



**Universitat
Pompeu Fabra**
Barcelona

Department
of Economics and Business

**Economic Working Paper
Series Working Paper No. 1775**

**Listing specs: The effect of
framing attributes on choice**

**Francesco Cerigioni
Simone Galperti**

April 2021

LISTING SPECS: THE EFFECT OF FRAMING ATTRIBUTES ON CHOICE

Francesco Cerigioni*

Simone Galperti†

March 25, 2021

ABSTRACT

Consistent evidence from various settings shows that individual decisions can depend on the order or emphasis used when presenting the attributes of available options. We introduce a model of such framing effects, which we characterize in terms of observable behavior. We show how strategic use of attribute framing affects competition in markets and outcomes of negotiations. In our framework, attribute framing can also cause choice to depend on the order in which options are listed and phenomena similar to the endowment effect. Finally, we use the model to discuss several approaches to welfare analysis.

KEYWORDS: Attribute, Framing, Order, Multi-Attribute Choice, Primacy, Recency, Emphasis, Salience.

JEL CLASSIFICATION: D01, D11, D90

*Department of Economics and Business, Universitat Pompeu Fabra, Ramon Trias Fargas 25-27, 08005, Barcelona, Spain and Barcelona GSE. E-mail: francesco.cerigioni@upf.edu. I am grateful to the Ministerio de Ciencia y Innovacion for the grant PGC2018-098949-B-I00 and financial support through the Severo Ochoa Programme for Centers of Excellence in R&D SEV-2015-0563.

†Department of Economics, University of California, San Diego, 9500 Gilman Dr., La Jolla, CA, 92093. E-mail: sgalperti@ucsd.edu.

In order to construct rich economic models one often needs a model of choice with frames. (Salant and Rubinstein (2008))

1 Introduction

Consistent evidence shows that decisions can depend on the *order* in which the attributes of available alternatives are presented. People’s willingness to pay for a medical treatment can depend on the presentation position of its price (Kjær et al. (2006)). Choices of health plans can depend on how attributes like copay, deductibles, and premium are presented (Ericson and Starc (2016)). Doctors’ diagnoses can depend on the order in which pieces of information are encountered (e.g., Bergus et al. (1995), Cunningham et al. (1997), Chapman and Elstein (2000)). Police investigations and jury decisions can depend on the presentation order of alibi and eyewitness evidence (Dahl et al. (2009)). Consumers’ evaluation of products can depend on the presentation order of their attributes (Kumar and Gaeth (1991), Levav et al. (2010) and references therein). Blake et al. (2018) show that the “purchase funnel”—the order of steps to buy a product—can affect consumers’ decisions.^{1,2}

Motivated by this evidence, we introduce an explicit, decision-theoretic, model of framing of choice items’ *attributes* and its effects on decisions. That is, we take the physical attributes of an item as the given information to be framed; different presentation orders correspond to different frames. We interpret the presentation position as the *observable* emphasis given to the attribute. As such, our model can be applied to other settings where emphasis is given by some graphical means, such as font size or color. We characterize our model axiomatically. We apply it to study how competing firms may strategically frame their products and negotiators may frame offers to strike a deal. Our model also provides a disciplined way to capture self-serving rationalizations and motivated reasoning about the value of items. We discuss several approaches to welfare analysis under our framing effects. The reasons and mechanisms behind them may be multiple and complex, but are of second-order importance for us. Our goal is to develop a model which is consistent with choice data that exhibits those framing effects.³

¹Other papers in marketing and psychology report related evidence, but are too many to have a complete list here. See, e.g., Cornelissen and Werner (2014) and Auspurg and Jäckle (2017) for recent reviews as well as Chrzan (1994), Day and Prades (2010), Day et al. (2012)).

²Even voters’ support for political candidates may depend on the presentation order of their “attributes.” For instance, in the 2016 U.S. Presidential race Hillary Clinton presented herself moving her maiden name (Rodham) to her middle name so as to *emphasize* her independence from her husband (Shafer (2017) and <https://www.theguardian.com/media/mind-your-language/2016/nov/11/winning-words-the-language-that-got-donald-trumpf-elected>).

³Some evidence suggests that experience may weaken framing effects (Levin and Gaeth (1988), Kumar and Gaeth (1991)). This does not undermine the importance of modeling such effects, as

To fix ideas, consider an example. Health plans are often presented in tables where each row is an attribute (copay, deductibles, premium, etc.) and each column is a plan. Let N be the number of attributes (hence, rows) and let $f(i)$ be the attribute in row i . The assignment f of attributes to rows is our frame. A plan is then a vector $x_f = (x_{f(i)})_{i=1}^N$, where $x_{f(i)}$ is the level of attribute $f(i)$. Mainstream choice theory assumes that f is irrelevant. We allow f to affect which plan a customer chooses (hence, each plan’s market share). For instance, this may change if the premium is moved from the first to the last row.

As a first pass at studying these framing effects, we introduce and characterize a baseline model called the *attribute-framing model*. If all available items have the same frame f , the decision-maker chooses the x_f that maximizes

$$\sum_{i=1}^N \alpha(i) u_{f(i)}(x_{f(i)}). \quad (1)$$

Each $u_{f(i)}$ is a utility function that captures the decision-maker’s (stable) underlying tastes for each attribute. The weight function α depends on the attribute’s position and is the heart of our model. Depending on its shape, we can capture several empirical regularities in how attribute orders influence choice by changing the marginal rate of substitution between attributes. We characterize which α give rise to *recency* (*primacy*) effects, whereby attributes presented later influence more (less) the evaluation of items than do earlier attributes.⁴ One interpretation is that α reveals whether the decision-maker perceives the attributes presented later or earlier as being emphasized. Primacy effects are consistent with the old adage “first impressions matter” and with the “leader-driven” effect: An item that starts ahead in terms of the first attribute is more likely to be chosen (Carlson et al. (2006)). We also characterize what it means to be more susceptible to these effects. Other forms of α are possible, depending on which presentation position carries *relatively* more weight for the decision-maker.

Of the many possible applications of our model, we present two. The first analyzes how firms can frame products to influence the competition they face. We show that by ordering their attributes—and thus giving each different emphasis—firms can create *fictitious* product differentiation that results in higher market prices and profits.⁵ Sometimes incumbents can also use framing to deter entry in a market, even though this involves trade-offs. In a nutshell, the incumbent has to make its product “look good, but not too much,” which results in lower profits than in uncontested markets. We relate our findings to the industrial-organization literature on obfuscation

many consequential decisions in life happen rarely and with no feedback (e.g., serving as a juror, buying a house, diagnosing a rare disease, or choosing a health or pension plan).

⁴For evidence on the primacy and recency effects, see Kardes and Herr (1990), Haugtvedt and Wegener (1994), Payne et al. (2000), Bond et al. (2007), Ge et al. (2011).

⁵The contract-theory literature has examined strategic framing in buyer-seller relationships, where framing is assumed to influence the buyer’s willingness to pay (see, e.g., Ostrizek and Shishkin (2018), Salant and Siegel (2018)). Our model can provide a foundation for how this influence works.

strategies.

Our second application analyzes framing in negotiations. Framing is often regarded as an important negotiation technique, which can help break an impasse and find solutions. Despite this, modeling framing in negotiations and how it is used has been challenging. We study negotiations that involve multiple attributes whose ideal levels differ between parties. We show how the proposing party chooses an offer and frames it so as to strike the best deal based on the attributes' importance and conflict with the other party. In our model, it can happen that the proposer strategically de-emphasizes some attribute, despite its being very important, so as to weaken the impact of the strong disagreement with the receiver on that attribute. Moreover, framing can emerge as a tool to break an impasse. Finally, we find that more susceptibility to framing not always benefits the proposer.

Our theory covers situations where items are presented using different frames. The idea is to formalize how the decision-maker responds to the frames in such situations as an hypothesis and use representation (1) to predict its observable implications. We can then test which hypothesis better describes behavior and adopt it in combination with (1). We discuss several illustrative hypotheses. The decision-maker may evaluate each item using its frame, or he may *reframe* all items using a common frame. This frame may depend on the menu of available items, connecting attribute framing with list orders effects. For example, it can be the first encountered frame or the most frequent frame in the menu.⁶ A key property here is that the decision-maker uses one of the existing frames, as opposed to forming a new one not already present in the menu. Thus, choices continue to be affected by the exogenously given frames.

The possibility that decision-makers reframe items in their mind can give rise to an interesting phenomenon, which may be called *self-serving rationalization* via framing. A large body of evidence finds that people often engage in such motivated reasoning (Bénabou and Tirole (2016)). We can capture this with our model in a disciplined manner: Consciously or not, our decision-maker can emphasize attributes by the order in which she considers them in her mind, but only to a certain degree pinned down by α . This process might happen *ex-post* to justify an action—like purchasing an expensive gadget—or *ex-ante* as a self-motivation to do something—like exercising. This general phenomenon can give new insights into well-known behavioral patterns similar to the endowment effect (Thaler (1980), Kahneman et al. (1991)). To rationalize the acquisition of a gadget, the owner may represent its attributes in her mind using the frame that maximizes how much she values it. Doing so will in general raise her willingness to accept above her initial willingness to pay for the gadget. We discuss consistent evidence in Section 5.

Self-serving rationalization offers a steppingstone to addressing the general ques-

⁶Krosnick and Alwin (1987) find that the first item on a menu “may establish a cognitive framework or standard of comparison that guides interpretation of later” items.

tion of how to do welfare analysis when frames affect choice. We can use our model, identified by choice data, to define the utility the decision-maker experiences once owning the item. Measuring welfare based on experienced utility—as opposed to decision utility—is a typical approach in behavioral economics (Kahneman et al. (1997); Kahneman et al. (1999)). Another possibility is to apply the choice-based approach of Bernheim and Rangel (2009) to our setting. Finally, we propose a welfare measure that averages across frames.

The axiomatic characterization of our attribute-framing model rests on the assumption that we can observe which items a decision-maker chooses as well as how their attributes are framed. We rely exclusively on choices from menus whose items are all framed in the same way. We consider the rich domain of lotteries over items to identify the weight function α . Our main axiom delivers this identification by considering appropriately constructed swaps of attribute positions.

Section 7 generalizes our theory in several ways. We consider a non-separable model where how much the decision-maker weighs attributes presented later depends on how good earlier attributes are. For example, she may overlook later attributes if the first ones already give her high utility. We axiomatize this model in the domain of stochastic choice, as this offers more structure for this task. It also allows us to showcase how to introduce framing effects in this domain, which is often used in practice. We discuss this in the context of a Luce framework, the perturbed-utility framework of Fudenberg et al. (2015), and the rational-inattention framework of Matějka and McKay (2015). The latter can accommodate both list-order and attribute-order effects, where the latter may drive the former.

Related Literature. The importance of framing for decision making has been recognized at least since the seminal work of Kahneman and Tversky (1979) and Tversky and Kahneman (1981). The literature has then evolved in two directions. Some papers developed general frameworks to think about framing (Salant and Rubinstein (2008), Bernheim and Rangel (2009), Salant (2011)). Others have focused on modeling specific ways in which frames can influence choice, especially for applications. Our paper belongs to this second strand.

This literature considers several forms of framing. Following Kahneman and Tversky (1979), many papers have studied presenting choices as gains or losses. Another form is ‘mental accounting’ in relation to saving and investment decisions (Thaler (1985), Thaler (1990)). Several papers have modeled salience, where the weights the decision-maker gives to attributes can depend on how each *stands out* from the others in a menu (e.g., Kőszegi and Szeidl (2012), Bordalo et al. (2012), Bordalo et al. (2013b), Bordalo et al. (2013a)). To the best of our knowledge, our paper is the first to study theoretically framing as the presentation order of items’ attributes. As such, it can be interpreted as a complementary theory of salience: While in those papers salience depends on how much an attribute varies across items and in relation to other

attributes, in our paper it depends on the presentation position of an attribute.⁷ Rubinstein and Salant (2006) considered the effects of items' position on a list. In some settings, attribute-order effects can be a driver of list-order effects (see Section 4).

The cognitive-science literature has studied how people seem to form their preferences at the moment of elicitation (Lichtenstein and Slovic (2006)). One interpretation of our model is that the decision-maker has well-defined tastes for each attribute. However, when it comes to combining them to evaluate of an item, she lets the attributes' framing influence her evaluation. In this way, elicitation methods can influence her choices right when she makes them. Although this may seem to undermine the discovery of decision-makers' true tastes, we show how observing choices across frames can overcome this issue.

2 The Model

The choice objects are called *items*. Each is described by the attributes in a set A . For example, cars are described by make, model, year, color, style, size, power train, etc.. We assume that $|A| = N$ is finite and $N \geq 2$. Each attribute $a \in A$ can take multiple levels, denoted by the set L_a . An item consists of a list of the level of all its attributes and is denoted by $x = (x_{a'}, \dots, x_{a''}) \in X = \times_{a \in A} L_a$.

We want to allow for the possibility that the order in which attributes are presented affects choice. To this end, we introduce the notion of *attribute-frame*, which we define as follows. Let F be the set of all bijections from $\{1, \dots, N\}$ to A . For every frame $f \in F$, $f(i)$ is the attribute presented in the i th position of the item description. We later discuss other interpretations of f , for instance in terms of observable emphasis. We denote an item x under frame f as

$$x_f = (x_{f(i)})_{i=1}^N,$$

where $x_{f(i)} \in L_{f(i)}$ for all i . The set of all items under frame f is

$$X_f = \times_{i \in N} L_{f(i)}.$$

Subsets of items are called menus. If all items in a menu are described according to f , we call it an f -*menu* and denote it by $M_f \subseteq X_f$. Hereafter, we will use the generic term 'menu' to refer to a general subset M of items with the properties that some items are not described using the same frame and no item appears more than once possibly under different frames.

Example 1 Suppose items are health plans described by copay, deductibles, and premium: $A = \{c, d, p\}$. Each attribute can be high or low: $L_a = \{h, l\}$. A frame is

⁷In their study of health-insurance decisions, Ericson and Starc (2016) argue that their evidence "suggests that theories of salience that only rely on the attributes of choice (rather than how they are presented) miss important elements of salience."

the order in which a plan description presents its attributes. This may be $\{d, c, p\}$ for frame f and $\{p, c, d\}$ for f' . Thus, the same plan with a high premium, a high copay, and low deductibles may be presented as $x_f = (l_d, h_c, h_p)$ or $x_{f'} = (h_p, h_c, l_d)$. We allow this order affect choice. Health plans are often presented in a table where, say, the rows are the attributes and the columns the plans. Viewed as menus, such tables always present all items using the same frame.

In the first part of the analysis we will consider only menus whose items are all presented using the same f . Such menus are interesting in themselves and widespread in practice. In online stores items are often organized in tables, whether they are health plans, investment products like ETFs, or electronic devices. Also, f -menus allow us to focus on the effects of the presentation order of attributes, removing other phenomena that may arise for general menus. For instance, choices may also depend on which frames appear in a menu. We will return to general menus in Section 4.

Our baseline model of frame-dependent choice is as follows. Section 7 provides its axiomatization. Let $c(M_f)$ be the set of choices from menu M_f . We assume that $c(M_f)$ is well-defined and nonempty for every M_f .

Definition 1 *An attribute-framing (AF) choice model is defined by a pair (α, u) , where $u = (u_a)_{a \in A}$, each $u_a : L_a \rightarrow \mathbb{R}$ is an attribute utility function, and $\alpha : \{1, \dots, N\} \rightarrow \mathbb{R}_{++}$ is a weight function that together satisfy, for all $f \in F$ and M_f ,*

$$c(M_f) = \arg \max_{x_f \in M_f} \sum_{i=1}^N \alpha(i) u_{f(i)}(x_{f(i)}).$$

The interpretation is that the decision-maker derives utility from each attribute, which he has to aggregate somehow. In the model, he does so linearly in a way that depends on the presentation order of the attributes through the weights α . Thus, attributes presented early can receive higher or lower weight than later attributes. This means that marginal rates of substitution between attributes can depend on their presentation position. For simplicity, hereafter we will always normalize α so that $\sum_{i=1}^N \alpha(i) = 1$.⁸ The additive structure of our AF model is intuitive and tractable. It is also widely used in studies of multi-attribute decision making.⁹ We will relax it in Section 7.

In this baseline model, α fully controls the effects of the presentation order of attributes. By varying the form of α we can capture various effects. We will focus on

⁸This model is related to the so-called “expectancy value model” of framing in psychology (e.g., Ajzen and Fishbein (1980); Nelson et al. (1997)).

⁹See, e.g., Lancaster (1966), McFadden (1973), Gorman (1980), Smith and Brynjolfsson (2001), Allen and Rehbeck (2016).

the main effects consistently found in the empirical literature: primacy and recency effects.¹⁰ We define them here and give a behavioral characterization in Section 7.

Definition 2 *Given the AF model (α, u) , the decision-maker exhibits primacy (recency) effects if α is strictly decreasing (increasing).*

We may want to compare decision-makers in terms of how susceptible they are to attribute framing. For this comparison to be meaningful, their tastes over attributes should be the same.

Definition 3 *Let (α^1, u^1) and (α^2, u^2) be AF models of decision-makers 1 and 2. Suppose for all $a \in A$, $u_a^1 = \gamma u_a^2 + \zeta_a$ for some $\gamma > 0$ and $\zeta_a \in \mathbb{R}$. Decision-maker 1 is more susceptible to recency (primacy) effects than decision-maker 2 is if*

$$\frac{\alpha^1(i+1)}{\alpha^1(i)} \geq (\leq) \frac{\alpha^2(i+1)}{\alpha^2(i)}, \quad i = 1, \dots, N-1.$$

In words, decision-maker 1 is more susceptible to recency (primacy) effects than decision-maker 2 is if α^1 increases (decreases) faster than does α^2 .

After discussing some aspects of this model, we immediately present several applications. We will later return to the foundations and generalizations to cover general menus and non-separable effects of attributes and frames on choice.

Discussion. It is worth highlighting a premise of our framework. When we study f -menus it might look as if we are assuming that the decision-maker rigidly follows the order in which attributes are presented. In fact, what we are assuming is that the exogenously given f influences the individual's preferences in a *consistent* way.

While our primary interpretation of frames is the order in which attributes of an item are physically presented, other interpretations are possible. One is to view each position $i = 1, \dots, N$ as a degree of emphasis or salience that the presentation of an item gives to its attributes (possibly in combination with the attributes' ordering). For instance, emphasizing may involve highlighting an attribute with color and font size or by placing it in a prominent position on an ad page. The key assumptions are that such attribute frames should (1) create an objective observable order and (2) work in terms of relative effects. That is, for instance, increasing all fonts proportionately does not change anything because the relative emphasis stays the same.¹¹ With these assumptions in mind, one can interpret and apply our model in a variety of settings where the emphasis given to attributes is part of the observable data.

¹⁰See Kardes and Herr (1990), Haugtvedt and Wegener (1994), Payne et al. (2000), Bond et al. (2007), Ge et al. (2011).

¹¹To see that salience is a relative concept, see Milosavljevic et al. (2012) and references therein.

3 Framing at Work: Illustrative Applications

3.1 Look Good, But Not Too Much: Strategic Framing and Market Competition

This section analyzes how firms can frame products to influence the competition they face. We show that by ordering their attributes—and thus giving each different emphasis—firms can create a *fictitious* product differentiation that results in higher profits. Sometimes incumbents can also use framing to deter entry, even though doing so involves some trade-offs. By describing how attribute orders affect choice, our model allows us to provide insights into when and how firms achieve these outcomes.

We start from a canonical model of vertical differentiation (Tirole (1988), Ch 7.5.1). Each of two firms, the incumbent and the entrant, manufactures a product. Their equal marginal cost is normalized to 0. Entry costs $K > 0$. Each consumer demands one product. The payoff of a product with intrinsic value $v > 0$ and price t is

$$\theta v - t.$$

The taste parameter θ is uniformly distributed across consumers between $\underline{\theta} \geq 0$ and $\bar{\theta} = \underline{\theta} + 1$. Note that the lower $\underline{\theta}$ is, the more heterogeneous the consumers are in relative terms. We assume that $\underline{\theta} < 1$. The payoff of buying nothing is 0.

We modify this model as follows. The products have three attributes: price (p), reliability (r), and build quality (b). The payoff of a product under frame f is

$$\theta[\alpha_f(r)x_r + \alpha_f(b)x_b] - \alpha_f(p)x_p,$$

where $\alpha_f(a) = \alpha(f^{-1}(a))$ for $a \in \{b, p, r\}$. Thus, the products' intrinsic value depends on the level of r and b . The presentation order affects each attribute's weight in the payoff.¹² We continue to assume that only θ differs across consumers, while α , u_b , and u_r are the same. These assumptions allow us to focus on the interaction between framing and vertical differentiation. We will analyze primacy effects: $\alpha(1) > \alpha(2) > \alpha(3)$. The analysis for other forms of α is analogous. To simplify notation, denote the incumbent's and the entrant's product by $I = (I_b, I_p, I_r)$ and $E = (E_b, E_p, E_r)$.¹³

Each product's reliability and build quality are exogenous. One interpretation is that the engineers of each firm have been able to develop its product to a certain degree and now the marketing team has to choose how to sell it. To make this interesting, we assume that $E_r > 0$, $I_b > 0$, and

$$I_r - E_r = E_b - I_b \equiv \delta > 0.$$

¹²Note that we can write this payoff in terms of our AF model in Definition 1 as $\sum_{i=1}^3 \alpha(i)u_{f(i)}(x_{f(i)}; \theta)$, where $u_r(x_r; \theta) = \theta x_r$, $u_b(x_b; \theta) = \theta x_b$, and $u_p(x_p; \theta) = -x_p$.

¹³We use this lighter notation rather than $x^I = (x_b^I, x_p^I, x_r^I)$ and $x^E = (x_b^E, x_p^E, x_r^E)$.

