything is framed as losses. QED

The effect of changing the reference point when all payoffs are framed as gains or as losses is clear. Yet, when some payoffs are framed as gains and others as losses, then the effect of changing the reference point is ambiguous. This is because one of the differential valuations of the form $f(\cdot) - f(\cdot)$ then contains both a loss and a gain.

5. Discussion

While we have analyzed loss aversion and the reflection effect for a valuation function with a constant rate risk aversion/risk lovingness, it is straightforward to derive results for a model where there is both loss aversion and a reflection effect. The simplest case is the cure scenario. Here, both loss aversion and the reflection effect prescribe that all payoffs should be framed as losses. From the perspective of the reflection effect, the patient should be made risk loving to make him willing to buy the risky curative treatment. Framing the best outcome as a gain will not make the patient more willing to buy this treatment, since the higher differential valuations caused by loss aversion then no longer apply. In the prevention scenario, the results are ambiguous. From the perspective of the reflection effect, all payoffs should be framed as gains, in order to make the patient risk averse and buy the safe preventive treatment. However, from the perspective of loss aversion, the lowest payoff should be framed as a loss, in order to make the patient fear incurring a low health status that could have been prevented. With a combined reflection effect and loss aversion effect, optimal framing therefore depends on the relative strength of these two effects.

The following questions are the subject of future research. *First*, we need to investigate optimal framing when the rate of risk aversion/risk lovingness of the patient is not constant. *Second*, further insights should be obtained if a broader range of scenarios for the choices

facing the patient are treated. *Third*, an aspect of prospect theory not treated in this paper is the overweighting of small probabilities (Kahneman and Tversky, 1992). Once a wider range of scenarios are treated, it can be investigated whether a physician could frame in order to take advantage of this effect as well. A *fourth* question deserving attention is at a more fundamental level. We have assumed that there is no limit on the extent to which a physician can frame. Yet, loss framing makes the patient unhappy, while gain framing makes the patient happy. If the patient avoids physicians who make him feel unhappy, then the ability of physicians to frame as losses will be reduced. This would suggest that persuading a patient to buy preventive treatment is harder than persuading a patient to buy curative treatment.

References

- Barigozzi, F. and Levaggi, R., 2008. Emotions in physician agency. *Health Policy* 88, 1-14.
- Calcott, P., 1999. Demand inducement as cheap talk. *Health Economics* 8, 721-733.
- De Jaegher, K. and Jegers, M. (2000) A model of physician behaviour with demand inducement. *Journal of Health Economics* 19, 231-258.
- De Jaegher, K. and Jegers, M. (2001) The physician-patient relationship as a game of strategic information transmission. *Health Economics* 10, 651-668.
- Dranove, D. (1988) Demand inducement and the physician/patient relationship. *Economic Inquiry* 26, 281-298.
- Evans, R.G., 1974. Supplier-induced demand: some empirical evidence and implications. In: Perlman, M. (ed.) *The Economics of Health and Medical Care*. Macmillan: London; 1974, pp.162-201.
- Frank R., 2001. Behavioral economics and health economics. NBER working paper no. 10881; 2004.
- Gerend, M.A. and Cullen, M., 2008. Effects of message framing and temporal context on college student drinking behavior. *Journal of Experimental Social Psychology* 44, 1167-1173.
- Just, D.R. and Wu, S., 2005. Loss aversion and reference points in contracts. Working paper.
- Kahneman, D. and Tversky, A., 1979. Prospect theory: an analysis of decision under risk. *Econometrica* 47, 263-291.
- Kenkel, D.S., 2000. Prevention. In: Culyer, J., Newhouse J.P. (Eds), *Handbook of health economics*. Elsevier Science: Amsterdam; 2000. p. 1675–1719.
- Levy, J.S., 2003. Applications of prospect theory to political science. *Synthese* 135, 215-241.
- Newhouse, J.P., 1970. A model of physician pricing. *Southern Economic Journal*, 174-183.
- Puppe, C., and Rosenkranz, S. Using suggested retail prices to achieve retail price maintenance by manipulating consumers' reference prices. *Economica*, forthcoming.
- Tversky, A. and Kahneman, D., 1981. The framing of decisions and the psychology of choice. *Science* 211, 453-458.
- Tversky, A. and Kahneman, D. 1992. Advances in prospect theory: cumulative representation of uncertainty. *Journal of Risk and Uncertainty* 5, 297–323.