That is, the differences in reliability and build quality between products offset each other. We add this property because it implies that the products are overall equivalent for consumers not affected by framing (i.e., if α is constant). In this case, entry would lead to standard Bertrand competition.

The timing is as follows: First, the incumbent chooses f , which also applies to the entrant's product. Second, the entrant decides whether to enter. If it does, the firms compete in prices à la Bertrand; otherwise, the incumbent sets its monopoly price. We assume that the incumbent controls how to frame both products because we are interested in how it can use framing to influence its competitive landscape. Also, since the incumbent is established in the industry, it alone may have the resources to run ads that fix f . In Appendix D we allow the entrant to choose f for its product. This raises the issue of how consumers choose if firms adopt different frames, which we tackle in Section 4. As expected, this extension can limit the incumbent's influence, but does not change the insights of this simpler setup.

Framing allows the incumbent to differentiate its product by emphasizing its superior reliability and de-emphasizing its inferior build quality. This is a realistic and expected strategy, of course. The point is that our model allows us to describe and analyze it. Let the difference in intrinsic value between the incumbent's and the entrant's product under f be

$$\Delta_f = [\alpha_f(r) - \alpha_f(b)]\delta,$$

which is positive if and only if f presents attribute r before b (i.e., $f^{-1}(r) < f^{-1}(b)$). As in Tirole (1988), we assume that after entry both firms have a positive market share in equilibrium.¹⁴ All proofs appear in Appendix A.

Lemma 1 (Framing-Driven Differentiation Equilibrium) *Fix f . After entry the equilibrium prices and profits (denoted by φ^o for oligopoly) satisfy the following properties. Let*

$$\begin{aligned} x_p &= \frac{|\Delta_f|}{3\alpha_f(p)}(2 + \underline{\theta}) & \varphi^o(x_f) &= \frac{2 + \underline{\theta}}{3}x_p \\ y_p &= \frac{|\Delta_f|}{3\alpha_f(p)}(1 - \underline{\theta}) & \varphi^o(y_f) &= \frac{1 - \underline{\theta}}{3}y_p. \end{aligned}$$

If $\Delta_f > 0$, then $x = I$ and $y = E$. If $\Delta_f < 0$, then $x = E$ and $y = I$.

¹⁴A sufficient condition is that

$$\delta \left(\frac{1}{\underline{\theta}} - 1 \right) \leq 3 \min_{f \in F} \frac{\alpha_f(r)E_r + \alpha_f(b)E_b}{|\alpha_f(r) - \alpha_f(b)|},$$

which holds if the products' intrinsic difference or the consumers' heterogeneity is sufficiently small (i.e., δ is low or $\underline{\theta}$ is high).

Lemma 1 offers several insights. First, the differentiation created by framing allows the incumbent to make higher profits, of course by presenting its product as superior to the competitor’s. In particular, the incumbent captures the top of the market (i.e., the consumers with high θ). Thus, by controlling the product frame, firms can not only boost their appeal with all consumers, but also capture the most profitable ones.

A second insight is that de-emphasizing prices can raise profits only if products are differentiated. Imagine that $\delta = 0$. Even if f presents p at the end, the equilibrium profits are zero—despite consumer heterogeneity in θ . Even if framing nudges them to weigh prices less, Bertrand competition neutralizes this by erasing any profit.

Several papers find consistent evidence on how changing the emphasis on, or salience of, prices affects product choice. In Lynch and Ariely (2000), consumers buy higher quality wine when prices are displayed not alongside product descriptions, but only later at checkout. Also, price elasticities are higher for undifferentiated wines (akin to small δ) independently of price presentation, and when it is harder to notice product differentiation (akin to small $|\Delta_f|$). In Blake et al. (2018), postponing purchase fees for concert tickets until checkout induces consumers to buy higher quality tickets and increases revenues from such tickets. In Smith and Brynjolfsson (2001), perceived differences between otherwise homogeneous goods help explain markups and price dispersion in online markets. Through the lens of primacy effects, postponing prices may be interpreted as akin to obfuscation strategies that weaken price sensitivity by creating search frictions. In Ellison and Ellison (2009), firms endogenously create such frictions to soften price competition and raise markups.¹⁵ A distinctive feature of our model is that framing creates the very product differentiation that allows firms to exploit these strategies.

A third insight is that the incumbent creates a positive *framing externality* on the entrant. By making its product “look better,” the incumbent weakens the competition from the entrant and charges higher prices. This leaves the bottom consumers exclusively for the entrant, which can then earn a profit. Thus, an incumbent faces a trade-off in emphasizing strengths and de-emphasizing weaknesses of its product.

¹⁵Ellison and Ellison (2009) argue that “obfuscation could [...] involve [...] altering [the consumers’] utility functions in a way that raises equilibrium profits.” They also find that obfuscation raises the price elasticity for low-quality products, but lowers it for high-quality products. In our model, if $\Delta_f > 0$, the price elasticities of the entrant’s and incumbent’s demand are (see Appendix A)

$$\frac{E_p}{I_p - E_p - \theta \frac{\Delta_f}{\alpha_f(p)}} \quad \text{and} \quad \frac{I_p}{\bar{\theta} \frac{\Delta_f}{\alpha_f(p)} - I_p + E_p}.$$

Note that lowering $\alpha_f(p)$ raises the first, but lowers the second. Ellison and Ellison (2009) note that it is hard to know what the elasticities would be absent obfuscation. Our model could provide such counterfactuals given estimates of α .

Doing so maximizes its value for all consumers—hence, the monopoly profits. However, it also emphasizes differences from potential competitors, thereby rendering entry more attractive for them. The best framing strategy may then differ between contested (low K) and uncontested (high K) markets.

To characterize the incumbent's optimal framing strategy, we need to know how it ranks frames as a monopolist. As we will see, it suffices to focus on three frames:

i	f^m	f^*	f_*
1	r	r	p
2	b	p	r
3	p	b	b

Letting φ^m denote the monopoly profits, we get (Lemma 2 in Appendix A)

$$\varphi^m(I_{f^m}) > \varphi^m(I_{f^*}) > \varphi^m(I_{f_*}).$$

We will focus on settings where the incumbent always prefers to remain a monopolist: $\varphi^m(I_{f^m}) > \max_{f \in F} \varphi^o(I_f)$, which is equivalent to

$$\frac{\alpha(1)I_r + \alpha(2)I_b}{\alpha(3)} > \delta h(\underline{\theta}) \max_{f \in F} \frac{|\alpha_f(r) - \alpha_f(b)|}{\alpha_f(p)} \quad \text{where} \quad h(\underline{\theta}) = \frac{4}{9} \left(\frac{2 + \underline{\theta}}{1 + \underline{\theta}} \right)^2 \quad (2)$$

and thus holds if the products' intrinsic difference or the consumers' heterogeneity is sufficiently small (i.e., δ is low or $\underline{\theta}$ is high). A monopolist simply uses framing to emphasize what its product delivers and de-emphasize what one has to pay for it.

By contrast, under the threat of competition framing becomes a tool to emphasize strengths and de-emphasize weaknesses *relative* to the entrant. Thus, the optimal strategy is more nuanced and depends on how consumers respond to framing. Define

$$\bar{\alpha}(2) = \frac{[\alpha(1)]^2 + [\alpha(3)]^2}{\alpha(1) + \alpha(3)} \quad \text{and} \quad \underline{\alpha}(2) = \alpha(1) - \alpha(3),$$

which satisfy $\bar{\alpha}(2) > \underline{\alpha}(2)$. We first characterize the cases where framing can never help deter entry. In these cases, either entry is not a threat and the incumbent uses f^m , or entry is inevitable and the incumbent uses f^m or f^* .

Proposition 1 (Optimal Framing without Entry Deterrence)

If $K > \varphi^o(E_{f^m})$, the incumbent chooses frame f^m and remains a monopolist. If $K \leq \min\{\varphi^o(E_{f^m}), \varphi^o(E_{f_})\}$, the incumbent cannot deter entry; it chooses f^m if $\alpha(2) < \underline{\alpha}(2)$ and f^* if $\alpha(2) > \underline{\alpha}(2)$.*

Note that $\varphi^o(E_{f_*}) < \varphi^o(E_{f^*})$, while $\varphi^o(E_{f_*}) > \varphi^o(E_{f^m})$ if and only if $\alpha(2) > \bar{\alpha}(2)$.

When deterring entry is impossible, the incumbent presents its strengths first, but may present its price *before* its weaknesses. If consumers underweight the second

attribute only a little ($\alpha(2) > \underline{\alpha}(2)$), the incumbent is forced to present its weakness after its price to optimally differentiate its product from the entrant's. If instead consumers underweight a lot the second attribute ($\alpha(2) < \underline{\alpha}(2)$), the incumbent can effectively de-emphasize both its weakness and price, thus presenting the price last.

Next, we describe when the incumbent uses framing to deter entry. This always involves frames that are suboptimal from the monopolist's viewpoint. The incumbent shows its strengths before its weaknesses, but may again emphasize its price by presenting it earlier—even first. In so doing, the incumbent forgoes some of its appeal to all consumers in exchange for saving its monopolistic position, by rendering the market less attractive for the entrant. Thus, framing is used to *weaken* the power of differentiation, should entry occur.¹⁶

Proposition 2 (Optimal Framing with Entry Deterrence)

I: If $\bar{\alpha}(2) > \alpha(2) > \underline{\alpha}(2)$, then $\varphi^o(E_{f^*}) > \varphi^o(E_{f^m}) > \varphi^o(E_{f_*})$. In this case, if $\varphi^o(E_{f^m}) \geq K > \varphi^o(E_{f_*})$ and

$$\frac{\alpha(2)I_r + \alpha(3)I_b}{\alpha(1)} \geq \left[\frac{\alpha(1)}{\alpha(2)} - \frac{\alpha(3)}{\alpha(2)} \right] \delta h(\underline{\theta}), \quad (3)$$

the incumbent chooses f_* and remains a monopolist. Otherwise, it chooses f^* and the entrant enters.

II: If $\alpha(2) < \underline{\alpha}(2)$, then $\varphi^o(E_{f^m}) > \varphi^o(E_{f^*}) > \varphi^o(E_{f_*})$. In this case, we have that

– if $\varphi^o(E_{f^m}) \geq K > \varphi^o(E_{f^*})$ and

$$\frac{\alpha(1)I_r + \alpha(3)I_b}{\alpha(2)} \geq \frac{\alpha(2)}{\alpha(3)} \left[\frac{\alpha(1)}{\alpha(2)} - 1 \right] \delta h(\underline{\theta}), \quad (4)$$

the incumbent chooses f^* and remains a monopolist;

– if $\varphi^o(E_{f^*}) \geq K > \varphi^o(E_{f_*})$ and

$$\frac{\alpha(2)I_r + \alpha(3)I_b}{\alpha(1)} \geq \frac{\alpha(2)}{\alpha(3)} \left[\frac{\alpha(1)}{\alpha(2)} - 1 \right] \delta h(\underline{\theta}), \quad (5)$$

the incumbent chooses f_* and remains a monopolist;

– otherwise, the incumbent chooses f^m and the entrant enters.

We conclude with how the incumbent's framing strategy depends on primitives of the market, in particular the consumers' susceptibility to framing.

¹⁶Note that the right-hand side of conditions (2), (3), (4), and (5) are equal. However, the left-hand side of (2) is strictly larger than the left-hand side of (3), (4), and (5).

Proposition 3 (Comparative Statics)

- *The incumbent is more likely to remain a monopolist and to use framing to deter entry when intrinsic product differences or consumer heterogeneity are smaller (i.e., lower δ or higher θ).*
- *A weaker susceptibility to primacy effects implies that $\varphi^o(E_{f^m})$ and $\varphi^o(E_{f^*})$ are lower and that the incumbent is more likely to use f_* to deter entry. Otherwise, it has ambiguous effects on $\varphi^o(E_{f^*})$ and the incumbent's use of f^* to deter entry.*

If the differentiation allowed by framing is smaller due to lower δ , the post-entry market is more competitive and less profitable. Thus, entry has to cost less to be a threat. The incumbent also has more to lose and so is more willing to deter entry, even if this requires forgoing some monopoly profit. A higher θ has similar effects, as it curbs the benefits from splitting the market between top and bottom consumers. Note that our results shed light not only on when incumbents use framing to defend their position, but also on how they do so.

Optimal framing depends in more intricate ways on the consumers' susceptibility to primacy effects. Weakening it curbs the frames' ability to create fictitious differentiation—lowering post-entry profits—but also to deter entry. Either way, weaker primacy effects can render entry less likely, as frames are less effective at stifling competition after entry *and* doing so benefits entrants less.

Our results offer some novel insights into advertisement. These complement the view that its function is to provide information about available products to consumers who have fixed tastes. Here, we keep that information fixed and change *how* it is framed, which is an important part of advertisement. The discussed benefits of controlling frames suggests another reason for why firms seek to be presented in prominent positions to consumers (like in web searches or e-commerce stores). The logic of our results is also related to the so-called *pioneering advantage*: Carpenter and Nakamoto (1989) find a gap between the market shares of pioneers and later entrants that cannot be explained by switching costs and seems to arise from the process whereby consumers form their preferences.

3.2 Break the Impasse: Framing in Negotiations

Framing is often regarded as an important negotiation technique (see, e.g., Donohue et al. (2011) and references therein). This is based on the fact that the way a party describes his offer strongly affects how others view it. Framing occurs in every negotiation whether parties are aware of it or not. The party controlling the framing process can define a negotiation to its advantage. Positioning a product advantageously at the outset of every negotiation is viewed as essential for consistently favorable outcomes.

Sometimes re-framing problems helps break an impasse. One framing technique often used involves actively focusing attention on some aspects of a problem and leaving others in the background, thereby shaping the other parties' frame of reference and what they pay attention to. Negotiators usually emphasize what they believe are important and advantageous aspects for them. They may also take others' viewpoints into account so as to offer solutions that reach win-win outcomes.

Despite its importance, modeling framing in negotiations and how it is used has been challenging. We believe our model provides a step forward in tractability. This illustration will offer some insights into how the proposing party may select which aspects are important and advantageous in framing a negotiation.

Two agents, called the proposer P and the receiver R , negotiate over a problem that involves several attributes. Let $L_a = \mathbb{R}$ for all $a \in A$. A specification of these attributes under f defines an item x_f in our model. Let agent j 's payoff from x_f be

$$\sum_{i=1}^N \alpha(i) u_{f(i)}(x_{f(i)}; \bar{x}_{f(i)}^j), \quad \text{where} \quad u_a(x_a; \bar{x}_a^j) = \beta_a - \gamma_a(x_a - \bar{x}_a^j)^2,$$

$\bar{x}_a^j, \beta_a \in \mathbb{R}$, and $\gamma_a > 0$ for all $a \in A$. We interpret $\bar{x}_a^j$ as agent j 's bliss point for a . Quadratic payoff functions are standard and convenient for tractability. We can also interpret $u_a(\cdot; \bar{x}_a^j)$ as a second-order Taylor approximation of a single-peaked function around the bliss point. We assume that $\bar{x}_a^P \neq \bar{x}_a^R$ for all $a \in A$. Importantly, β_a and γ_a can differ across attributes. Note that we can replace each u_a with

$$\hat{u}_a(x_a; \bar{x}_a^j) = \frac{\beta_a}{\sum_{a' \in A} \beta_{a'}} \left[1 - \frac{\gamma_a}{\beta_a} (x_a - \bar{x}_a^j)^2 \right] = \hat{\beta}_a [1 - \hat{\gamma}_a (x_a - \bar{x}_a^j)^2],$$

without changing the agents' preferences. Thus, β_a is directly related to the relative importance of attribute a for the agents and γ_a to their sensitivity to deviations from the bliss point. Finally, we assume that α is strictly decreasing.

The negotiation proceeds as follows. The proposer chooses frame $f \in F$ and an item x_f to maximize her payoff. The receiver accepts x_f if and only if her payoff is at least as large as the reservation utility $\bar{u}^R \in \mathbb{R}$. We assume that there are proposals other than the bliss item $\bar{x}^R$ which the receiver would accept, but she would not accept the proposer's bliss item $\bar{x}^P$ under any frame:

$$\max_{f \in F} \sum_{i=1}^N \alpha(i) \beta_{f(i)} > \bar{u}^R > \max_{f \in F} \sum_{i=1}^N \alpha(i) [\beta_{f(i)} - \gamma_{f(i)} (\bar{x}_{f(i)}^P - \bar{x}_{f(i)}^R)^2]. \quad (6)$$

For now, let the proposer's reservation utility be $\bar{u}^P = -\infty$.

Before we solve the model, a few remarks are in order. First, we assume that the agents disagree only in their bliss points, but have otherwise the same preferences (in terms of β_a or γ_a). This is because we want to focus on how the proposer exploits differences across attributes to frame the problem by emphasizing some over others,

which is the core of our paper. This aspect would be obfuscated by differences between the agents' preferences. Second, we can interpret our model as the first period of a repeated bargaining with alternating offers. In this case, $\bar{u}^R$ is the payoff the first-period receiver expects from rejecting, calculated by backward induction. This interpretation also justifies assuming that preferences are symmetric between the agents (except their bliss points): The only difference between them is that one is selected first to make an offer. Another interpretation of why the proposer's payoff depends on framing is that he is a third party who acts on behalf of a client and hence internalizes how she will perceive the outcome based on its presentation.

We start from the proposer's offer given any frame. By standard steps (provided in Appendix B), the optimal level of each attribute is

$$x_{f(i)}(\lambda) = \frac{1}{1+\lambda} \bar{x}_{f(i)}^P + \frac{\lambda}{1+\lambda} \bar{x}_{f(i)}^R, \quad (7)$$

where λ is the Lagrange multiplier of the receiver's participation constraint. Assumption (6) implies $\lambda > 0$. The proposer offers a compromise between bliss points for every attribute. The easier it is to convince the receiver to accept—as captured by a lower λ —the more this compromise caters to the proposer's bliss point. Indeed, $x_f(\lambda)$ also has to satisfy

$$\sum_{i=1}^N \alpha(i) [\beta_{f(i)} - \gamma_{f(i)} (x_{f(i)}(\lambda) - \bar{x}_{f(i)}^R)^2] = \bar{u}^R, \quad (8)$$

where the left-hand side increases in λ . This is where the choice of f matters, as it can help satisfy (8) and thus lower λ .

To examine the proposer's optimal framing, we proceed as follows (see Appendix B for details). Using (7) and (8), we can substitute $x_f(\lambda)$ and λ in the proposer's payoff function and derive a reduced objective that depends only on f :

$$U^P(f) = B(f) - \left[\sqrt{\Gamma(f)} - \sqrt{B(f) - \bar{u}^R} \right]^2,$$

where

$$\Gamma(f) = \sum_{i=1}^N \alpha(i) \gamma_{f(i)} (\bar{x}_{f(i)}^P - \bar{x}_{f(i)}^R)^2 \quad \text{and} \quad B(f) = \sum_{i=1}^N \alpha(i) \beta_{f(i)}.$$

Crucially, $U^P(f)$ increases as $\Gamma(f)$ decreases and as $B(f)$ increases. Thus, the proposer faces a trade-off between emphasizing important attributes (high β_a) and attributes that involve little disagreement (low $\gamma_a |\bar{x}_a^R - \bar{x}_a^P|$). We may then conclude the following.

Proposition 4 (Optimal Framing in Negotiations) *If attributes a and a' satisfy $\beta_a > \beta_{a'}$ and $\gamma_a |\bar{x}_a^R - \bar{x}_a^P| < \gamma_{a'} |\bar{x}_{a'}^R - \bar{x}_{a'}^P|$, then every optimal frame f presents a before a' (i.e., $f^{-1}(a) < f^{-1}(a')$).*

The takeaway here is that the proposer should present earlier attributes which are important *and* involve little disagreement and should present later attributes which are unimportant *and* involve significant disagreement. Things are more subtle for important but highly conflictual attributes, which should be presented in middle positions. The point, however, is that it is possible that the proposer strategically de-emphasizes some attribute of the negotiation, despite its being very important, so as to weaken the impact of the disagreement with the receiver on that attribute. For this to be the case, the gain through $\Gamma(f)$ has to dominate on the loss through $B(f)$.

In our model, framing can also emerge as a tool to break an impasse. Consider again the interpretation of agent P as acting on behalf of a client. Suppose the client's reservation utility now satisfies $\bar{u}^P > -\infty$. By assumption (6) the proposer can always find a deal that the receiver accepts. Yet, this deal may be unacceptable for the proposer's client, unless framed in the right way. Consider the generic case in which framing matters:

$$U^P(\underline{f}^P) = \min_{f \in F} U^P(f) < \max_{f \in F} U^P(f) = U^P(\bar{f}^P).$$

Corollary 1 (Breaking the Impasse) *If the proposer's reservation utility satisfies $U^P(\underline{f}^P) < \bar{u}^P < U^P(\bar{f}^P)$, then using frame $\underline{f}^P$ leads to an impasse, while using $\bar{f}^P$ leads to an agreement.*

Thus, our model captures the common intuition that successful negotiators are those who also have the skill of finding the right way to frame things. Moreover, by being explicit about how framing works, the model can offer insights into useful strategies to break an impasse.

Finally, does more susceptibility to framing effects benefit proposers in negotiations? Consider two proposer-receiver pairs that differ only in α , denoted by α^1 for the first pair and α^2 for the second. Suppose α^1 exhibits more susceptibility to primacy effects than α^2 (Definition 3).

Corollary 2 *Fix $\bar{u}^R$ and suppose more important attributes also involve less disagreement (i.e., $\beta_a > \beta_{a'}$ if and only if $\gamma_a|\bar{x}_a^R - \bar{x}_a^P| < \gamma_{a'}|\bar{x}_{a'}^R - \bar{x}_{a'}^P|$ for all $a, a' \in A$). Then, more susceptibility to framing always benefits the proposer.*

Without this inverse relation between importance and disagreement across attributes, more susceptibility to framing can harm the proposer as shown next.

Example 2 *There are two attributes, a and a' , which satisfy $\beta_a = 0$, $\beta_{a'} = 1$, $\gamma_a(\bar{x}_a^R - \bar{x}_a^P)^2 = 1$, and $\gamma_{a'}(\bar{x}_{a'}^R - \bar{x}_{a'}^P)^2 = 1 + z$ where $z > 0$. Suppose $\bar{u}^R = 0$ (this is just a normalization). Slightly abusing notation, let $\alpha(1) = \alpha \in (\frac{1}{2}, 1)$ and $\alpha(2) = 1 - \alpha$.*

Thus, more susceptibility to primacy effects here means higher α . Let f present first a and f' present first a' . The proposer's payoff is either

$$U^P(f) = 1 - \alpha - \left[\sqrt{1 + (1 - \alpha)z} - \sqrt{1 - \alpha} \right]^2$$

or

$$U^P(f') = \alpha - \left[\sqrt{1 + \alpha z} - \sqrt{\alpha} \right]^2.$$

We will show that there exist z and $\bar{\alpha}$ such that each payoff is strictly decreasing in α for $\alpha > \bar{\alpha}$. In this case, the proposer is worse off when he and the receiver are more susceptible to framing. By simple steps, $\frac{\partial U^P(f)}{\partial \alpha} < 0$ and $\frac{\partial U^P(f')}{\partial \alpha} < 0$ if and only if

$$\frac{\frac{1}{\alpha} + z}{\sqrt{\frac{1}{\alpha} + z} - 1} < z < \frac{\frac{1}{1-\alpha} + z}{\sqrt{\frac{1}{1-\alpha} + z} - 1}.$$

Note that the left term is decreasing in α , the right term is increasing in α , and

$$\lim_{\alpha \rightarrow 1} \frac{\frac{1}{1-\alpha} + z}{\sqrt{\frac{1}{1-\alpha} + z} - 1} = +\infty.$$

Evaluated at $\alpha = 1$, the first inequality holds if

$$\frac{1}{z} + 2 < \sqrt{1 + z},$$

hence for sufficiently large but finite z . Given this z , there exists $\bar{\alpha}$ such that both derivatives are strictly negative for $\alpha > \bar{\alpha}$.

4 General Menus and Frame-Driven List Effects

Sometimes decision-makers face items whose attributes are presented with different orders—for example, on the shelves of grocery stores or when sellers independently choose how to best frame their products.¹⁷ This leads to general menus where different frames are present simultaneously. How do frames affect choice in such cases?

The question is non-trivial, but our model can help organize the discussion and provide some answers. The idea is to let the data speak. We can identify our model using f -menus only (see Section 7). We can then use it to predict choices under different hypotheses on how the decision-maker responds to general menus—we present a few shortly. We can test these hypotheses against the data and select which we judge to be the best fit. Formally, let h be an hypothesis and $\mathbf{M}$ a collection of menus. The decision-maker's choices from these menus form the actual dataset $c(\mathbf{M})$. Using our model with (α, u) calibrated to this decision-maker, we can calculate his utility from

¹⁷See the extension of the entry game of Section 3.1 in Appendix D.

each item and his choices under h , which form the predicted dataset $\hat{c}(\mathbf{M}; h)$. We can then compare $c(\mathbf{M})$ and $\hat{c}(\mathbf{M}; h)$ in any standard way. For example, h can be falsified if $\hat{c}(\mathbf{M}; h) \neq c(\mathbf{M})$, or, more realistically, if $d(\hat{c}(\mathbf{M}; h), c(\mathbf{M})) > \tau$ for some distance function d and tolerance threshold $\tau > 0$. Among multiple hypotheses, we may select the one that minimizes $d(\hat{c}(\mathbf{M}; h), c(\mathbf{M}))$.¹⁸

To illustrate this approach, we consider three hypotheses.

Own-frame hypothesis (h_1): Suppose that, when facing a general menu, Bob chooses as if he evaluates each item following its own order of attributes. That is, for every M , the predicted choice is

$$\hat{c}(M; h_1) = \arg \max_{x_f \in M} \sum_{i=1}^N \alpha(i) u_{f(i)}(x_{f(i)}).$$

Note that for f -menus this boils down to the model in Definition 1.

Single-reframe hypothesis (h_2): A second possibility is that, when facing a general menu, Bob chooses as if he reframes all its items using the same f . Formally, let φ be a function that maps every M to some $f \in F$. Let $M_{\varphi(M)}$ be the $\varphi(M)$ -menu that contains all items in the original M presented with frame $\varphi(M)$. In this case, Bob's choice from M should coincide with that predicted by the AF model from $M_{\varphi(M)}$:

$$\hat{c}(M; h_2) = \arg \max_{x_{\varphi(M)} \in M_{\varphi(M)}} \sum_{i=1}^N \alpha(i) u_{\varphi(M)(i)}(x_{\varphi(M)(i)}).$$

This hypothesis is very flexible. Bob may adopt different frames for different menus. Alternatively, he may adopt the same frame for all general menus. That is, h_2 covers the possibility that general menus cause framing effects to disappear. This may happen if the effort to organize and make sense of the various items causes the different emphasis put on attributes to wane.

Anchor-frame hypothesis (h_3): A special case of h_2 is that Bob chooses as if he uses the frame of one item in the menu as an anchor (see Krosnick and Alwin (1987) for consistent evidence). For example, the f of the first listed item, the last listed item, or the most frequent f could cue Bob to use f to compare all items in the menu. This anchoring may introduce links between attribute-order framing effects and list framing effects, which we can formalize and test with our model (see Section 7.4 for further discussion). In fact, list framing effects are sometimes viewed as the outcome of particular ways of processing the attributes of the listed items (see, e.g., Rubinstein and Salant (2006) and references therein). Knowing what determines framing anchors

¹⁸Standard methods can be used to define d . For example, one can use the *swap index* in Apesteguia and Ballester (2015). Given h , the model generates a preference relation over the items in each menu. One can measure the distance between the actual and predicted choices by the number of swaps in the preference relation needed to make the actual choice preferred to the predicted choice. One can then aggregate this measure across menus and choose the hypothesis that minimizes it.

can be valuable for sellers. For instance, if the anchor is the item listed first, sellers have an incentive to try to put their product in that position and frame it in the most favorable way. This mechanism may contribute to explaining why firms pay a premium to be listed first, say, by search engines.

The role of items as cues for how to reframe other items in general menus renders choice and preferences menu dependent. This can lead to failures of standard axioms, such as the independence of irrelevant alternatives (IIA). This is intuitive. Consider the menus $\{z_{f''}, x_f, y_{f'}\}$ and $\{x_f, y_{f'}\}$. Suppose for general menus Bob uses the item listed first to reframe the others. Then, it is possible that for him x dominates y and z under frame f'' , but y dominates x under f . Thus, through the lens of our model we can understand violations of IIA as resulting from attribute-framing effects. Of course, IIA may fail for many other reasons. The example also illustrates that, due to the same mechanism, larger menus may increase the likelihood that some item is chosen by cuing a frame that favors it. This violates other regularity axioms that characterize standard choice models.

5 Self-Serving Rationalization via Framing

A large body of evidence shows that people often engage in motivated reasoning, rationalization, self-deception, self-justification, and reduction of cognitive dissonance by strategically presenting to themselves situations and decisions in the most favorable *perspective* (Bénabou and Tirole (2016)). One way is to emphasize some of their aspects over others. Such habits can be conscious or automatic, affective (to feel better) or functional (to achieve goals), and depend on emotions. For instance, rationalization can serve to avoid disappointment, guilt, or regret. Cognitive dissonance may result in a strategy called minimization, namely, reducing the importance of elements of dissonance (Lindsey-Mullikin (2003), Beasley and Joslyn (2001)). Self-serving justification aims to make questionable behaviors appear less unethical. It can occur *ex ante*—to paint violations as excusable in the eye of one’s moral self—or *ex post*—to lessen the experienced threat for one’s moral self (Shalvi et al. (2015)).

A key question is how to capture self-serving perspective manipulations in a disciplined manner. We argue that our framework can provide a way. Our premise is that, when making decisions, some individuals are susceptible to frames set by someone else, like salespeople or experimenters. In a similar logic, the choosing self of such individuals may also be influenced by frames set by their rationalizing self. This dual-self view is consistent with leading models of motivated reasoning (Bénabou and Tirole (2016)). Imagine we can describe the choosing self with our AF model. The rationalizing self can set f to manipulate the perspective under which the choosing self makes decisions, emphasizing certain aspects with their presentation order. Introspection suggests that when facing a decision—especially new and complex ones—we

first try to organize its aspects, thereby forming a specific presentation order. This order may depend on our motivations and affect our choice.

We distinguish between two scenarios: *ex-ante* and *ex-post* self-serving framing. In the first, f is set *before* the choosing self makes a decision; in the second, f is set *after* a decision. The rationalizing self may want to maximize or minimize the evaluation of an item depending on her motivation in the situation of the moment (Bénabou and Tirole (2016)). For every item $x \in X$, define $\bar{f}_x$ and $\underline{f}_x$ as

$$\bar{f}_x \in \arg \max_{f \in F} \sum_{i=1}^N \alpha(i) u_{f(i)}(x_{f(i)}) \quad \text{and} \quad \underline{f}_x \in \arg \min_{f \in F} \sum_{i=1}^N \alpha(i) u_{f(i)}(x_{f(i)}).$$

Ex ante, Ann may adopt $\bar{f}_x$ to motivate herself to do x —say, exercise—or to justify doing x —say, violating some norm; she may adopt $\underline{f}_x$ if for example she is going to bargain on the price of x . Ex post, Bob may adopt $\bar{f}_x$ if he bought x and wants to justify the expenditure to himself, or $\underline{f}_x$ if he could not get x and wants to lessen his feeling of regret or disappointment.

Importantly, our model imposes trade-offs and constraints on the rationalizing self. An individual cannot deceive herself without limits. Emphasizing some aspects requires de-emphasizing others: It is not possible to simply increase or decrease the weight on all attributes. Also, de-emphasizing has bounded effects: It is not possible to ‘forget’ bad aspects since $\alpha > 0$. Finally, our model assigns a precise meaning to frames, namely, the order in which the rationalizing self describes the aspects of an item.

The possibility that a decision-maker may frame items in a self-serving manner is consistent with the possibility that he is influenced by frames set by others. Bob’s choosing self may buy x under the influence of some f in the store, which can differ from the frame $\bar{f}_x$ her rationalizing self sets once at home (recall Proposition 2). This relates to and offers a formalization of the distinction between *decision utility* and *experienced utility* (Kahneman et al. (1997); Kahneman et al. (1999)). The first is the utility that drives decisions in the heat of the moment—for instance, in the store under the f crafted by a skillful salesperson. The second is the hedonic utility experienced in the cold state of the rationalizing moment—for instance, at home after calmly thinking about the bought item. For the above reasons, the experienced utility may be determined by $\bar{f}_x$. Our model provides a tool to calculate both decision and experienced utility knowing (α, u) .

5.1 A Framing Perspective on the Endowment Effect

To illustrate the logic of self-serving framing, we apply it in the context of a well-known phenomenon: the endowment effect (Thaler (1980)). This phenomenon relates the willingness to pay (*WTP*) for acquiring an object and the willingness to accept

(WTA) for giving up possession of the same object. Standard choice theory predicts that $WTA = WTP$. Yet, evidence shows that subjects often exhibit $WTA > WTP$ (Kahneman et al. (1991)). Here we sketch one angle to think about this phenomenon, which may complement the leading explanation based on expectation-based reference dependence (Kőszegi and Rabin (2006)).

Imagine the following situation. The choice items of interest have multiple attributes. Ann can be described by an AF model (α, u) , where α is not constant. Her value of having no item is zero. We offer Ann the possibility of acquiring x under frame f , which determines her decision utility for x_f . Assuming quasi-linearity in money, we can define Ann's willingness to pay for x as

$$WTP(x) = \sum_{i=1}^N \alpha(i) u_{f(i)}(x_{f(i)}).$$

If when asked to give up x Ann evaluates it under f , we observe $WTA(x) = WTP(x)$. But she may frame x differently at this stage, which leads to $WTP(x) \neq WTA(x)$.

Even so, why should reframing of acquired items systematically lead to $WTA \geq WTP$? We already mentioned reasons for self-serving framing suggested by cognitive science and psychology that may lead Ann to try to avoid negative feelings ex post. In addition, according to Beggan (1992) the desire to see oneself favorably may induce people to overvalue objects associated with the self, namely, owned objects. Thus, Ann may tend to use $\bar{f}_x$ when considering giving up x . This implies

$$WTA(x) = \sum_{i=1}^N \alpha(i) u_{\bar{f}_x(i)}(x_{\bar{f}_x(i)}) \geq WTP(x)$$

for every initial f , with strict inequality for some f . The avoidance of negative feelings seems to be a potential cause of the endowment effect (Zhang and Fishbach (2005)). Other evidence shows that the longer an individual owns the item, the bigger is the WTA - WTP gap (Strahilevitz and Loewenstein (1998)). Presumably, the longer ownership allows Ann to figure out the best frame for x .

It is worth noting that even if sellers can select f to maximize profits, a gap between WTP and WTA may still arise. Section 3.1 showed that in competitive settings it may not be optimal for sellers to select $f = \bar{f}_x$.

One final insight of our model is that experience should eliminate the endowment effect. If Ann remembers how she reframed x after experiencing it a few times, the WTP - WTA gap should disappear because she cannot be manipulated again by changing the presentation of x . Consistent with this, evidence shows that market experience seems to eliminate the endowment effect (List (2003)) and that the effects of attribute-order framing disappear for subjects who had experience with the choice items (Levin and Gaeth (1988), Kumar and Gaeth (1991)). This does not reduce the importance of studying framing effects, as many and consequential choices in life happen infrequently and with little to no feedback.

6 Welfare Analysis

How do we measure the welfare of agents who are sensitive to attribute-framing? We are certainly not the first to consider welfare analysis in the presence of framing effects (Bernheim and Rangel (2009); Rubinstein and Salant (2011)). Our framework—by being explicit about what frames are and how they work—provides several ways to approach this thorny question. Each has some merits and flaws. Since these are often well-known, we limit ourselves to presenting what these ways are.

6.1 Choice-Based Welfare

One way is to apply the choice-based approach proposed by Bernheim and Rangel (2009). They define a generalized choice situation as a constraint set paired with an ancillary condition. In our setting, the constraint set corresponds to a set of choice items $D \subset X$; the ancillary condition is the collection $\mathbf{f}$ of all frames f with which the items are presented. Following Bernheim and Rangel (2009), we say that it is possible to strictly improve upon $x \in \hat{D}$ if there exists $x' \in \hat{D}$ such that, for all $(D, \mathbf{f})$ which satisfy $x, x' \in D$, the decision-maker never chooses x . We can then say that x is a weak individual welfare optimum when a strict improvement is not possible.¹⁹ Bernheim and Rangel (2009) show that this welfare criterion is the most discerning criterion that never overrules choice.

Bernheim and Rangel’s criterion has specific implications for our model (see also their Theorem 3). Suppose that for every $\mathbf{f}$ Ann compares $x, x' \in D$ using a common frame $f \in \mathbf{f}$ (like in the case of f -menus or hypothesis h_2 above). Then, she never chooses x when x' is available if and only if for all $f \in F$

$$\sum_{i=1}^N \alpha(i) u_{f(i)}(x'_{f(i)}) > \sum_{i=1}^N \alpha(i) u_{f(i)}(x_{f(i)}).$$

Therefore, x is (weakly) welfare optimal if for all x' there exists *some* f such that

$$\sum_{i=1}^N \alpha(i) u_{f(i)}(x_{f(i)}) \geq \sum_{i=1}^N \alpha(i) u_{f(i)}(x'_{f(i)}).$$

In other words, x is welfare optimal if for every x' Ann prefers x_f to x'_f for some appropriately chosen f based on x' . Note that this f need not be the one that maximizes the utility from x (i.e., $\bar{f}_x$).

¹⁹We refer the reader to Bernheim and Rangel (2009) for the definitions of weak improvement and strict welfare optimum.

6.2 Experienced-Utility Welfare

Another typical approach of behavioral economics to welfare analysis involves the distinction between decision utility and experienced (or true) utility, where the latter should be used to measure well-being (Kahneman et al. (1997); Kahneman et al. (1999); Bernheim and Rangel (2009)). As noted in Section 5, our model provides a way and a rationale for defining experienced utility for decision-makers affected by attribute-framing. Suppose that, for the reasons discussed above, Ann re-frames each owned item x according to $\bar{f}_x$ after acquiring it. Then, her experienced utility is

$$\bar{U}(x) = \sum_{i=1}^N \alpha(i) u_{\bar{f}_x}(x).$$

This welfare measure is based on the idea that the important source of well-being for Ann is the utility she experiences once owning x , not the utility she used to choose x .

It is worth noting a couple of properties of the experienced utility $\bar{U}$. First, it defines a ranking over items that is frame-independent. The original dependence is removed by considering the best frame $\bar{f}_x$ for each x . Second, although the underlying AF model is additively separable across attributes for every f , the induced $\bar{U}$ need not be separable as the whole x determines the best frame $\bar{f}_x$. Thus, there can be interdependences between attributes that are exclusively driven by self-serving framing considerations. This structural difference between decision utility and experienced utility may suggest a way to test this approach to welfare analysis.

6.3 Frame-Free Welfare

The last alternative we discuss uses the properties of our model to entirely remove frames from welfare measures. It is based on the premise that frames should not matter for decision-makers' choices and hence, a fortiori, for a planner's welfare analysis.

The idea is to exploit our model's identification of the tastes for each attribute. We can define the frame-free welfare generated by an item as the sum of the utilities of its attributes. That is, given a decision-maker described by (α, u) , this measure is

$$U^o(x) = \sum_{a \in A} u_a(x_a), \quad x \in X.$$

Of course, this way of removing framing effects involves some degree of paternalism. For another interpretation of this approach, note that U^o is equivalent, in terms of ranking, to taking the average across all frames of the total utility of an item. Indeed, since each attribute can be presented in each position, we have

$$\frac{\sum_{f \in F} \sum_{i=1}^N \alpha(i) u_{f(i)}(x_{f(i)})}{|F|} = \left(\frac{\sum_{i=1}^N \alpha(i)}{|A|} \right) \sum_{a \in A} u_a(x_a).$$

7 Axiomatizations and Extensions

7.1 Behavioral Characterization of AF Models

To characterize our model, we enrich the choice domain by allowing for (simple) *lotteries* over items. This provides enough structure for the task. The idea is that each item involves some risk: Its attributes are presented in a specific order, but (the consequences of) their levels can be uncertain at the time of choice. For instance, when choosing between a new sedan or a used SUV, a buyer may not know which will better serve his needs. We will rely only on lotteries whose support involves items all framed in the same way. Such lotteries belong to $\Delta(X_f)$ and are denoted by p_f , q_f , and r_f . To simplify notation, we denote binary lotteries that yield x_f with probability p and y_f with probability $1 - p$ by

$$(x_f, y_f; p)$$

An f -menu M_f is a subset of $\Delta(X_f)$. We assume that $|A| = N \geq 3$; one can allow for $N = 2$ at the cost of stronger separability axioms.

As the primitive data, we assume that we can observe the decision-maker's choices from all f -menus. This choice set is denoted by $c(M_f)$ and has the usual interpretation. Note that the frames of each item in a menu are part of the dataset.

Our basic assumption is that as long as all items in a menu are framed in the same way, the decision-maker satisfies standard rationality assumptions. That is, we can describe her choices as the maximization of some utility function which can depend at most on the frame. We go one step further and assume that she is an expected-utility maximizer. We present these properties directly as an assumption because they follow from standard axioms.

Assumption 1 (f -EU Representation) *For every $f \in F$, there exists a function $w_f : X_f \rightarrow \mathbb{R}$ such that for every M_f*

$$c(M_f) = \arg \max_{q_f \in M_f} v_f(q_f), \quad \text{where} \quad v_f(q_f) = \sum_{x_f \in \text{supp } q_f} w_f(x_f) q_f(x_f).$$

To characterize our AF model in Definition 1, we then need to find properties of c that correspond to each w_f taking the form

$$w_f(x_f) = \sum_{i=1}^N \alpha(i) u_{f(i)}(x_{f(i)})$$

for some $\alpha : \{1, \dots, N\} \rightarrow \mathbb{R}_{++}$ and $u_a : L_a \rightarrow \mathbb{R}$ for all $a \in A$. We organize these properties in four axioms.

Axiom 1 is a simple non-triviality condition: For no attribute the decision-maker is indifferent between all its possible levels. To state this and later properties formally, let $x_{f(-i)}$ be the description of item x_f excluding position i .²⁰

Axiom 1 (Non-triviality) *For every $a \in A$, there exists $x_a, y_a \in L_a$ such that, if $f(1) = a$ and $x_{f(-1)} = y_{f(-1)}$, then $c(x_f, y_f) = \{x_f\}$.*

Axiom 2 is inspired by standard separability axioms as in Debreu (1960): How the decision maker trades off the levels of any two attributes does not depend on the levels of other attributes. We relax this in Section 7.2.

Axiom 2 (Separability) *Fix $f \in F$ and any $j, k \in \{1, \dots, N\}$. For all x_f, x'_f, y_f, y'_f that satisfy $x_{f(i)} = y_{f(i)}$ and $x'_{f(i)} = y'_{f(i)}$ for $i = j, k$ and $x_{f(i)} = x'_{f(i)}$ and $y_{f(i)} = y'_{f(i)}$ for all $i \neq j, k$, we have*

$$c(x_f, x'_f) = c(y_f, y'_f).$$

Axiom 3 captures the property that the decision-maker's tastes for each attribute do not depend on the position in which the attributes are presented.

Axiom 3 (Taste Framing Independence) *For every $i, j = 1, 2, \dots, N$, $a \in A$, and $f, f' \in F$ such that $f(i) = a$ and $f'(j) = a$, the following holds: If $p_{f(i)}, q_{f(i)} \in \Delta(L_{f(i)})$, $\hat{p}_{f'(j)} = p_{f(i)}$, $\hat{q}_{f'(j)} = q_{f(i)}$, $x_{f(-i)} = y_{f(-i)}$, and $\hat{x}_{f'(-j)} = \hat{y}_{f'(-j)}$, then*

$$c((p_{f(i)}, x_{f(-i)}), (q_{f(i)}, y_{f(-i)})) = c((\hat{p}_{f'(j)}, \hat{x}_{f'(-j)}), (\hat{q}_{f'(j)}, \hat{y}_{f'(-j)})).$$

Axiom 4 exploits the cardinality of expected utility to identify how the decision-maker weights attributes based on their presentation. The idea is to consider three items that differ in only two attributes. One item is better in both attributes, one is worse in both attributes, and one is in between. Consider a lottery between the best and worst item that the decision-maker deems indifferent to the in-between item. The requirement is that this lottery depends at most on the presentation positions of the two different attributes. To formalize this, for every $a \in A$ we say that x_a is strictly preferred to y_a —written $x_a \succ y_a$ —if $\{x_f\} = c(x_f, y_f)$ whenever $f(1) = a$, $x_{f(1)} = x_a$, $y_{f(1)} = y_a$, and $x_{f(-1)} = y_{f(-1)}$. Now fix any $i = 2, \dots, N$. Given any x_f, y_f , and z_f such that $x_{f(1)} \succ z_{f(1)} = y_{f(1)}$, $x_{f(i)} = z_{f(i)} \succ y_{f(i)}$, and $x_{f(j)} = y_{f(j)} = z_{f(j)}$ for all $j \neq 1, i$, define p_{xyz_f} as the probability that satisfies

$$\{(x_f, y_f; p_{xyz_f}), z_f\} = c((x_f, y_f; p_{xyz_f}), z_f).$$

²⁰We write $c(p_f, \dots, q_f)$ for $c(\{p_f, \dots, q_f\})$ to simplify notation.

Axiom 4 (Position Dependence) Let $f, f' \in F$ satisfy $f(1) = f'(1)$ and $f(i) = f'(j)$. Let p_{xyz_f} and $p_{xyz_{f'}}$ be as in Definition 4. If $x_{f(i)} = x_{f'(j)}$, $y_{f(i)} = y_{f'(j)}$, and $z_{f(i)} = z_{f'(j)}$, then

$$\frac{p_{xyz_f}}{1 - p_{xyz_f}} = g(i, j) \frac{p_{xyz_{f'}}}{1 - p_{xyz_{f'}}}$$

where $g(i, j) \in \mathbb{R}$ can depend only on i and j .

This axiom certainly imposes significant structure on c . However, note that a standard model without framing effects would require $g(i, j) = 1$ for all i, j .

Theorem 1 (AF Representation) Under Assumption 1, Axioms 1–4 hold if and only if c has an AF representation: There exists $\alpha : \{1, \dots, N\} \rightarrow \mathbb{R}_{++}$ and non-constant $u_a : L_a \rightarrow \mathbb{R}$ for every $a \in A$ that satisfy, for all $f \in F$ and M_f ,

$$c(M_f) = \arg \max_{x_f \in M_f} \sum_{i=1}^N \alpha(i) u_{f(i)}(x_{f(i)}).$$

The proofs of this section appear in Appendix C.

7.1.1 Behavioral Characterization of Primacy and Recency Effects

We now characterize the attribute-framing effects and the comparison between decision-makers in Definitions 2 and 3. To this end, we define primacy and recency effects in terms of observable data using specific lotteries.

Definition 4 (Calibration Lottery) Fix any $i = 1, \dots, N - 1$. Let x_f, y_f , and z_f be such that $x_{f(i)} = z_{f(i)} \succ y_{f(i)}$, $x_{f(i+1)} \succ z_{f(i+1)} = y_{f(i+1)}$, and $x_{f(j)} = z_{f(j)} = y_{f(j)}$ for $j \neq i, i + 1$. Obtain $x_{f'}, y_{f'}$, and $z_{f'}$ by swapping only the attributes in positions i and $i + 1$. Given this, define $p_{xyz_f}^i$ and $p_{xyz_{f'}}^i$ as the probabilities that satisfy

$$\{(x_f, y_f; p_{xyz_f}^i), z_f\} = c((x_f, y_f; p_{xyz_f}^i), z_f)$$

and

$$\{(x_{f'}, y_{f'}; p_{xyz_{f'}}^i), z_{f'}\} = c((x_{f'}, y_{f'}; p_{xyz_{f'}}^i), z_{f'}).$$

Intuitively, z_f dominates y_f in an earlier attribute, while $z_{f'}$ dominates $y_{f'}$ in a later attribute. Thus, $p_{xyz_f}^i$ should be higher (lower) than $p_{xyz_{f'}}^i$ if the decision-maker is affected by primacy (recency) effects.

Definition 5 (Revealed Primacy/Recency Effects) For $i = 1, \dots, N - 1$, define $p_{xyz_f}^i$ and $p_{xyz_{f'}}^i$ as above. Then, c exhibits primacy (recency) effects if

$$p_{xyz_f}^i > (<) p_{xyz_{f'}}^i \quad \text{for all } i = 1, \dots, N - 1.$$

Proposition 5 *Let (α, u) be an AF representation of c . Then, c exhibits a primacy (recency) effect if and only if α is strictly decreasing (increasing).*

We now turn to comparing individuals' susceptibility to attribute framing. For this to be meaningful, we should only compare individuals who have the same tastes for attributes.

Definition 6 (Revealed Same Tastes) *Decision-maker 1 and 2 exhibit the same tastes for attributes if c^1 and c^2 have the following property. For every $a \in A$, $f \in F$ that satisfies $f(1) = a$, and $p_{f(1)}, q_{f(1)} \in \Delta(L_{f(1)})$,*

$$c^1((p_{f(1)}, x_{f(-1)}), (q_{f(1)}, y_{f(-1)})) = c^2((p_{f(1)}, x_{f(-1)}), (q_{f(1)}, y_{f(-1)})).$$

This explains Definition 3 because if c^1 and c^2 have this property, u_a^1 and u_a^2 represent the same vN-M preference over $\Delta(L_a)$, hence $u_a^1 = \gamma_a u_a^2 + \zeta_a$ for $\gamma_a > 0$ and $\zeta_a \in \mathbb{R}$.

Definition 7 (Revealed Comparative Primacy) *Suppose that decision-makers 1 and 2 exhibit the same tastes for attributes. Decision-maker 1 is more susceptible to primacy effect than decision-maker 2 is if, for all $i = 1, \dots, N - 1$,*

$$\{(x_f, y_f; p_{xy_f}^i), z_f\} = c^2((x_f, y_f; p_{xy_f}^i), z_f) \Rightarrow z_f \in c^1((x_f, y_f; p_{xy_f}^i), z_f)$$

and

$$\{(x_{f'}, y_{f'}; p_{xy_{f'}}^i), z_{f'}\} = c^2((x_{f'}, y_{f'}; p_{xy_{f'}}^i), z_{f'})$$

$$\Rightarrow (x_{f'}, y_{f'}; p_{xy_{f'}}^i) \in c^1((x_{f'}, y_{f'}; p_{xy_{f'}}^i), z_{f'}).$$

Definition 8 (Revealed Comparative Recency) *Suppose that decision-makers 1 and 2 exhibit the same tastes for attributes. Decision maker 1 is more susceptible to recency effect than decision-maker 2 is if, for all $i = 1, \dots, N - 1$,*

$$\{(x_f, y_f; p_{xyz_f}^i), z_f\} = c^2((x_f, y_f; p_{xyz_f}^i), z_f) \Rightarrow (x_f, y_f; p_{xyz_f}^i) \in c^1((x_f, y_f; p_{xyz_f}^i), z_f)$$

and

$$\{(x_{f'}, y_{f'}; p_{xyz_{f'}}^i), z_{f'}\} = c^2((x_{f'}, y_{f'}; p_{xyz_{f'}}^i), z_{f'}) \Rightarrow z_{f'} \in c^1((x_{f'}, y_{f'}; p_{xyz_{f'}}^i), z_{f'}).$$

The next result maps these behavioral comparisons into properties of our AF representation, thereby providing the foundations for Definition 3.

Proposition 6 *Suppose decision-makers 1 and 2 can be represented by AF models (α^1, u^1) and (α^2, u^2) and exhibit the same tastes for attributes. Decision-maker 1 is more susceptible to primacy (recency) effect than decision-maker 2 is if and only if*

$$\frac{\alpha^1(i)}{\alpha^1(i+1)} \geq (\leq) \frac{\alpha^2(i)}{\alpha^2(i+1)}, \quad i = 1, \dots, N - 1.$$

7.2 Framing without Separability

In this section we propose a way to relax the additive structure of AF models. To this end, it helps to transition to a random-choice framework. This also allows us to develop a model that can be more easily applied to empirical analysis, which often examines framing in terms of how it affects the probability of choosing an item. In the vast literature on random choice, several papers include observable attributes of choice items (Lancaster (1966), McFadden (1973), Gorman (1980), Allen and Rehbeck (2016)). We also include their framing as part of the dataset to study its effects.²¹

Our primitive data is therefore as follows. For every finite $M_f \subset \Delta(X_f)$, we assume to observe the probability that the decision-maker chooses each $q_f \in M_f$, denoted by

$$\pi(q_f, M_f).$$

This has the usual interpretation of the random-choice literature.²² The item frames are again assumed to be part of the dataset.

We continue to assume that as long as all items in a menu are framed in the same way, we can describe behavior using a standard model. We start from a canonical Luce representation and later consider more general models of random choice. We again introduce this representation directly as an assumption, because it follows from well-known axioms.

Assumption 2 (*f*-EU Luce Representation) *For every $f \in F$, there exists a function $w_f : X_f \rightarrow \mathbb{R}$ such that for every finite $M_f \subset \Delta(X_f)$*

$$\pi(q_f, M_f) = \frac{e^{v_f(q_f)}}{\sum_{q'_f \in M_f} e^{v_f(q'_f)}}, \quad \text{where} \quad v_f(q_f) = \sum_{x_f \in \text{supp } q_f} w_f(x_f) q_f(x_f). \quad (9)$$

The basic premise of this paper is that people often encounter attributes of items in an exogenous order and this may affect their choices. One way to keep this premise while relaxing additivity is to allow the weight Ann assigns to the utility of an attribute to depend on its presentation position as well as the attributes that come before it. That is, Ann may aggregate the utilities across attributes as follows:

$$w_f(x_f) = \sum_{i=1}^N u_{f(i)}(x_{f(i)}) Q(i, x_{f(i-1)}, \dots, x_{f(1)}).$$

²¹Gul et al. (2014) propose a related, but different, approach where the decision-maker *subjectively* frames multi-attribute items. Their elegant analysis identifies how she treats items as more or less substitute based on the subjective similarity of attributes. This approach is silent about the role of objective and exogenous frames. It seems possible that exogenous and subjective frames interact, opening an interesting connection between our and their work. For a study of how ordering of *alternatives* might affect choice in the Luce model, see Tserenjigmid (2021).

²²See, e.g., Luce (1959), Block and Marschak (1960), Marschak (1974), Gul and Pesendorfer (2006), Manzini and Mariotti (2014), and Apesteguia and Ballester (2018).

On practical grounds, it is valuable to impose more structure on the dependence of Q on earlier attributes. We therefore introduce and characterize the form

$$Q(i, x_{f(i-1)}, \dots, x_{f(1)}) = \alpha(i) \exp \left\{ \sum_{k=1}^{i-1} \phi_k(u_{f(k)}(x_{f(k)})) \right\}, \quad (10)$$

where $\alpha : \{1, \dots, N\} \rightarrow \mathbb{R}_{++}$, $\phi_i : \mathcal{U} \rightarrow \mathbb{R}$ for all $i = 1, \dots, N$ with $\mathcal{U} = \cup_{a \in A} u_a(L_a)$ (i.e., the union of the ranges of all attribute-specific utility functions), and by convention $\sum_{k=1}^0 \phi_k(u_{f(k)}(x_{f(k)})) \equiv 0$. We refer to this model by the triplet (α, u, ϕ) , where $u = (u_a)_{a \in A}$ and $\phi = (\phi_i)_{i=1}^N$. If each ϕ_i is constant, we obtain a random-choice version of our baseline AF model.²³

The idea behind expression (10) is that the *utility* from attributes presented earlier affects the weight assigned to later attributes. The first impression left by the early attributes matters also because it affects the responsiveness to later impressions. For example, suppose each ϕ_k is decreasing. Then, the more Ann likes early attributes, the less she weighs later attributes. Put differently, she may underweight later attributes not just because they come later, but also because earlier attributes are already pretty good. Other possible interpretations are that Ann pays less attention to later attributes if earlier ones are good enough; if instead early attributes are not good, she may look for reasons to like an item by carefully inspecting later attributes (decreasing ϕ_k), or she may become suspicious and lose interest (increasing ϕ_k). In this way, the model allows for smooth forms of satisficing across attributes.

Our characterization involves four axioms. First, suppose items x and y differ only in attribute a and Ann prefers x_a to y_a . Then, we would expect that she also prefers x to y —in probabilistic terms, she is more likely to choose x than y —no matter what the frame is.

Axiom 5 (Attribute Monotonicity) *For every $i = 1, \dots, N - 1$, $a \in A$, and $f, f' \in F$ such that $f(N) = a$ and $f'(i) = a$, the following holds: If $x_{f(N)} = \hat{x}_{f'(i)} = x_a$, $y_{f(N)} = \hat{y}_{f'(-i)} = y_a$, $x_{f(-N)} = y_{f(-N)}$, and $\hat{x}_{f'(-i)} = \hat{y}_{f'(-i)}$, then*

$$\pi(x_f | \{x_f, y_f\}) \geq \frac{1}{2} \quad \Rightarrow \quad \pi(\hat{x}_{f'} | \{\hat{x}_{f'}, \hat{y}_{f'}\}) \geq \frac{1}{2}.$$

This intuitive property rules out some predictions that are possible under expression (10) without further restrictions, but are highly unrealistic. If attribute a appears in position $k < N$ and ϕ_k is decreasing, the better x_a reduces more the weight Ann assigns to later attributes than does y_a . If this effect is strong enough, Ann's overall value of x may be smaller than that of y , leading her to choose y more often. Such violations of simple dominance seem implausible.

²³The model defined by (10) is related to Epstein (1983), to which we owe significant inspiration for our characterization.

Axiom 6 considers the comparison of items whose attributes are identical *up to* some position i . It states that the levels of such attributes do not affect how Ann trades off the attributes after position i . Given any x_f , let $x_f^i = (x_{f(1)}, \dots, x_{f(i)})$. Note that for $p_f \in \Delta(\times_{k=i+1}^N L_{f(k)})$, the object (x_f^i, p_f) defines a lottery in $\Delta(X_f)$.

Axiom 6 (Common-Root Independence) *Fix any $f \in F$ and $i = 1, \dots, N - 1$. For all (x_f^i, p_f) , (y_f^i, p_f) , (x_f^i, q_f) , and (y_f^i, q_f) in $\Delta(X_f)$, we have*

$$\pi((x_f^i, p_f) | \{(x_f^i, p_f), (x_f^i, q_f)\}) = \pi((y_f^i, p_f) | \{(y_f^i, p_f), (y_f^i, q_f)\}).$$

Axiom 7 considers the comparison of items whose attributes are identical *after* some position. It requires that how these identical attributes are ordered does not affect Ann's choice.

Axiom 7 (Tail Frame Invariance) *Fix $i \geq 2$ and any $f, f' \in F$ that satisfy $f(k) = f'(k)$ for $k \leq i - 1$. Let x_f and y_f satisfy $x_{f(k)} = y_{f(k)}$ for $k \geq i$. Let $\hat{x}_{f'}$ and $\hat{y}_{f'}$ satisfy $x_{f(k)} = \hat{x}_{f'(k)}$ and $y_{f(k)} = \hat{y}_{f'(k)}$ for $k \leq i - 1$ and $x_{f'(k)} = \hat{x}_{f'(k)}$ and $y_{f'(k)} = \hat{y}_{f'(k)}$ for $k \geq i$. Then, the following holds*

$$\pi(x_f | \{x_f, y_f\}) = \pi(\hat{x}_{f'} | \{\hat{x}_{f'}, \hat{y}_{f'}\}).$$

Finally, similarly to Axiom 4 for the AF model, Axiom 8 exploits comparisons between frames to identify their effects. It allows the effect of postponing an attribute in the presentation order to depend on the level of the attributes that precede it.

Axiom 8 (Lasting Impressions) *For all $f, f' \in F$ that satisfy $f(i) = f'(1)$ for $i \neq 1$, the following holds: If $x_{f(i)} \neq \hat{x}_{f(i)}$, $x_{f(i)} = y_{f'(1)}$, $\hat{x}_{f(i)} = \hat{y}_{f'(1)}$, $x_{f(-i)} = \hat{x}_{f(-i)}$, $y_{f'(-1)} = \hat{y}_{f'(-1)}$, and $\pi(y_{f'}, \{y_{f'}, \hat{y}_{f'}\}) \neq \pi(\hat{y}_{f'}, \{y_{f'}, \hat{y}_{f'}\})$, then*

$$\frac{\pi(x_f, \{x_f, \hat{x}_f\})}{\pi(\hat{x}_f, \{x_f, \hat{x}_f\})} = r(i, x_{f(1)}, \dots, x_{f(i-1)}) \frac{\pi(y_{f'}, \{y_{f'}, \hat{y}_{f'}\})}{\pi(\hat{y}_{f'}, \{y_{f'}, \hat{y}_{f'}\})}.$$

This is where we relax additive separability. In fact, if we required r to depend at most on i , we would obtain that each ϕ_i is constant and hence a random-choice version of our AF model.

Before stating our result, we introduce the restrictions on (α, u, ϕ) implied by our axioms (in particular Axiom 5). For all $f \in F$, $x_f \in X_f$, and $i = 1, \dots, N - 1$, let

$$R_{\alpha, u, \phi}^i(x_f) = \sum_{j=i+1}^N u_{f(j)}(x_{f(j)}) \alpha(j) \exp \left\{ \sum_{k=i+1}^{j-1} \phi_k(u_{f(k)}(x_{f(k)})) \right\},$$

which is the residual value of x_f after position i . For all $a \in A$ and $x_a \in L_a$, let

$$\sigma_{u, \phi}^i(x_a, y_a) = - \frac{e^{\phi_i(u_a(x_a))} - e^{\phi_i(u_a(y_a))}}{u_a(x_a) - u_a(y_a)}.$$

This measures the relative strength of the framing and direct effect on an item's value of changing x_a with y_a in position i . The next condition ensures that this direct effect always dominates, taking into account the residual value of an item.

Definition 9 (Regularity) *The model (α, u, ϕ) is regular if for all $i = 1, \dots, N$, $f \in F$, and $x_f \in X_f$*

$$\alpha(i) \geq \sup_{y_{f(i)} \in L_{f(i)}} \sigma_{u, \phi}^i(x_{f(i)}, y_{f(i)}) R_{\alpha, u, \phi}^i(x_f).$$

This condition looks complex due to its generality, but is also intuitive. It exactly characterizes our model as the next theorem shows. Note that it holds automatically if all ϕ_i are increasing and all u_a are positive, or if all ϕ_i are decreasing and all u_a are negative. In applications, it is easy to select regular (α, u, ϕ) . If all u_a are bounded and all ϕ_i are differentiable, we can ensure regularity by assuming an appropriate bound on each derivative ϕ'_i .

Theorem 2 *Axioms 5–8 hold if and only if there exist regular (α, u, ϕ) such that for every $f \in F$ the function w_f in expression (9) satisfies*

$$w_f(x_f) = \sum_{i=1}^N u_{f(i)}(x_{f(i)}) \alpha(i) \exp \left\{ \sum_{k=1}^{i-1} \phi_k(u_{f(k)}(x_{f(k)})) \right\}.$$

7.3 Framing in Perturbed-Utility Models

We briefly discuss how to connect our theory of attribute framing with the perturbed-utility model of Fudenberg et al. (2015). This provides a more general way of modeling stochastic choice influenced by such framing effects.

Fudenberg et al. (2015) describe the choice probabilities as resulting from a maximization problem: Ann maximizes the expected utility of her choices net of some cost that is convex in probabilities. Formally, for every M_f

$$(\pi(x_f|M_f), \dots, \pi(y_f|M_f)) \in \arg \max_{\pi \in \Delta(M_f)} \left\{ \sum_{z_f \in M_f} v(z_f) \pi(z_f) - \chi(\pi(z_f)) \right\},$$

where χ is a perturbation function that may reward Ann for randomizing.

We connect the two models building on Fudenberg et al.'s (2015) elegant analysis. Given any continuous and strictly increasing function $h : (0, 1) \rightarrow \mathbb{R}_+$, define the marginal cost as

$$\chi'(\pi) = \ln(h(\pi)).$$

Fudenberg et al. (2015) show that the utility of any two items x_f and y_f satisfies

$$v_f(x_f) - v_f(y_f) = \chi'(\pi(x_f|\{x_f, y_f\})) - \chi'(\pi(y_f|\{x_f, y_f\})).$$

Our previous characterizations involved specifying properties of payoff differences of the form $v_f(x_f) - v_f(y_f)$ through our axioms. Thus, to specify similar properties in the perturbed-utility framework, we only need to reformulate our axioms in terms of the “rescaled” probabilities $h(\pi(x_f|\{x_f, y_f\}))$.

7.4 Framing in Rational-Inattention Models

We can also connect our theory of attribute framing with the random-choice model based on rational inattention proposed by Matějka and McKay (2015). Recall that we can think of M_f as a table, where attributes correspond to the rows $i = 1, \dots, N$ and items to the columns $j = 1, \dots, |M_f|$. Then, choosing an item leads to the consequence of getting the attributes $(x_{f(i)}^j)_{i=1}^N$ in the corresponding column j .

The rational-inattention model is based on the idea that the decision-maker is uncertain about the consequences of his choices and spends costly attention to learn about them. In our case, suppose Bob is uncertain about the entries of the table (i.e., M_f) and so about the levels of the attributes obtained by selecting a specific item (i.e., column). Let G be his prior about the entries of menus. As in Matějka and McKay (2015), suppose Bob allocates attention to the items in a menu incurring a cost in the form of entropy reduction. Following Matějka and McKay (2015), the solution to Bob’s optimal attention-allocation problem leads to the choice probabilities

$$\pi(x_f^j|M_f) = \frac{\bar{\pi}^j e^{v_f(x_f^j)}}{\sum_{j'} \bar{\pi}^{j'} e^{v_f(x_f^{j'})}}, \quad \text{where } \bar{\pi}^i = \mathbb{E}_G[\pi(x_f^i|M_f)] \quad (11)$$

for every M_f . Thus, for every realization of the entries in table M_f , the probability that Bob chooses the item in column j is similar to our Luce model (9), except for the additional weights $\bar{\pi}^j$. Each $\bar{\pi}^j$ equals the ex-ante probability of choosing the item in column j averaging over all realizations of M_f .

We can connect our theory with this model as follows. Note that expression (11) implies that, for fixed column j ,

$$v_f(x_f^j) - v_f(y_f^j) = \ln(\pi(x_f^j|M_f)) - \ln(\pi(y_f^j|M_f)).$$

This again suggests a simple way to adapt our axioms to specialize the function v_f in the present context as we did in the previous characterizations.

The flexibility of Matějka and McKay’s (2015) framework allows for several extensions of our theory of attribute-framing effects. First, it allows to overcome some of the well-known limitations of the Luce model (like unrealistic responses to item copies). More interestingly for us, it allows for interactions between attribute-order and list-order effects. To illustrate the point, suppose that as Bob goes through the columns from left to right, he gets tired and pays less attention to later columns. In this case, no matter what, he is overall less likely to choose items in later columns (i.e.,

$\bar{\pi}^j > \bar{\pi}^{j+1}$ for all j). This can be formalized by assuming a prior G specifying that later columns are very likely to have sufficiently bad attribute realizations. Importantly, the average weights $\bar{\pi}^j$ have to be consistent with the actual choice frequencies, which are affected by the attribute frames. Therefore, those frames ultimately influence the list-order effects. We leave studying such interactions for future research.

8 Final Remarks

We introduced a model of framing effects that explicitly takes into account how alternatives are presented to people. The order or emphasis given to the attributes of available items can influence which is chosen. This is at odds with mainstream choice theory for which the presentation of the attributes should be irrelevant, but is in line with rich empirical evidence suggesting that such effects should be taken into account when studying choice behavior.

The model provides a first theoretical structure to understand such attribute-framing effects. It provides testable predictions and the possibility to compare framing effects across individuals. It can be easily generalized to allow for richer framing effects. In particular, it may open a bridge between attribute-order effects and list-order effects. Moreover, the model has several interesting implications for competition among firms and negotiation.

Many other applications wait to be written. We mention a couple that we find intriguing. The first is to formally study rhetoric and its concerns with how to arrange the points of an argument in the most persuasive manner—what the classics called *dispositio*. This could offer a novel angle on persuasion that is fundamentally different from strategic information provision, because here the *what* stays the same but the *how* changes. The second application involves relating attribute-order effects and present bias. We can view consumption streams as items and per-period consumption levels as their attributes. One wonders whether presenting streams in chronological or reversed chronological order changes the displayed present bias. If so, framing may emerge as another way to address this bias.

Acknowledgements: We thank for comments and suggestions Jose Apesteguia, Yeon-Koo Che, Drew Fudenberg, Botond Köszegi, Bart Lipman, Fabio Maccheroni, Mark Machina, Pietro Ortoleva, Charles Sprenger, Steve Tadelis, and all seminar participants at BRIC 2019, Barcelona GSE Summer Forum on Stochastic Choice 2019, and ESSET 2019. Francesco Cerigioni is grateful to the Ministerio de Ciencia y Innovación for the grant PGC2018-098949-B-I00 and financial support through the Severo Ochoa Programme for Centers of Excellence in R&D SEV-2015-0563. All remaining errors are ours.

Appendix

A Strategic Framing: Proofs

A.1 Proof of Lemma 1

Fix f and suppose $\Delta_f > 0$ —the other case follows similarly. Let θ_f^* identify the type of consumer indifferent between I_f and E_f :

$$\theta_f^*[\alpha_f(r)I_r + \alpha_f(b)I_b] - \alpha_f(p)I_p = \theta_f^*[\alpha_f(r)E_r + \alpha_f(b)E_b] - \alpha_f(p)E_p$$

and therefore

$$\theta_f^* = \frac{I_p - E_p}{\Delta_f} \alpha_f(p).$$

The demand for the incumbent's and the entrant's product is then respectively

$$\bar{\theta} - \frac{I_p - E_p}{\Delta_f} \alpha_f(p) \quad \text{and} \quad \frac{I_p - E_p}{\Delta_f} \alpha_f(p) - \underline{\theta}.$$

The firms' profit-maximization problems are

$$\max_{I_p} \left(\bar{\theta} - \frac{I_p - E_p}{\Delta_f} \alpha_f(p) \right) I_p \quad \text{and} \quad \max_{E_p} \left(\frac{E_p - I_p}{\Delta_f} \alpha_f(p) - \underline{\theta} \right) E_p.$$

They result in the following best response functions:

$$I_p = \frac{1}{2} \left[E_p + \frac{\bar{\theta} \Delta_f}{\alpha_f(p)} \right] \quad \text{and} \quad E_p = \frac{1}{2} \left[I_p - \frac{\underline{\theta} \Delta_f}{\alpha_f(p)} \right].$$

Solving this system of equations leads to the claimed equilibrium prices, which we can substitute in the profit functions to derive $\varphi^o(I)$ and $\varphi^o(E)$.

A.2 Lemma 2

Lemma 2 *Under monopoly (i.e., $K = +\infty$), the incumbent's optimal frame is f^m and $\varphi^m(I_{f^m}) > \varphi^m(I_{f^*}) > \varphi^m(I_{f_*})$.*

Proof. Suppose $K = +\infty$ and fix f . Given I_p , the type of consumers that is indifferent between buying I_f or nothing is

$$\theta_f^m = \frac{\alpha_f(p)}{\alpha_f(r)I_r + \alpha_f(b)I_b} I_p.$$

Thus, the monopolist maximizes

$$I_p \left(\bar{\theta} - \frac{\alpha_f(p)}{\alpha_f(r)I_r + \alpha_f(b)I_b} I_p \right),$$

which leads to the optimal monopolistic price and profit

$$I_p = \left[\frac{\alpha_f(r)}{\alpha_f(p)} I_r + \frac{\alpha_f(b)}{\alpha_f(p)} I_b \right] \frac{\bar{\theta}}{2} \quad \text{and} \quad \varphi^m(I_f) = \left[\frac{\alpha_f(r)}{\alpha_f(p)} I_r + \frac{\alpha_f(b)}{\alpha_f(p)} I_b \right] \frac{\bar{\theta}^2}{4}.$$

It is easy to see that f^m maximizes $\varphi^m(I_f)$ and $\varphi^m(I_{f^m}) > \varphi^m(I_{f^*}) > \varphi^m(I_{f_*})$. ■

A.3 Proof of Propositions 1, 2, and 3

Recall that $\varphi^o(E_f)$ is proportional to $\varphi^o(I_f)$, so the incumbent and the entrant rank frames in the same way.

If $K > \varphi^o(E_{f^m})$, then the incumbent can choose the monopoly-optimal frame f^m and deter entry. By condition 2, this is the best strategy for the incumbent.

If instead $K \leq \min\{\varphi^o(E_{f^m}), \varphi^o(E_{f_*})\}$, then for every choice of f we have $\varphi^o(E_f) \geq K$. Thus, the incumbent cannot prevent entry. In this case, by Lemma 1 the incumbent will always choose f such that $\Delta_f > 0$: The optimal f always presents attribute r before attribute b . Given this, to maximize $\varphi^o(I_f)$, $f \in \{f^m, f^*, f_*\}$ has to maximize

$$\frac{\Delta_f}{\alpha_f(p)} = \left[\frac{\alpha_f(r)}{\alpha_f(p)} - \frac{\alpha_f(b)}{\alpha_f(p)} \right] \delta.$$

This implies that $\varphi^o(I_{f_*}) > \varphi^o(I_{f^*})$, so the optimal frame is either f^m or f^* . However, $\varphi^o(I_{f^m}) > \varphi^o(I_{f^*})$ if and only if

$$\frac{\alpha(1)}{\alpha(3)} - \frac{\alpha(2)}{\alpha(3)} > \frac{\alpha(1)}{\alpha(2)} - \frac{\alpha(3)}{\alpha(2)} \quad \Leftrightarrow \quad \alpha(2) < \alpha(1) - \alpha(3) = \underline{\alpha}(2).$$

This completes the proof of Proposition 1.

To prove Proposition 2, note that if $\alpha(2) > \underline{\alpha}(2)$, then $\varphi^o(I_{f^m}) > \varphi^o(I_{f_*})$ if and only if

$$\frac{\alpha(1)}{\alpha(2)} - \frac{\alpha(2)}{\alpha(3)} > \frac{\alpha(2)}{\alpha(1)} - \frac{\alpha(3)}{\alpha(1)} \quad \Leftrightarrow \quad \alpha(2) < \frac{[\alpha(1)]^2 + [\alpha(3)]^2}{\alpha(1) + \alpha(3)} = \bar{\alpha}(2).$$

Thus, if $\bar{\alpha}(2) > \alpha(2) > \underline{\alpha}(2)$, then $\varphi^o(E_{f_*}) > \varphi^o(E_{f^m}) > \varphi^o(E_{f^*})$ and the optimal frame under oligopoly is f^* . Given this, the only case where the incumbent can use framing to deter entry is if $K \in (\varphi^o(E_{f_*}), \varphi^o(E_{f^m})]$, which requires to switch to frame f_* —every other frame in F yields $\varphi^o(E_f) \geq \varphi^o(E_{f^m})$ and hence cannot deter entry. The incumbent prefers deterring entry with f_* to allowing entry by choosing f^* if and only if $\varphi^m(I_{f_*}) \geq \varphi^o(I_{f^*})$, which is equivalent to condition (3).

If instead $\alpha(2) < \underline{\alpha}(2)$, then $\varphi^o(E_{f^m}) > \varphi^o(E_{f^*}) > \varphi^o(E_{f_*})$ and the optimal frame under oligopoly is f^m . Thus, there are two cases in which the incumbent can use framing to deter entry. The first is if $K \in (\varphi^o(E_{f^*}), \varphi^o(E_{f^m})]$, which implies that best entry-deterrent frame is f^* . Every other frame in F either yields $\varphi^o(E_f) \geq \varphi^o(E_{f^m})$ —hence cannot deter entry—or it yields $\varphi^m(I_f) < \varphi^m(I_{f^*})$. Given this, the incumbent prefers

detering entry with f^* to allowing entry by choosing f^m if and only if $\varphi^m(I_{f^*}) \geq \varphi^o(I_{f^m})$, which is equivalent to condition (4). The second case is if $K \in (\varphi^o(E_{f^*}), \varphi^o(E_{f^m})]$, which implies that best entry-deterrent frame is f^* . Again, every other frame in F either yields $\varphi^o(E_f) \geq \varphi^o(E_{f^*})$ —hence cannot deter entry—or it yields $\varphi^m(I_f) < \varphi^m(I_{f^*})$. Given this, the incumbent prefers deterring entry with f^* to allowing entry by choosing f^m if and only if $\varphi^m(I_{f^*}) \geq \varphi^o(I_{f^m})$, which is equivalent to condition (5).

Consider now Proposition 3. It is easy to see that if either δ is lower or $\underline{\theta}$ is higher, then $\varphi^o(E_{f^*})$, $\varphi^o(E_{f^m})$, and $\varphi^o(E_{f^*})$ are all lower. Thus, the range of entry costs where the incumbent faces no entry threat (i.e., $K > \varphi^o(E_{f^m})$) expands, and the range of costs where the incumbent can never prevent entry (i.e., $K \leq \min\{\varphi^o(E_{f^m}), \varphi^o(E_{f^*})\}$) shrinks. In addition, the right-hand side of conditions (3), (4), and (5) are all smaller if either δ is lower or $\underline{\theta}$ is higher. Thus, for $K \in (\varphi^o(E_{f^*}), \max\{\varphi^o(E_{f^m}), \varphi^o(E_{f^*})\}]$, the incumbent is more likely to use framing to deter entry.

Finally, consider α and α' such that $\frac{\alpha(1)}{\alpha(2)} \geq \frac{\alpha'(1)}{\alpha'(2)}$ and $\frac{\alpha(2)}{\alpha(3)} \geq \frac{\alpha'(2)}{\alpha'(3)}$. It is easy to see that this implies that $\varphi^o(E_{f^m})$ and $\varphi^o(E_{f^*})$ are both lower under α than under α' . Moreover, the left-hand side of conditions (3) and (5) is higher under α' than under α , while the right-hand side of conditions (3) and (5) is lower under α' than under α . Thus, whenever K falls in the region where deterring entry requires to use f^* , if the incumbent finds this optimal under α , it also finds it optimal under α' . Regarding $\varphi^o(E_{f^*})$ and condition 4, their ranking under α and α' is ambiguous.

B Framing in Negotiations: Proofs

B.1 Proof of Proposition 4

Given f , agent P solves is

$$\begin{aligned} \max_x \quad & \sum_{i=1}^N \alpha(i) [\beta_{f(i)} - \gamma_{f(i)}(x_{f(i)} - \bar{x}_{f(i)}^P)^2] \\ \text{s.t.} \quad & \sum_{i=1}^N \alpha(i) [\beta_{f(i)} - \gamma_{f(i)}(x_{f(i)} - \bar{x}_{f(i)}^R)^2] \geq \bar{u} \end{aligned} \quad (12)$$

That leads to the following Lagrangian, with standard positivity and slack conditions:

$$\max_{x, \lambda} \quad \sum_{i=1}^N \alpha(i) [\beta_{f(i)} - \gamma_{f(i)}(x_{f(i)} - \bar{x}_{f(i)}^P)^2] + \lambda \left(\sum_{i=1}^N \alpha(i) [\beta_{f(i)} - \gamma_{f(i)}(x_{f(i)} - \bar{x}_{f(i)}^R)^2] \right) \quad (13)$$

Thus we get the following necessary and sufficient FOC for $i = 1, \dots, N$:

$$\begin{aligned} x_{f(i)} : \quad & -2\alpha(i) \left[\gamma_{f(i)}(x_{f(i)} - \bar{x}_{f(i)}^P) + \gamma_{f(i)}\lambda(x_{f(i)} - \bar{x}_{f(i)}^R) \right] = 0 \\ \lambda : \quad & \sum_{i=1}^N \alpha(i) [\kappa_{f(i)} - \gamma_{f(i)}(x_{f(i)} - \bar{x}_{f(i)}^R)^2] = \bar{u}^R \end{aligned} \quad (14)$$

So we get

$$x_{f(i)} = \frac{1}{1+\lambda} \bar{x}_{f(i)}^P + \frac{\lambda}{1+\lambda} \bar{x}_{f(i)}^R. \quad (15)$$

Substituting the optimal x_f into agent P 's objective function and agent R 's participation constraint we obtain

$$\begin{aligned} & \sum_{i=1}^N \alpha(i) \left[\beta_{f(i)} - \left(\frac{\lambda}{1+\lambda} \right)^2 \gamma_{f(i)} (\bar{x}_{f(i)}^R - \bar{x}_{f(i)}^P)^2 \right] \\ = & \sum_{i=1}^N \alpha(i) \beta_{f(i)} - \left(\frac{\lambda}{1+\lambda} \right)^2 \sum_{i=1}^N \alpha(i) \gamma_{f(i)} (\bar{x}_{f(i)}^R - \bar{x}_{f(i)}^P)^2 \end{aligned}$$

and

$$\begin{aligned} & \left(\frac{1}{1+\lambda} \right)^2 \sum_{i=1}^N \alpha(i) \gamma_{f(i)} (\bar{x}_{f(i)}^P - \bar{x}_{f(i)}^R)^2 = \sum_{i=1}^N \alpha(i) \beta_{f(i)} - \bar{u}^R \\ \Rightarrow \lambda &= \frac{\sqrt{\sum_{i=1}^N \alpha(i) \gamma_{f(i)} (\bar{x}_{f(i)}^P - \bar{x}_{f(i)}^R)^2} - \sqrt{\sum_{i=1}^N \alpha(i) \beta_{f(i)} - \bar{u}^R}}{\sqrt{\sum_{i=1}^N \alpha(i) \beta_{f(i)} - \bar{u}^R}} \\ &= \sqrt{\frac{\Gamma(f)}{B(f) - \bar{u}^R}} - 1. \end{aligned}$$

The assumption on $\bar{u}^R$ implies that $\Gamma(f) > B(f) - \bar{u}^R$ for all f and hence $\lambda > 1$. We can then use this last conditions to replace λ in agent P 's objective and obtain

$$\begin{aligned} U^P(f) &= \sum_{i=1}^N \alpha(i) \beta_{f(i)} - \left(\frac{\lambda}{1+\lambda} \right)^2 \sum_{i=1}^N \alpha(i) \gamma_{f(i)} (\bar{x}_{f(i)}^R - \bar{x}_{f(i)}^P)^2 \\ &= B(f) - \lambda^2 (B(f) - \bar{u}^R) \\ &= B(f) - \left[\sqrt{\Gamma(f)} - \sqrt{B(f) - \bar{u}^R} \right]^2. \end{aligned}$$

Since the quantity in squared brackets is always positive, $U^P(f)$ is increasing in $B(f)$ and decreasing in $\Gamma(f)$. Proposition 4 follows.

B.2 Proof of Corollary 2

We can view α^1 and α^2 as probability distributions over the positions $\{1, \dots, N\}$. Then, the condition of Definition 3 can be read as α^2 MLR dominates α^1 , which in turn implies that α^2 FOSD α^1 . Given the optimal framing strategy f^* in Proposition 4, we have that $\beta_{f^*(i)}$ is an decreasing function of i and $\gamma_{f^*(i)} (\bar{x}_{f^*(i)}^R - \bar{x}_{f^*(i)}^P)$ is an increasing function of i under both α^1 and α^2 . Standard results imply that $\Gamma^2(f^*) > \Gamma^1(f^*)$ and $B^2(f^*) < B^1(f^*)$. Therefore, using the expression of U^P , we get that the proposer is better off when payoffs are defined by α^1 than by α^2 .

C Characterization of the AF Model

C.1 Proof of Theorem 1

We will prove sufficiency of Axioms 1–4; necessity is easy to verify and is thus omitted. Given Assumption 1, the condition $c(\{x_f, x'_f\}) = c(\{y_f, y'_f\})$ implies that

$$v_f(x_f) > (=) v_f(x'_f) \Leftrightarrow v_f(y_f) > (=) v_f(y'_f).$$

Given the restrictions on $x_f, x'_f, y_f,$ and y'_f in Axiom 2, this means that how v_f ranks the attributes in positions j and k is independent of the other attributes' levels. By Axiom 1 each position of f can matter for choice. By standard arguments (Debreu (1960)), we can write v_f in an additive form in positions and the attribute assigned by f to each position:

$$v_f(x_f) = \sum_{i=1}^N w_{i,f(i)}^f(x_{f(i)}), \quad (16)$$

where for every $f \in F$ there exist non-constant $w_{i,f(i)}^f : L_{f(i)} \rightarrow \mathbb{R}$ for every $i \in \{1, \dots, N\}$. Additive forms are unique up to positive affine transformations, which in this case can depend on f : If we have two such representations v_f and $\hat{v}_f$ of choice under f , we must have $v_f = \beta_f \hat{v}_f + \xi_f$, where $\beta_f > 0$ and $\xi_f \in \mathbb{R}$.

Given this, Axiom 3 implies that for every $a \in A$, if $f(i) = f'(j) = a$, then $w_{i,f(i)}^f$ and $w_{j,f'(j)}^{f'}$ represent the same vN-M utility function over L_a . Therefore, for every $a \in A$, fix any $f_a \in F$ such that $f_a(1) = a$. Let $u_a = w_{1,f_a}^{f_a}$. For any other $f \in F$ and $i = 1, \dots, N$ such that $f(i) = a$, we have that $w_{i,a}^f = \gamma_i^f u_a + \zeta_i^f$, where $\gamma_i^f > 0$ and $\zeta_i^f \in \mathbb{R}$.²⁴ Letting $\gamma_1^{f_a} = 1$ and $\zeta_1^{f_a} = 0$, we can write

$$v_f(x_f) = \sum_{i=1}^N w_{i,f(i)}^f(x_{f(i)}) = \sum_{i=1}^N \gamma_i^f u_{f(i)}(x_{f(i)}) + \sum_{i=1}^N \zeta_i^f.$$

By affine uniqueness of v_f , for all $f \in F$ we can let $\gamma_1^f = 1$ and $\zeta_i^f = 0$ for all $i = 1, \dots, N$. Therefore, we obtain the representation

$$v_f(x_f) = u_{f(1)}(x_{f(1)}) + \sum_{i=2}^N \gamma_i^f u_{f(i)}(x_{f(i)}).$$

Next, we want to show that γ_i^f depends only on i . By definition of $x_f, y_f,$ and z_f in Axiom 4, we have

$$p_{xyz_f}[u_{f(1)}(x_{f(1)}) + \gamma_i^f u_{f(i)}(x_{f(i)})] + (1 - p_{xyz_f})[u_{f(1)}(y_{f(1)}) + \gamma_i^f u_{f(i)}(y_{f(i)})]$$

²⁴Note that γ_i^f cannot also depend on $f(i)$ in addition to i and f because there is no $\gamma_{i,b}^f$ for $b \neq f(i)$. If any, the superscript already allows for the dependence on $f(i)$.

$$= u_{f(1)}(y_{f(1)}) + \gamma_i^f u_{f(i)}(x_{f(i)}).$$

This implies that

$$\frac{p_{xyz_f}}{1 - p_{xyz_f}} = \gamma_i^f \frac{u_{f(i)}(x_{f(i)}) - u_{f(i)}(y_{f(i)})}{u_{f(1)}(x_{f(1)}) - u_{f(1)}(y_{f(1)})}.$$

Similarly,

$$\frac{p_{xyz_{f'}}}{1 - p_{xyz_{f'}}} = \gamma_j^{f'} \frac{u_{f'(j)}(x_{f'(j)}) - u_{f'(j)}(y_{f'(j)})}{u_{f'(1)}(x_{f'(1)}) - u_{f'(1)}(y_{f'(1)})} = \gamma_j^{f'} \frac{u_{f(i)}(x_{f(i)}) - u_{f(i)}(y_{f(i)})}{u_{f(1)}(x_{f(1)}) - u_{f(1)}(y_{f(1)})}.$$

Therefore,

$$\frac{p_{xyz_f}}{1 - p_{xyz_f}} = \frac{\gamma_i^f}{\gamma_j^{f'}} \cdot \frac{p_{xyz_{f'}}}{1 - p_{xyz_{f'}}}.$$

By Axiom 4, we have

$$\frac{\gamma_i^f}{\gamma_j^{f'}} = g(i, j)$$

and therefore $\gamma_i^f = \gamma_i > 0$ and $\gamma_j^{f'} = \gamma_j > 0$ for all f, f' .

We conclude that for every $f \in F$ and $x_f \in X_f$

$$v_f(x_f) = u_{f(1)}(x_{f(1)}) + \sum_{i=2}^N \gamma_i u_{f(i)}(x_{f(i)}).$$

C.2 Proof of Proposition 5

Consider primacy effect—the argument is the same for recency effect. Recall that $x_a \succ_a y_a$ if and only if $u_a(x_a) > u_a(y_a)$. Fix any $i = 1, \dots, N-1$. Using the AF representation, we have that

$$\{(x_f, y_f; p_{xyz_f}^i), z_f\} = c((x_f, y_f; p_{xyz_f}^i), z_f)$$

is equivalent to

$$\sum_{j=1}^N \alpha(j) u_{f(j)}(z_{f(j)}) = p_{xyz_f}^i \left\{ \sum_{j=1}^N \alpha(j) u_{f(j)}(x_{f(j)}) \right\} + (1 - p_{xyz_f}^i) \left\{ \sum_{j=1}^N \alpha(j) u_{f(j)}(y_{f(j)}) \right\}.$$

Since $x_{f(j)} = y_{f(j)} = z_{f(j)}$ for all $j \neq i, i+1$, this condition becomes

$$\begin{aligned} & \alpha(i) u_{f(i)}(z_{f(i)}) + \alpha(i+1) u_{f(i+1)}(z_{f(i+1)}) = \\ & p_{xyz_f}^i \{ \alpha(i) u_{f(i)}(x_{f(i)}) + \alpha(i+1) u_{f(i+1)}(x_{f(i+1)}) \} \\ & + (1 - p_{xyz_f}^i) \{ \alpha(i) u_{f(i)}(y_{f(i)}) + \alpha(i+1) u_{f(i+1)}(y_{f(i+1)}) \}. \end{aligned}$$

Using $x_{f(i)} = z_{f(i)}$ and $y_{f(i+1)} = z_{f(i+1)}$, we obtain

$$\frac{p_{xyz_f}^i}{1 - p_{xyz_f}^i} = \frac{\alpha(i) [u_{f(i)}(x_{f(i)}) - u_{f(i)}(y_{f(i)})]}{\alpha(i+1) [u_{f(i+1)}(x_{f(i+1)}) - u_{f(i+1)}(y_{f(i+1)})]}.$$

By similar calculations, using $x_{f'(i+1)} = z_{f'(i+1)}$ and $y_{f'(i)} = z_{f'(i)}$, we have

$$\frac{p_{xyz_{f'}}^i}{1 - p_{xyz_{f'}}^i} = \frac{\alpha(i+1)[u_{f'(i+1)}(x_{f'(i+1)}) - u_{f'(i+1)}(y_{f'(i+1)})]}{\alpha(i)[u_{f'(i)}(x_{f'(i)}) - u_{f'(i)}(y_{f'(i)})]}.$$

Since $x_{f'}$, $y_{f'}$, and $z_{f'}$ are obtained from x_f , y_f , and z_f by swapping the attributes in positions i and $i+1$, we have

$$\frac{p_{xyz_{f'}}^i}{1 - p_{xyz_{f'}}^i} = \frac{\alpha(i+1)[u_{f(i)}(x_{f(i)}) - u_{f(i)}(y_{f(i)})]}{\alpha(i)[u_{f(i+1)}(x_{f(i+1)}) - u_{f(i+1)}(y_{f(i+1)})]}.$$

It follows that

$$\frac{p_{xyz_f}^i}{1 - p_{xyz_f}^i} = \left[\frac{\alpha(i)}{\alpha(i+1)} \right]^2 \cdot \frac{p_{xyz_{f'}}^i}{1 - p_{xyz_{f'}}^i}.$$

This implies that $\alpha(i) > \alpha(i+1)$ if and only if $p_{xyz_f}^i > p_{xyz_{f'}}^i$ as desired.

C.3 Proof of Proposition 6

We will prove the result for primacy effect—the argument is similar for recency effect.

Part 1: “only if”. As in the proof of Proposition 5,

$$\{(x_f, y_f; p_{xyz_f}^i), z_f\} = c^2((x_f, y_f; p_{xyz_f}^i), z_f)$$

is equivalent to

$$\frac{p_{xyz_f}^i}{1 - p_{xyz_f}^i} = \frac{\alpha^2(i)[u_{f(i)}^2(x_{f(i)}) - u_{f(i)}^2(y_{f(i)})]}{\alpha^2(i+1)[u_{f(i+1)}^2(x_{f(i+1)}) - u_{f(i+1)}^2(y_{f(i+1)})]}.$$

The condition $\{(x_{f'}, y_{f'}; p_{xyz_{f'}}^i), z_{f'}\} = c^2((x_{f'}, y_{f'}; p_{xyz_{f'}}^i), z_{f'})$ is equivalent to

$$\frac{p_{xyz_{f'}}^i}{1 - p_{xyz_{f'}}^i} = \frac{\alpha^2(i+1)[u_{f(i)}^2(x_{f(i)}) - u_{f(i)}^2(y_{f(i)})]}{\alpha^2(i)[u_{f(i+1)}^2(x_{f(i+1)}) - u_{f(i+1)}^2(y_{f(i+1)})]}.$$

By similar calculations, $z_f \in c^1((x_f, y_f; p_{xyz_f}^i), z_f)$ is equivalent to

$$\frac{p_{xyz_f}^i}{1 - p_{xyz_f}^i} \leq \frac{\alpha^1(i)[u_{f(i)}^1(x_{f(i)}) - u_{f(i)}^1(y_{f(i)})]}{\alpha^1(i+1)[u_{f(i+1)}^1(x_{f(i+1)}) - u_{f(i+1)}^1(y_{f(i+1)})]}.$$

The condition $(x_{f'}, y_{f'}; p_{xyz_{f'}}^i) \in c^1((x_{f'}, y_{f'}; p_{xyz_{f'}}^i), z_{f'})$ is equivalent to

$$\frac{p_{xyz_{f'}}^i}{1 - p_{xyz_{f'}}^i} \geq \frac{\alpha^1(i+1)[u_{f(i)}^1(x_{f(i)}) - u_{f(i)}^1(y_{f(i)})]}{\alpha^1(i)[u_{f(i+1)}^1(x_{f(i+1)}) - u_{f(i+1)}^1(y_{f(i+1)})]}.$$

Therefore, these conditions are equivalent to

$$\left[\frac{\alpha^2(i)}{\alpha^2(i+1)} \right]^2 = \frac{\frac{p_{xyz_f}^i}{1-p_{xyz_f}^i}}{\frac{p_{xyz_{f'}}^i}{1-p_{xyz_{f'}}^i}} \leq \left[\frac{\alpha^1(i)}{\alpha^1(i+1)} \right]^2$$

for all $i = 1, \dots, N-1$.

Part 2: “if”. Let $\bar{p}_{xyz_f}^i$ and $\bar{p}_{xyz_{f'}}^i$ be the indifference probabilities for decision-maker 1:

$$\frac{\bar{p}_{xyz_f}^i}{1 - \bar{p}_{xyz_f}^i} = \frac{\alpha^1(i)[u_{f(i)}^1(x_{f(i)}) - u_{f(i)}^1(y_{f(i)})]}{\alpha^1(i+1)[u_{f(i+1)}^1(x_{f(i+1)}) - u_{f(i+1)}^1(y_{f(i+1)})]},$$

$$\frac{\bar{p}_{xyz_{f'}}^i}{1 - \bar{p}_{xyz_{f'}}^i} = \frac{\alpha^1(i+1)[u_{f(i)}^1(x_{f(i)}) - u_{f(i)}^1(y_{f(i)})]}{\alpha^1(i)[u_{f(i+1)}^1(x_{f(i+1)}) - u_{f(i+1)}^1(y_{f(i+1)})]}.$$

Using $u_a^1 = \gamma_a u_a^2 + \zeta_a$ for all $a \in A$ and $\frac{\alpha^1(i)}{\alpha^1(i+1)} \geq \frac{\alpha^2(i)}{\alpha^2(i+1)}$ for all $i = 1, \dots, N-1$, we have

$$\begin{aligned} \frac{\bar{p}_{xyz_f}^i}{1 - \bar{p}_{xyz_f}^i} &= \frac{\alpha^1(i)\gamma_{f(i)}[u_{f(i)}^2(x_{f(i)}) - u_{f(i)}^2(y_{f(i)})]}{\alpha^1(i+1)\gamma_{f(i+1)}[u_{f(i+1)}^2(x_{f(i+1)}) - u_{f(i+1)}^2(y_{f(i+1)})]} \\ &\geq \frac{\alpha^2(i)\gamma_{f(i)}[u_{f(i)}^2(x_{f(i)}) - u_{f(i)}^2(y_{f(i)})]}{\alpha^2(i+1)\gamma_{f(i+1)}[u_{f(i+1)}^2(x_{f(i+1)}) - u_{f(i+1)}^2(y_{f(i+1)})]} \\ &= \frac{\gamma_{f(i)}}{\gamma_{f(i+1)}} \frac{p_{xyz_f}^i}{1 - p_{xyz_f}^i}. \end{aligned}$$

Similar calculations imply that

$$\frac{\bar{p}_{xyz_{f'}}^i}{1 - \bar{p}_{xyz_{f'}}^i} \leq \frac{\gamma_{f(i)}}{\gamma_{f(i+1)}} \frac{p_{xyz_{f'}}^i}{1 - p_{xyz_{f'}}^i}.$$

There are two cases to consider. First, suppose $\gamma_f(i) \geq \gamma_{f(i+1)}$. It follows that

$$\frac{\bar{p}_{xyz_{f'}}^i}{1 - \bar{p}_{xyz_{f'}}^i} \leq \frac{p_{xyz_{f'}}^i}{1 - p_{xyz_{f'}}^i},$$

which implies that $(x_{f'}, y_{f'}; p_{xyz_{f'}}^i) \in c^1((x_{f'}, y_{f'}; p_{xyz_{f'}}^i), z_{f'})$.

Now suppose that $z_f \notin c^1((x_f, y_f; p_{xyz_f}^i), z_f)$. This means that

$$\begin{aligned} \frac{p_{xyz_f}^i}{1 - p_{xyz_f}^i} &> \frac{\alpha^1(i)[u_{f(i)}^1(x_{f(i)}) - u_{f(i)}^1(y_{f(i)})]}{\alpha^1(i+1)[u_{f(i+1)}^1(x_{f(i+1)}) - u_{f(i+1)}^1(y_{f(i+1)})]} \\ &= \frac{\alpha^1(i)\gamma_{f(i)}[u_{f(i)}^2(x_{f(i)}) - u_{f(i)}^2(y_{f(i)})]}{\alpha^1(i+1)\gamma_{f(i+1)}[u_{f(i+1)}^2(x_{f(i+1)}) - u_{f(i+1)}^2(y_{f(i+1)})]} \end{aligned}$$

$$= \frac{\alpha^1(i)\gamma_{f(i)}\alpha^2(i+1)p_{xyz_f}^i}{\alpha^1(i+1)\gamma_{f(i+1)}\alpha^2(i)(1-p_{xyz_f}^i)},$$

which implies that

$$\frac{\gamma_{f(i)}}{\gamma_{f(i+1)}} < \frac{\alpha^2(i)/\alpha^2(i+1)}{\alpha^1(i)/\alpha^1(i+1)} \leq 1.$$

This contradicts $\gamma_f(i) \geq \gamma_{f(i+1)}$ and thus proves $z_f \in c^1((x_f, y_f; p_{xyz_f}^i), z_f)$.

For the second case, suppose that $\gamma_f(i) < \gamma_{f(i+1)}$. It follows that

$$\frac{\bar{p}_{xyz_f}^i}{1 - \bar{p}_{xyz_f}^i} > \frac{p_{xyz_f}^i}{1 - p_{xyz_f}^i},$$

which implies that $z_f \in c^1((x_f, y_f; p_{xyz_f}^i), z_f)$. Now suppose that

$$(x_{f'}, y_{f'}; p_{xyz_{f'}}^i) \notin c^1((x_{f'}, y_{f'}; p_{xyz_{f'}}^i), z_{f'}).$$

This means that

$$\begin{aligned} \frac{p_{xyz_{f'}}^i}{1 - p_{xyz_{f'}}^i} &< \frac{\alpha^1(i+1)[u_{f(i)}^1(x_{f(i)}) - u_{f(i)}^1(y_{f(i)})]}{\alpha^1(i)[u_{f(i+1)}^1(x_{f(i+1)}) - u_{f(i+1)}^1(y_{f(i+1)})]} \\ &= \frac{\alpha^1(i+1)\gamma_{f(i)}[u_{f(i)}^2(x_{f(i)}) - u_{f(i)}^2(y_{f(i)})]}{\alpha^1(i)\gamma_{f(i+1)}[u_{f(i+1)}^2(x_{f(i+1)}) - u_{f(i+1)}^2(y_{f(i+1)})]} \\ &= \frac{\alpha^1(i+1)\gamma_{f(i)}\alpha^2(i)p_{xyz_{f'}}^i}{\alpha^1(i)\gamma_{f(i+1)}\alpha^2(i+1)(1-p_{xyz_{f'}}^i)}, \end{aligned}$$

which implies that

$$\frac{\gamma_{f(i)}}{\gamma_{f(i+1)}} > \frac{\alpha^2(i+1)/\alpha^2(i)}{\alpha^1(i+1)/\alpha^1(i)} \geq 1.$$

This contradicts $\gamma_f(i) < \gamma_{f(i+1)}$ and thus proves $(x_{f'}, y_{f'}; p_{xyz_{f'}}^i) \in c^1((x_{f'}, y_{f'}; p_{xyz_{f'}}^i), z_{f'})$.

References

- Ajzen, I. and M. Fishbein (1980). Understanding attitudes and predicting social behavior. Englewood cliffs NJ: Prentice–Hall.
- Allen, R. and J. Rehbeck (2016). Revealed stochastic choice with attributes. Working Paper, Ohio State University.
- Apesteagua, J. and M. A. Ballester (2015). A measure of rationality and welfare. Journal of Political Economy 123(6), 1278–1310.
- Apesteagua, J. and M. A. Ballester (2018). Monotone stochastic choice models: The case of risk and time preferences. Journal of Political Economy 126(1), 74–106.
- Auspurg, K. and A. Jäckle (2017). First equals most important? order effects in vignette-based measurement. Sociological Methods & Research 46(3), 490–539.
- Beasley, R. K. and M. R. Joslyn (2001). Cognitive dissonance and post-decision attitude change in six presidential elections. Political psychology 22(3), 521–540.
- Beggan, J. K. (1992). On the social nature of nonsocial perception: The mere ownership effect. Journal of Personality and Social Psychology 62(2), 229.
- Bénabou, R. and J. Tirole (2016). Mindful economics: The production, consumption, and value of beliefs. Journal of Economic Perspectives 30(3), 141–64.
- Bergus, G., G. Chapman, C. Gjerde, and A. Elstein (1995). Clinical reasoning about new symptoms despite preexisting disease: sources of error and order effects. Family medicine 27(5), 314–320.
- Bernheim, B. D. and A. Rangel (2009). Beyond revealed preference: choice-theoretic foundations for behavioral welfare economics. Quarterly Journal of Economics 124(1), 51–104.
- Blake, T., S. Moshary, K. Sweeney, and S. Tadelis (2018). Price salience and product choice. Marketing Science, forthcoming.
- Block, H. D. and J. Marschak (1960). Random orderings and stochastic theories of response. Contributions To Probability And Statistics. Stanford University Press.
- Bond, S. D., K. A. Carlson, M. G. Meloy, J. E. Russo, and R. J. Tanner (2007). Information distortion in the evaluation of a single option. Organizational Behavior and Human Decision Processes 102(2), 240–254.
- Bordalo, P., N. Gennaioli, and A. Shleifer (2012). Salience theory of choice under risk. Quarterly Journal of Economics 127(3), 1243–1285.
- Bordalo, P., N. Gennaioli, and A. Shleifer (2013a). Salience and asset prices. American Economic Review 103(3), 623–28.

- Bordalo, P., N. Gennaioli, and A. Shleifer (2013b). Salience and consumer choice. Journal of Political Economy 121(5), 803–843.
- Carlson, K. A., M. G. Meloy, and J. E. Russo (2006). Leader-driven primacy: Using attribute order to affect consumer choice. Journal of Consumer Research 32(4), 513–518.
- Carpenter, G. S. and K. Nakamoto (1989). Consumer preference formation and pioneering advantage. Journal of Marketing Research, 285–298.
- Chapman, G. B. and A. S. Elstein (2000). Cognitive processes and biases in medical decision making. Decision making in health care: Theory, psychology, and applications, 183–210.
- Chrzan, K. (1994). Three kinds of order effects in choice-based conjoint analysis. Marketing Letters 5(2), 165–172.
- Cornelissen, J. P. and M. D. Werner (2014). Putting framing in perspective: A review of framing and frame analysis across the management and organizational literature. The Academy of Management Annals 8(1), 181–235.
- Cunnington, J. P., J. M. Turnbull, G. Regher, M. Marriott, and G. R. Norman (1997). The effect of presentation order in clinical decision making. Academic Medicine.
- Dahl, L. C., C. E. Brimacombe, and D. S. Lindsay (2009). Investigating investigators: How presentation order influences participant-investigators’ interpretations of eyewitness identification and alibi evidence. Law and Human Behavior 33(5), 368.
- Day, B., I. J. Bateman, R. T. Carson, D. Dupont, J. J. Louviere, S. Morimoto, R. Scarpa, and P. Wang (2012). Ordering effects and choice set awareness in repeat-response stated preference studies. Journal of Environmental Economics and Management 63(1), 73–91.
- Day, B. and J.-L. P. Prades (2010). Ordering anomalies in choice experiments. Journal of Environmental Economics and Management 59(3), 271–285.
- Debreu, G. (1960). Topological methods in cardinal utility theory. Mathematical Methods in the Social Sciences 1959, 16–26.
- Donohue, W., R. G. Rogan, and S. Kaufman (2011). Framing matters: Perspectives on negotiation research and practice in communication. Peter Lang Publishing.
- Ellison, G. and S. F. Ellison (2009). Search, obfuscation, and price elasticities on the internet. Econometrica 77(2), 427–452.
- Epstein, L. G. (1983). Stationary cardinal utility and optimal growth under uncertainty. Journal of Economic Theory 31(1), 133–152.
- Ericson, K. M. M. and A. Starc (2016). How product standardization affects choice: Evidence from the massachusetts health insurance exchange. Journal of Health Economics 50, 71–85.

- Fudenberg, D., R. Iijima, and T. Strzalecki (2015). Stochastic choice and revealed perturbed utility. Econometrica 83(6), 2371–2409.
- Ge, X., G. Häubl, and T. Elrod (2011). What to say when: influencing consumer choice by delaying the presentation of favorable information. Journal of Consumer Research 38(6), 1004–1021.
- Gorman, W. M. (1980). A possible procedure for analysing quality differentials in the egg market. Review of Economic Studies 47(5), 843–856.
- Gul, F., P. Natenzon, and W. Pesendorfer (2014). Random choice as behavioral optimization. Econometrica 82(5), 1873–1912.
- Gul, F. and W. Pesendorfer (2006). Random expected utility. Econometrica 74(1), 121–146.
- Haugtvedt, C. P. and D. T. Wegener (1994). Message order effects in persuasion: An attitude strength perspective. Journal of Consumer Research 21(1), 205–218.
- Kahneman, D. et al. (1999). Objective happiness. Well-being: The foundations of hedonic psychology 3(25), 1–23.
- Kahneman, D., J. L. Knetsch, and R. H. Thaler (1991). Anomalies: The endowment effect, loss aversion, and status quo bias. Journal of Economic Perspectives 5(1), 193–206.
- Kahneman, D. and A. Tversky (1979). Prospect theory: An analysis of decision under risk. Econometrica, 263–291.
- Kahneman, D., P. P. Wakker, and R. Sarin (1997). Back to bentham? explorations of experienced utility. Quarterly Journal of Economics 112(2), 375–406.
- Kardes, F. R. and P. M. Herr (1990). Order effects in consumer judgment, choice, and memory: The role of initial processing goals. ACR North American Advances.
- Kjær, T., M. Bech, D. Gyrd-Hansen, and K. Hart-Hansen (2006). Ordering effect and price sensitivity in discrete choice experiments: need we worry? Health economics 15(11), 1217–1228.
- Kőszegi, B. and M. Rabin (2006). A model of reference-dependent preferences. Quarterly Journal of Economics 121(4), 1133–1165.
- Kőszegi, B. and A. Szeidl (2012). A model of focusing in economic choice. Quarterly journal of economics, qjs049.
- Krosnick, J. A. and D. F. Alwin (1987). An evaluation of a cognitive theory of response-order effects in survey measurement. Public Opinion Quarterly 51(2), 201–219.
- Kumar, V. and G. J. Gaeth (1991). Attribute order and product familiarity effects in decision tasks using conjoint analysis. International Journal of Research in Marketing 8(2), 113–124.

- Lancaster, K. J. (1966). A new approach to consumer theory. Journal of Political Economy 74(2), 132–157.
- Levav, J., M. Heitmann, A. Herrmann, and S. S. Iyengar (2010). Order in product customization decisions: Evidence from field experiments. Journal of Political Economy 118(2), 274–299.
- Levin, I. P. and G. J. Gaeth (1988). How consumers are affected by the framing of attribute information before and after consuming the product. Journal of Consumer Research 15(3), 374–378.
- Lichtenstein, S. and P. Slovic (2006). The construction of preference. Cambridge University Press.
- Lindsey-Mullikin, J. (2003). Beyond reference price: understanding consumers’ encounters with unexpected prices. Journal of Product & Brand Management.
- List, J. A. (2003). Does market experience eliminate market anomalies? Quarterly Journal of Economics 118(1), 41–71.
- Luce, R. D. (1959). Individual choice behavior: A theoretical analysis. John Wiley.
- Lynch, J. G. J. and D. Ariely (2000). Wine online: Search costs affect competition on price, quality, and distribution. Marketing Science 19(1), 83–103.
- Manzini, P. and M. Mariotti (2014). Stochastic choice and consideration sets. Econometrica 82(3), 1153–1176.
- Marschak, J. (1974). Binary-choice constraints and random utility indicators (1960). In Economic Information, Decision, and Prediction, pp. 218–239. Springer.
- Matějka, F. and A. McKay (2015). Rational inattention to discrete choices: A new foundation for the multinomial logit model. American Economic Review 105(1), 272–98.
- McFadden, D. (1973). Conditional logit analysis of qualitative choice behavior.
- Milosavljevic, M., V. Navalpakkam, C. Koch, and A. Rangel (2012). Relative visual saliency differences induce sizable bias in consumer choice. Journal of Consumer Psychology 22(1), 67–74.
- Nelson, T. E., Z. M. Oxley, and R. A. Clawson (1997). Toward a psychology of framing effects. Political Behavior 19(3), 221–246.
- Ostrizek, F. and D. Shishkin (2018). Screening with frames. Working paper. Princeton University.
- Payne, J. W., D. A. Schkade, W. H. Desvousges, and C. Aultman (2000). Valuation of multiple environmental programs. Journal of Risk and Uncertainty 21(1), 95–115.

- Rubinstein, A. and Y. Salant (2006). A model of choice from lists. Theoretical Economics 1(1), 3–17.
- Rubinstein, A. and Y. Salant (2011). Eliciting welfare preferences from behavioural data sets. Review of Economic Studies 79(1), 375–387.
- Salant, Y. (2011). Procedural analysis of choice rules with applications to bounded rationality. American Economic Review 101(2), 724–48.
- Salant, Y. and A. Rubinstein (2008). (a, f): choice with frames. Review of Economic Studies 75(4), 1287–1296.
- Salant, Y. and R. Siegel (2018). Contracts with framing. American Economic Journal: Microeconomics 10(3), 315–46.
- Shafer, E. F. (2017). Hillary rodham versus hillary clinton: Consequences of surname choice in marriage. Gender Issues 34(4), 316–332.
- Shalvi, S., F. Gino, R. Barkan, and S. Ayal (2015). Self-serving justifications: Doing wrong and feeling moral. Current Directions in Psychological Science 24(2), 125–130.
- Smith, M. D. and E. Brynjolfsson (2001). Consumer decision-making at an internet shopbot: Brand still matters. Journal of Industrial Economics 49(4), 541–558.
- Strahilevitz, M. A. and G. Loewenstein (1998). The effect of ownership history on the valuation of objects. Journal of Consumer Research 25(3), 276–289.
- Thaler, R. (1980). Toward a positive theory of consumer choice. Journal of Economic Behavior & Organization 1(1), 39–60.
- Thaler, R. (1985). Mental accounting and consumer choice. Marketing Science 4(3), 199–214.
- Thaler, R. H. (1990). Anomalies: Saving, fungibility, and mental accounts. The Journal of Economic Perspectives 4(1), 193–205.
- Tirole, J. (1988). The Theory of Industrial Organization. MIT press.
- Tserenjigmid, G. (2021). The order dependent luce model.
- Tversky, A. and D. Kahneman (1981). The framing of decisions and the psychology of choice. Science 211, 30.
- Zhang, Y. and A. Fishbach (2005). The role of anticipated emotions in the endowment effect. Journal of Consumer Psychology 15(4), 316–324.

Supplemental Material (For Online Publication Only)

D Extension of Strategic Framing: Both Firms Choose Frames

We extend the analysis in Section 3.1 by letting the entrant choose the frame under which to present its product. The consumers then compare products using one of the two frames before buying one. Letting f^I and f^E be the incumbent's and entrant's frames, we assume that each consumer uses f^I with probability μ and f^E with probability $1 - \mu$, where $\frac{1}{2} \leq \mu \leq 1$. One interpretation is that each consumer uses the frame of whichever of the two items he or she encounters first (as under hypothesis h_3 in Section 4), and this is more likely to be the incumbent's product. The timing is as follows:

1. The incumbent chooses f^I for its product.
2. The entrant decides whether to enter at cost $K > 0$.
3. If the entrant enters, it chooses f^E for its product.
4. If the entrant enters, the firms compete in prices; otherwise, the incumbent sets its monopoly price.
5. Consumers make their comparisons (if any) and purchase decisions.

For every $f \in F$, recall that $\Delta_f = [\alpha_f(r) - \alpha_f(b)]\delta$ and the consumer who is indifferent between products under each f is characterized by

$$\theta_f^* = \frac{I_p - E_p}{\Delta_f} \alpha_f(p).$$

Clearly, the firms will choose f to try to get the *high end* of the market: $\Delta_{f^I} > 0$ and $\Delta_{f^E} < 0$. Thus, the demand for the incumbent and the entrant will be

$$\begin{aligned} D^I &= \mu \left(\bar{\theta} - \frac{I_p - E_p}{\Delta_{f^I}} \alpha_{f^I}(p) \right) + (1 - \mu) \left(\frac{E_p - I_p}{\Delta_{f^E}} \alpha_{f^E}(p) - \underline{\theta} \right), \\ D^E &= \mu \left(\frac{I_p - E_p}{\Delta_{f^I}} \alpha_{f^I}(p) - \underline{\theta} \right) + (1 - \mu) \left(\bar{\theta} - \frac{E_p - I_p}{\Delta_{f^E}} \alpha_{f^E}(p) \right). \end{aligned}$$

Maximization of profits leads to the following best response functions:

$$\begin{aligned} I_p &= \frac{1}{2} \left[(\mu \bar{\theta} - (1 - \mu) \underline{\theta}) \frac{\Delta_{f^I} \Delta_{f^E}}{\mu \alpha_{f^I}(p) \Delta_{f^E} + (1 - \mu) \alpha_{f^E}(p) \Delta_{f^I}} + E_p \right] \\ E_p &= \frac{1}{2} \left[((1 - \mu) \bar{\theta} - \mu \underline{\theta}) \frac{\Delta_{f^I} \Delta_{f^E}}{\mu \alpha_{f^I}(p) \Delta_{f^E} + (1 - \mu) \alpha_{f^E}(p) \Delta_{f^I}} + I_p \right] \end{aligned}$$

Hence, we obtain the following optimal prices:

$$I_p = \frac{1}{3} \frac{\Delta_{fI} \Delta_{fE}}{\mu \alpha_{fI}(p) \Delta_{fE} + (1 - \mu) \alpha_{fE}(p) \Delta_{fI}} [(1 + \mu) \bar{\theta} - (2 - \mu) \underline{\theta}]$$

$$E_p = \frac{1}{3} \frac{\Delta_{fI} \Delta_{fE}}{\mu \alpha_{fI}(p) \Delta_{fE} + (1 - \mu) \alpha_{fE}(p) \Delta_{fI}} [(2 - \mu) + (1 - 2\mu) \underline{\theta}]$$

As expected, for $\mu = 1$ we obtain the previous optimal prices.

Given these prices, we can rewrite the resulting demands as follows:

$$D^I = \frac{2 - \mu}{3} + \frac{2\mu - 1}{3} \bar{\theta} \quad \text{and} \quad D^E = \frac{2 - \mu}{3} (2\bar{\theta} - 1) - \underline{\theta}$$

Thus, the profits are

$$\varphi^o(I_f) = \frac{1}{9} \frac{\Delta_{fI} \Delta_{fE}}{\mu \alpha_{fI}(p) \Delta_{fE} + (1 - \mu) \alpha_{fE}(p) \Delta_{fI}} [(1 + \mu) \bar{\theta} - (2 - \mu) \underline{\theta}] [(2 - \mu) + (2\mu - 1) \bar{\theta}]$$

$$\varphi^o(E_f) = \frac{1}{9} \frac{\Delta_{fI} \Delta_{fE}}{\mu \alpha_{fI}(p) \Delta_{fE} + (1 - \mu) \alpha_{fE}(p) \Delta_{fI}} [(2 - \mu) + (1 - 2\mu) \underline{\theta}] [(2 - \mu) (2\bar{\theta} - 1) - 3\underline{\theta}].$$

Now note that, when setting f , both firms have the same incentives and will hence choose the same f . This implies that $\Delta_{fI} = -\Delta_{fE} = \Delta$ and $\alpha_{fI}(p) = \alpha_{fE}(p) = \alpha_f(p)$. We can then rewrite the profits as

$$\varphi^o(I_f) = \frac{1}{9} \frac{\Delta}{\alpha_f(p)} [(1 + \mu) \bar{\theta} - (2 - \mu) \underline{\theta}] [(2 - \mu) + (2\mu - 1) \bar{\theta}]$$

$$\varphi^o(E_f) = \frac{1}{9} \frac{\Delta}{\alpha_f(p)} [(2 - \mu) + (1 - 2\mu) \underline{\theta}] [(2 - \mu) (2\bar{\theta} - 1) - 3\underline{\theta}]$$

Thus, given the results of Section 3.1, the incumbent makes higher profits when alone in setting the framing under oligopoly if

$$(2\bar{\theta} - \underline{\theta})^2 \geq [(1 + \mu) \bar{\theta} - (2 - \mu) \underline{\theta}] [(2 - \mu) + (2\mu - 1) \bar{\theta}]$$

This holds if $2 + \underline{\theta} \geq (1 + \mu)(1 + \underline{\theta}) - (2 - \mu) \underline{\theta}$ and $2 + \underline{\theta} \geq (2 - \mu) + (2\mu - 1)(1 + \underline{\theta})$. Since both conditions are satisfied, the incumbent has less market power when the entrant can also influence the frame. This leads to less margins to deter entry, but qualitatively the same analysis of Section 3.1 applies.

E Proof of Theorem 2

The proof proceeds in five steps. We seek to obtain a regular representation of the form

$$w_f(x_f) = \sum_{i=1}^N \alpha(i) u_{f(i)}(x_{f(i)}) \prod_{k=1}^{i-1} B_k(u_{f(k)}(x_{f(k)})), \quad (17)$$

where $\prod_{k=1}^0 B_k(u_{f(k)}(x_{f(k)})) \equiv 1$ and $B_i : \mathcal{U} \rightarrow \mathbb{R}_{++}$ for every $i = 1, \dots, N$ and $\mathcal{U} = \cup_{a \in A} u_a(L_a)$. The representation in Theorem 2 follows from the change of variables

$$\phi_k(u_{f(k)}(x_{f(k)})) = \ln\{B_k(u_{f(k)}(x_{f(k)}))\}.$$

For every (q, M) , let

$$\ell(q, M) = \ln(\pi(q, M)).$$

Step 1. Axiom 6 implies that for all $x^i p_f$, $y^i p_f$, $x^i q_f$, and $y^i q_f$ in $\Delta(X_f)$, we have that

$$\frac{\pi((x_f^i, p_f) | \{(x_f^i, p_f), (x_f^i, q_f)\})}{\pi((x_f^i, q_f) | \{(x_f^i, p_f), (x_f^i, q_f)\})} = \frac{\pi((y_f^i, p_f) | \{(y_f^i, p_f), (y_f^i, q_f)\})}{\pi((y_f^i, q_f) | \{(y_f^i, p_f), (y_f^i, q_f)\})}.$$

Therefore,

$$\ell((x_f^i, p_f) | \{(x_f^i, p_f), (x_f^i, q_f)\}) - \ell((x_f^i, q_f) | \{(x_f^i, p_f), (x_f^i, q_f)\}) = v_f(x_f^i, p_f) - v_f(x_f^i, q_f)$$

is equal to

$$\ell((y_f^i, p_f) | \{(y_f^i, p_f), (y_f^i, q_f)\}) - \ell((y_f^i, q_f) | \{(y_f^i, p_f), (y_f^i, q_f)\}) = v_f(y_f^i, p_f) - v_f(y_f^i, q_f).$$

It follows that

$$v_f(x_f^i, p_f) \geq v_f(x_f^i, q_f) \Leftrightarrow v_f(y_f^i, p_f) \geq v_f(y_f^i, q_f).$$

This implies that $v_f(x_f^i, \cdot)$ and $v_f(y_f^i, \cdot)$ represent the same preference over $\Delta(\times_{k=i+1}^N L_{f(k)})$ for all $x_f^i, y_f^i \in \times_{k=1}^i L_{f(k)}$ and all $i = 1, \dots, N-1$.

To unpack the consequences of this property, consider first $i = 1$ and fix any level $\bar{x}_{f(1)} \in L_{f(1)}$. By the uniqueness properties of vN-M utility representations, there exists $u_{f(1)}^f(x_{f(1)}; \bar{x}_{f(1)}) \in \mathbb{R}$ and $B_{f(1)}^f(x_{f(1)}; \bar{x}_{f(1)}) > 0$ such that, for all $x_{f(1)} \in L_{f(1)}$ and all $p_f \in \Delta(\times_{k=2}^N L_{f(k)})$, we have

$$v_f(x_f^1, p_f) = u_{f(1)}^f(x_{f(1)}; \bar{x}_{f(1)}) + B_{f(1)}^f(x_{f(1)}; \bar{x}_{f(1)})v_f(\bar{x}_{f(1)}, p_f).$$

Therefore, clearly, $u_{f(1)}^f(\bar{x}_{f(1)}; \bar{x}_{f(1)}) = 0$ and $B_{f(1)}^f(\bar{x}_{f(1)}; \bar{x}_{f(1)}) = 1$.

Now consider $i = 2$ and focus on the elements (x_f^2, p_f) with the property that $x_{f(1)} = \bar{x}_{f(1)}$. Using Axiom 6, we conclude that $v_f(\bar{x}_{f(1)}, x_{f(2)}, \cdot)$ and $v_f(\bar{x}_{f(1)}, y_{f(2)}, \cdot)$ represent the same EU preference over $\Delta(\times_{k=3}^N L_{f(k)})$. Fix any level $\bar{x}_{f(2)} \in L_{f(2)}$. By the same uniqueness argument as before, there exists $u_{f(2)}^f(x_{f(2)}; \bar{x}_{f(1)}, \bar{x}_{f(2)}) \in \mathbb{R}$ and $B_{f(2)}^f(x_{f(2)}; \bar{x}_{f(1)}, \bar{x}_{f(2)}) > 0$ such that, for all $x_{f(2)} \in L_{f(2)}$ and all $p_f \in \Delta(\times_{k=3}^N L_{f(k)})$, we have

$$v_f(\bar{x}_{f(1)}, x_{f(2)}, p_f) = u_{f(2)}^f(x_{f(2)}; \bar{x}_{f(1)}, \bar{x}_{f(2)}) + B_{f(2)}^f(x_{f(2)}; \bar{x}_{f(1)}, \bar{x}_{f(2)})v_f(\bar{x}_{f(1)}, \bar{x}_{f(2)}, p_f).$$

If we now replace in the expression for $v_f(x_f^1, p_f)$, we have

$$v_f(x_f^2, p_f) = u_{f(1)}^f(x_{f(1)}; \bar{x}_{f(1)}) + B_{f(1)}^f(x_{f(1)}; \bar{x}_{f(1)}) \left\{ u_{f(2)}^f(x_{f(2)}; \bar{x}_{f(1)}, \bar{x}_{f(2)}) \right.$$

$$\begin{aligned}
& + B_{f(2)}^f(x_{f(2)}; \bar{x}_{f(1)}, \bar{x}_{f(2)})v_f(\bar{x}_{f(1)}, \bar{x}_{f(2)}, p_f) \Big\} \\
& = u_{f(1)}^f(x_{f(1)}; \bar{x}_f^1) + B_{f(1)}^f(x_{f(1)}; \bar{x}_f^1)u_{f(2)}^f(x_{f(2)}; \bar{x}_{f(1)}^2) \\
& \quad + B_{f(1)}^f(x_{f(1)}; \bar{x}_f^1)B_{f(2)}^f(x_{f(2)}; \bar{x}_f^2)v_f(\bar{x}_f^2, p_f).
\end{aligned}$$

Iteratively repeating this argument, we obtain that for all $x_f \in X_f$

$$v_f(x_f) = u_{f(1)}^f(x_{f(1)}; \bar{x}_f^1) + \sum_{k=2}^N u_{f(k)}^f(x_{f(k)}; \bar{x}_f^k) \prod_{j=1}^{k-1} B_{f(j)}^f(x_{f(j)}; \bar{x}_f^j),$$

which becomes after suppressing the dependence on the arbitrary $\bar{x}_f$,

$$v_f(x_f) = u_{f(1)}^f(x_{f(1)}) + \sum_{k=2}^N u_{f(k)}^f(x_{f(k)}) \prod_{j=1}^{k-1} B_{f(j)}^f(x_{f(j)}). \quad (18)$$

Step 2. Now consider Axiom 5. Note that $\pi(x_f|\{x_f, y_f\}) \geq \frac{1}{2}$ is equivalent to

$$\ell(x_f|\{x_f, y_f\}) - \ell(y_f|\{x_f, y_f\}) = v_f(x_f) - v_f(y_f) \geq 0.$$

Using the representation in (18), we have that

$$v_f(x_f) \geq v_f(y_f) \Leftrightarrow u_a^f(x_a) \prod_{j=1}^{i-1} B_{f(j)}^f(x_{f(j)}) \geq u_a^f(y_a) \prod_{j=1}^{i-1} B_{f(j)}^f(y_{f(j)}) \Leftrightarrow u_a^f(x_a) \geq u_a^f(y_a),$$

because $\prod_{j=1}^{i-1} B_{f(j)}^f(x_{f(j)}) = \prod_{j=1}^{i-1} B_{f(j)}^f(y_{f(j)}) > 0$.

By similar reasoning, the axiom says that $v_{f'}(\hat{x}_{f'}) \geq v_{f'}(\hat{y}_{f'})$ for every specification of $\hat{x}_{f'(k)} = \hat{y}_{f'(k)}$. Also, recall that there exists $\bar{x}_{f'} \in X_{f'}$ such that $u_{f'(k)}^{f'}(\bar{x}_{f'(k)}) = 0$ for all k . Therefore, letting $\hat{x}_{f'(k)} = \hat{y}_{f'(k)} = \bar{x}_{f'(k)}$ for all $k \neq i$, we have

$$\begin{aligned}
v_{f'}(\hat{x}_{f'}) \geq v_{f'}(\hat{y}_{f'}) & \Leftrightarrow u_a^{f'}(x_a) \prod_{j=1}^{i-1} B_{f'(j)}^{f'}(\bar{x}_{f'(j)}) \geq u_a^{f'}(y_a) \prod_{j=1}^{i-1} B_{f'(j)}^{f'}(\bar{x}_{f'(j)}) \\
& \Leftrightarrow u_a^{f'}(x_a) \geq u_a^{f'}(y_a),
\end{aligned}$$

because $\prod_{j=1}^{i-1} B_{f'(j)}^{f'}(\bar{x}_{f'(j)}) > 0$. We conclude that u_a^f and $u_a^{f'}$ represent the same ranking over L_a . In particular, this means that $u_a^f(x_a) = u_a^f(y_a)$ if and only if $u_a^{f'}(x_a) = u_a^{f'}(y_a)$.

Towards our goal of showing that each $B_{f(k)}^f$ depends on $x_{f(k)}$ only via $u_{f(k)}^f(x_{f(k)})$, consider first the simple case where $u_{f(k)}^f(x_{f(k)}) \neq u_{f(k)}^f(y_{f(k)})$ for all $x_{f(k)}, y_{f(k)} \in L_{f(k)}$ (i.e., $u_{f(k)}^f$ is injective). Then, we can just re-define

$$\hat{B}_{f(k)}^f(u_{f(k)}^f(x_{f(k)})) = B_{f(k)}^f((u_{f(k)}^f)^{-1}(u_{f(k)}^f(x_{f(k)}))).$$

Thus, the desired property of the weights holds trivially.

Now consider the less immediate case where $u_{f(k)}^f(x_{f(k)}) = u_{f(k)}^f(y_{f(k)})$ for some $x_{f(k)}, y_{f(k)} \in L_{f(k)}$, where $f(k) = a \in A$ and $k < N$. Then, by the previous argument, for every frame f' with $f'(N) = a$, we must have $u_a^{f'}(x_{f(k)}) = u_a^{f'}(y_{f(k)})$ and so

$$v_{f'}(\hat{x}_{f'(1)}, \dots, \hat{x}_{f'(N-1)}, x_{f(k)}) = v_{f'}(\hat{x}_{f'(1)}, \dots, \hat{x}_{f'(N-1)}, y_{f(k)}).$$

By the axiom, we must also have

$$v_f(\bar{x}_{f(1)}, \dots, \bar{x}_{f(k-1)}, x_{f(k)}, z_{f(k+1)}, \dots, z_{f(N)}) = v_f(\bar{x}_{f(1)}, \dots, \bar{x}_{f(k-1)}, y_{f(k)}, z_{f(k+1)}, \dots, z_{f(N)}),$$

for every $(z_{f(k+1)}, \dots, z_{f(N)}) \in \times_{j=k+1}^N L_{f(j)}$. Therefore, after simplifying the term $\prod_{j=1}^{k-1} B_{f(j)}^f(\bar{x}_{f(j)}) > 0$ and using $u_{f(j)}^f(\bar{x}_{f(j)}) = 0$ for $j < k$, we have

$$\begin{aligned} 0 &= u_{f(k)}^f(x_{f(k)}) - u_{f(k)}^f(y_{f(k)}) \\ &\quad + \left[B_{f(k)}^f(x_{f(k)}) - B_{f(k)}^f(y_{f(k)}) \right] \sum_{j=k+1}^N u_{f(j)}^f(z_{f(j)}) \prod_{i=k+1}^{j-1} B_{f(i)}^f(z_{f(i)}) \\ &= \left[B_{f(k)}^f(x_{f(k)}) - B_{f(k)}^f(y_{f(k)}) \right] \sum_{j=k+1}^N u_{f(j)}^f(z_{f(j)}) \prod_{i=k+1}^{j-1} B_{f(i)}^f(z_{f(i)}), \end{aligned}$$

where $\prod_{i=k+1}^k B_{f(i)}^f(z_{f(i)}) \equiv 1$. Since $u_{f(j)}^f(z_{f(j)}) \neq 0$ for some $z_{f(j)} \in L_{f(j)}$ and some $j \geq k+1$, we must have

$$u_{f(k)}^f(x_{f(k)}) = u_{f(k)}^f(y_{f(k)}) \quad \Rightarrow \quad B_{f(k)}^f(x_{f(k)}) = B_{f(k)}^f(y_{f(k)}).$$

This implies that $B_{f(k)}^f$ cannot depend on $x_{f(k)}$ other than through $u_{f(k)}^f(x_{f(k)})$, as desired.

To recap, we now have the following representation: For all $f \in F$ and $x_f \in X_f$,

$$v_f(x_f) = u_{f(1)}^f(x_{f(1)}) + \sum_{j=2}^N u_{f(j)}^f(x_{f(j)}) \prod_{k=1}^{j-1} B_{f(k)}^f(u_{f(k)}^f(x_{f(k)})).$$

We concluded earlier that u_a^f for $f(N) = a$ and $u_a^{f'}$ for any other $f' \in F$ —where $f'(N)$ may be different from a —represent the same ranking over L_a . By Axiom 6, u_a^f is also a vN-M utility function over L_a . Therefore, there exists $\gamma_a^{f'} > 0$ and $\zeta_a^{f'} \in \mathbb{R}$ such that, for every f' different from a fixed f^* with $f^*(N) = a$, we must have

$$u_a^{f'} = \gamma_a^{f'} u_a + \zeta_a^{f'},$$

where we define $u_a = u_a^{f^*}$. Note that this implies that without loss of generality each $B_{f(k)}^f$ function depends on $x_{f(k)}$ only through $u_{f(k)}$: We can simply define

$$\hat{B}_{f(k)}^f(u_{f(k)}) = B_{f(k)}^f(\gamma_{f(k)}^f u_{f(k)} + \zeta_{f(k)}^f).$$

Therefore, (simplifying notation) we have

$$\begin{aligned}
v_f(x_f) &= \gamma_{f(1)}^f u_{f(1)}(x_{f(1)}) + \zeta_{f(1)}^f + \sum_{j=2}^N [\gamma_{f(j)}^f u_{f(j)}(x_{f(j)}) + \zeta_{f(j)}^f] \prod_{k=1}^{j-1} B_{f(k)}^f(u_{f(k)}(x_{f(k)})) \\
&= \gamma_{f(1)}^f u_{f(1)}(x_{f(1)}) + \zeta_{f(1)}^f + \sum_{j=2}^N \gamma_{f(j)}^f u_{f(j)}(x_{f(j)}) \prod_{k=1}^{j-1} B_{f(k)}^f(u_{f(k)}(x_{f(k)})) \\
&\quad + \sum_{j=2}^N \zeta_{f(j)}^f \prod_{k=1}^{j-1} B_{f(k)}^f(u_{f(k)}(x_{f(k)})).
\end{aligned}$$

Step 3. We now would like to show that each $B_{f(k)}^f$ depends only on the position k for all $f \in F$. To this end, we use Axiom 8, which implies the following. First, note that

$$\ell(x_f, \{x_f, \hat{x}_f\}) - \ell(\hat{x}_f, \{x_f, \hat{x}_f\}) = [u_{f(i)}^f(x_{f(i)}) - u_{f(i)}^f(\hat{x}_{f(i)})] \prod_{k=1}^{i-1} B_{f(k)}^f(u_{f(k)}(x_{f(k)}))$$

and

$$\ell(y_{f'}, \{y_{f'}, \hat{y}_{f'}\}) - \ell(\hat{y}_{f'}, \{y_{f'}, \hat{y}_{f'}\}) = u_{f'(1)}^{f'}(x_{f(i)}) - u_{f'(1)}^{f'}(\hat{x}_{f(i)}).$$

Using this, we have

$$\ell(x_f, \{x_f, \hat{x}_f\}) - \ell(\hat{x}_f, \{x_f, \hat{x}_f\}) = [u_{f(i)}(x_{f(i)}) - u_{f(i)}(\hat{x}_{f(i)})] \prod_{k=1}^{i-1} \gamma_{f(i)}^f B_{f(k)}^f(u_{f(k)}(x_{f(k)}))$$

and

$$\ell(y_{f'}, \{y_{f'}, \hat{y}_{f'}\}) - \ell(\hat{y}_{f'}, \{y_{f'}, \hat{y}_{f'}\}) = [u_{f'(1)}(x_{f(i)}) - u_{f'(1)}(\hat{x}_{f(i)})] \gamma_{f(1)}^{f'}.$$

The axiom requires that

$$\frac{\prod_{k=1}^{i-1} \gamma_{f(i)}^f B_{f(k)}^f(u_{f(k)}(x_{f(k)}))}{\gamma_{f(1)}^{f'}} = r(i, x_{f(1)}, \dots, x_{f(i-1)}).$$

This implies that $\gamma_{f(1)}^{f'} = \gamma_1$ for all $f, f' \in F$ and some $\gamma_1 > 0$, $\gamma_{f(i)}^f = \gamma_i$ for all $f \in F$ and some $\gamma_i > 0$, and $B_{f(k)}^f(u_{f(k)}(x_{f(k)})) = B_k(u_{f(k)}(x_{f(k)}))$ for all $f \in F$ and some real number $B_k(u_{f(k)}(x_{f(k)})) > 0$. Thus, we can define $\mathcal{U} = \cup_{a \in A} u_a(L_a)$ and the function $B_k : \mathcal{U} \rightarrow \mathbb{R}_{++}$ as taking the values just defined.

These steps refine the representation of v_f to the following:

$$\begin{aligned}
v_f(x_f) &= \gamma_1 u_{f(1)}(x_{f(1)}) + \zeta_{f(1)}^f + \sum_{j=2}^N \gamma_j u_{f(j)}(x_{f(j)}) \prod_{k=1}^{j-1} B_k(u_{f(k)}(x_{f(k)})) \\
&\quad + \sum_{j=2}^N \zeta_{f(j)}^f \prod_{k=1}^{j-1} B_k(u_{f(k)}(x_{f(k)})).
\end{aligned}$$

Step 4. By the uniqueness of v_f as a Luce value up to adding constants, we can set $\zeta_{f(1)}^f = 0$ for every f without loss. We would like to also show that $\zeta_{f(j)}^f = 0$ for every f and $j > 1$. To this end, we exploit Axiom 7 to further refine the representation as follows. For $i = 2$, its conclusion is equivalent to the equality between

$$\begin{aligned}\ell(x_f|\{x_f, y_f\}) - \ell(y_f|\{x_f, y_f\}) &= v_f(x_f) - v_f(y_f) \\ &= \gamma_1[u_{f(1)}(x_{f(1)}) - u_{f(1)}(y_{f(1)})] \\ &\quad + \zeta_{f(2)}^f [B_1(u_{f(1)}(x_{f(1)})) - B_1(u_{f(1)}(y_{f(1)}))] \end{aligned}$$

and

$$\begin{aligned}\ell(x'_{f'}|\{x'_{f'}, y'_{f'}\}) - \ell(y'_{f'}|\{x'_{f'}, y'_{f'}\}) &= v_{f'}(x'_{f'}) - v_{f'}(y'_{f'}) \\ &= \gamma_1[u_{f(1)}(x_{f(1)}) - u_{f(1)}(y_{f(1)})] \\ &\quad + \zeta_{f'(2)}^{f'} [B_1(u_{f(1)}(x_{f(1)})) - B_1(u_{f(1)}(y_{f(1)}))] . \end{aligned}$$

This implies that

$$[\zeta_{f(2)}^f - \zeta_{f'(2)}^{f'}] [B_1(u_{f(1)}(x_{f(1)})) - B_1(u_{f(1)}(y_{f(1)}))] = 0.$$

Since B_1 is not constant, we must have $\zeta_{f(2)}^f = \zeta_{f'(2)}^{f'} = \zeta_{f(1)}$ for all f and f' that satisfy $f(1) = f'(1)$. Now suppose that, for all $k = 2, \dots, j$, we have $\zeta_{f(k)}^f = \zeta_{f'(k)}^{f'} = \zeta_{f(1), \dots, f(k-1)}$ for all f and f' that satisfy $f(m) = f'(m)$ for $m \leq k-1$. Let $i = j+1$ in Axiom 7. Its conclusion is equivalent to the equality between

$$\begin{aligned}&\ell(x_f|\{x_f, y_f\}) - \ell(y_f|\{x_f, y_f\}) \\ &= v_f(x_f) - v_f(y_f) \\ &= \sum_{k=1}^j \gamma_k u_{f(k)}(x_{f(k)}) \prod_{m=1}^{k-1} B_m(u_{f(m)}(x_{f(m)})) \\ &\quad - \sum_{k=1}^j \gamma_k u_{f(k)}(y_{f(k)}) \prod_{m=1}^{k-1} B_m(u_{f(m)}(y_{f(m)})) \\ &\quad + \sum_{k=2}^j \zeta_{f(1), \dots, f(k-1)} \left\{ \prod_{m=1}^{k-1} B_m(u_{f(m)}(x_{f(m)})) - \prod_{m=1}^{k-1} B_m(u_{f(m)}(y_{f(m)})) \right\} \\ &\quad + \zeta_{f(j+1)}^f \left\{ \prod_{m=1}^j B_m(u_{f(m)}(x_{f(m)})) - \prod_{m=1}^j B_m(u_{f(m)}(y_{f(m)})) \right\} \end{aligned}$$

and

$$\begin{aligned}&\ell(x'_{f'}|\{x'_{f'}, y'_{f'}\}) - \ell(y'_{f'}|\{x'_{f'}, y'_{f'}\}) \\ &= v_{f'}(x'_{f'}) - v_{f'}(y'_{f'}) \\ &= \sum_{k=1}^j \gamma_k u_{f(k)}(x_{f(k)}) \prod_{m=1}^{k-1} B_m(u_{f(m)}(x_{f(m)})) \end{aligned}$$

$$\begin{aligned}
& - \sum_{k=1}^j \gamma_k u_{f(k)}(y_{f(k)}) \prod_{m=1}^{k-1} B_m(u_{f(m)}(y_{f(m)})) \\
& + \sum_{k=2}^j \zeta_{f(1), \dots, f(k-1)} \left\{ \prod_{m=1}^{k-1} B_m(u_{f(m)}(x_{f(m)})) - \prod_{m=1}^{k-1} B_m(u_{f(m)}(y_{f(m)})) \right\} \\
& + \zeta_{f'(j+1)}^{f'} \left\{ \prod_{m=1}^j B_m(u_{f(m)}(x_{f(m)})) - \prod_{m=1}^j B_m(u_{f(m)}(y_{f(m)})) \right\}.
\end{aligned}$$

This implies that

$$\left[\zeta_{f(j+1)}^f - \zeta_{f'(j+1)}^{f'} \right] \left\{ \prod_{m=1}^j B_m(u_{f(m)}(x_{f(m)})) - \prod_{m=1}^j B_m(u_{f(m)}(y_{f(m)})) \right\} = 0.$$

Since the quantity in brackets is again not constant, we must have $\zeta_{f(j+1)}^f = \zeta_{f'(j+1)}^{f'} = \zeta_{f(1), \dots, f(j)}$ for all f and f' that satisfy $f(m) = f'(m)$ for $m \leq j$. By induction, we can extend this property to every $j = 2, \dots, N$.

Now consider any f and any x_f that satisfies $x_{f(i)} = \bar{x}_{f(i)}$ for all $i \geq 2$, so that $\gamma_i u_{f(i)}(x_{f(i)}) + \zeta_{f(1), \dots, f(i-1)} = 0$. In this case, we have that $v_f(x_{f(1)}, \bar{x}_{f(-1)})$ is a vN-M utility function over $L_{f(1)}$ and it takes the form

$$v_f(x_{f(1)}, \bar{x}_{f(-1)}) = \gamma_1 u_{f(1)}(x_{f(1)}) + \zeta_{f(1)} B_1(u_{f(1)}(x_{f(1)})).$$

Since $u_{f(1)}$ is also a vN-M utility function over $L_{f(1)}$, we must have

$$v_f(x_{f(1)}, \bar{x}_{f(-1)}) = \hat{\gamma}_{f(1)} u_{f(1)}(x_{f(1)}) + \hat{\zeta}_{f(1)}.$$

This implies that

$$[\hat{\gamma}_{f(1)} - \gamma_1] u_{f(1)}(x_{f(1)}) + \hat{\zeta}_{f(1)} = \zeta_{f(1)} B_1(u_{f(1)}(x_{f(1)})).$$

There are several cases to consider, which all yield $\zeta_{f(1)} = 0$. First, if $\hat{\gamma}_{f(1)} = \gamma_1$, then the equality can hold if and only if $\hat{\zeta}_{f(1)} = \zeta_{f(1)} = 0$ because B_1 is not constant. Given this, suppose that $\hat{\gamma}_{f(1)} > \gamma_1$ without loss of generality. If B_1 is not a linear function of $u_{f(1)}$, the equality can only hold if and only if $\zeta_{f(1)} = 0$. Finally, suppose that B_1 is indeed linear in $u_{f(1)}$, that is, there exist $\bar{\gamma}_1 > 0$ and $\bar{\zeta}_1 \in \mathbb{R}$ such that

$$B_1(u_{f(1)}) = \bar{\gamma}_1 u_{f(1)} + \bar{\zeta}_1.$$

In this case, we have

$$\begin{aligned}
v_f(x_{f(1)}, \bar{x}_{f(-1)}) &= \gamma_1 u_{f(1)}(x_{f(1)}) + \zeta_{f(1)} [\bar{\gamma}_1 u_{f(1)}(x_{f(1)}) + \bar{\zeta}_1] \\
&= [\gamma_1 + \zeta_{f(1)} \bar{\gamma}_1] u_{f(1)}(x_{f(1)}) + \zeta_{f(1)} \bar{\zeta}_1.
\end{aligned}$$

By the uniqueness properties of any Luce value function, we can let $\bar{\zeta}_1 = 0$ without loss of generality. Finally, by Axiom 8, $\gamma_1 + \zeta_{f(1)} \bar{\gamma}_1$ cannot depend on $f(1)$ but only on the position $i = 1$. Thus, without loss of generality, we can let $\zeta_{f(1)} = 0$ and adjust γ_1 accordingly.

Now, suppose we established that $\zeta_{f(1), \dots, f(j-1)} = 0$ for all $j = 1, \dots, i-1$. Consider any f and any x_f that satisfies $x_{f(j)} = \bar{x}_{f(j)}$ for all $j \neq i$, so that $\gamma_j u_{f(j)}(x_{f(j)}) + \zeta_{f(1), \dots, f(j-1)} = 0$. Again, we have that $v_f(x_{f(i)}, \bar{x}_{f(-i)})$ is a vN-M utility function over $L_{f(i)}$ and it takes the form

$$v_f(x_{f(i)}, \bar{x}_{f(-i)}) = [\gamma_i u_{f(i)}(x_{f(i)}) + \zeta_{f(1), \dots, f(i-1)} B_i(u_{f(i)}(x_{f(i)}))] \bar{B}_i,$$

where $\bar{B}_i = \prod_{k=1}^{i-1} B_k(u_{f(k)}(\bar{x}_{f(k)})) > 0$. Since $u_{f(i)}$ is also a vN-M utility function over $L_{f(i)}$, we must have

$$v_f(x_{f(i)}, \bar{x}_{f(-i)}) = \hat{\gamma}_{f(i)} u_{f(i)}(x_{f(i)}) + \hat{\zeta}_{f(i)}.$$

This implies that

$$[\hat{\gamma}_{f(i)} - \gamma_i \bar{B}_i] u_{f(i)}(x_{f(i)}) + \hat{\zeta}_{f(i)} = \zeta_{f(1), \dots, f(i-1)} B_i(u_{f(i)}(x_{f(i)})).$$

Repeating the previous reasoning, we again conclude that $\zeta_{f(1), \dots, f(i-1)} = 0$ without loss of generality.

Step 5. Combining all these steps, we obtain the representation

$$w_f(x_f) = \sum_{i=1}^N \alpha(i) u_{f(i)}(x_{f(i)}) \prod_{k=1}^{i-1} B_k(u_{f(k)}(x_{f(k)})),$$

where the function α is defined by $\alpha(i) = \gamma_i$ for all $i = 1, \dots, N$.

We conclude by showing that (α, u, ϕ) must satisfy regularity. Using the representation, for any x'_f and y'_f such that $x_{f'(-N)} = y_{f'(-N)}$ we have that $\pi(x'_f | \{x'_f, y'_f\}) \geq \frac{1}{2}$ if and only if $u_{f'(N)}(x_{f'(N)}) \geq u_{f'(N)}(y_{f'(N)})$. By Axiom 5, we must have $\pi(x_f | \{x_f, y_f\}) \geq \frac{1}{2}$ for any f that presents attribute $f'(N)$ in any position $i < N$. This holds if $u_{f(i)}(x_{f(i)}) = u_{f(i)}(y_{f(i)})$, so hereafter assume that $u_{f(i)}(x_{f(i)}) > u_{f(i)}(y_{f(i)})$. In this case, $\pi(x_f | \{x_f, y_f\}) \geq \frac{1}{2}$ if and only if

$$\sum_{j=i}^N u_{f(j)}(x_{f(j)}) \alpha(j) \prod_{k=1}^{j-1} B_k(u_{f(k)}(x_{f(k)})) \geq \sum_{j=i}^N u_{f(j)}(y_{f(j)}) \alpha(j) \prod_{k=1}^{j-1} B_k(u_{f(k)}(y_{f(k)})).$$

Since $\prod_{k=1}^{i-1} B_k(u_{f(k)}(x_{f(k)})) = \prod_{k=1}^{i-1} B_k(u_{f(k)}(y_{f(k)})) > 0$, the last condition is equivalent to

$$\begin{aligned} & u_{f(i)}(x_{f(i)}) \alpha(i) + \sum_{j=i+1}^N u_{f(j)}(x_{f(j)}) \alpha(j) \prod_{k=i}^{j-1} B_k(u_{f(k)}(x_{f(k)})) \\ & \geq u_{f(i)}(y_{f(i)}) \alpha(i) + \sum_{j=i+1}^N u_{f(j)}(y_{f(j)}) \alpha(j) \prod_{k=i}^{j-1} B_k(u_{f(k)}(y_{f(k)})). \end{aligned}$$

Using now $\prod_{k=i+1}^{j-1} B_k(u_{f(k)}(x_{f(k)})) = \prod_{k=i+1}^{j-1} B_k(u_{f(k)}(y_{f(k)}))$ and $u_{f(i)}(x_{f(i)}) > u_{f(i)}(y_{f(i)})$, we obtain

$$\alpha(i)$$

$$\geq -\frac{B_i(u_{f(i)}(x_{f(i)})) - B_i(u_{f(i)}(y_{f(i)}))}{u_{f(i)}(x_{f(i)}) - u_{f(i)}(y_{f(i)})} \sum_{j=i+1}^N u_{f(j)}(x_{f(j)}) \alpha(j) \prod_{k=i+1}^{j-1} B_k(u_{f(k)}(x_{f(k)})).$$

Note that this same condition is required if we started with $u_{f(i)}(x_{f(i)}) < u_{f(i)}(y_{f(i)})$. Since this has to hold for all $f \in F$, $x_f \in X_f$, and $y_{f(i)}$, it is equivalent to the condition in Definition 9.