

Imperative-or-declarative and practical triviality¹

Yichi (Raven) ZHANG — *Heinrich Heine University*

Abstract. We draw attention to a novel puzzle concerning imperative-or-declarative disjunctions. To solve this puzzle, we introduce a notion of practical triviality: a sentence is practically trivial if it produces no effect on the addressee’s preference structure. We spell out this analysis formally using a new formulation of preference semantics (e.g., Starr, 2020) according to which preference relations are defined over worlds rather than over propositions.

Keywords: imperative, disjunction, triviality, information oddness, preference semantics.

1. Introduction

Our point of departure is a puzzle concerning imperative-or-declarative disjunctions (cf. Clark, 1993; Schwager, 2004; Franke, 2005; Kaufmann, 2012):

- (1) a. Bet heads or you will lose. (So do bet heads!)
- b. #Bet heads or you will win. (So don’t bet heads!)

Following Schwager (2004), we shall refer to such constructions as IoDs. As (1) shows, IoDs can only be used positively with an undesirable second disjunct to incentivize the addressee to carry out the action dictated by the first disjunct, but not negatively with a desirable second disjunct to deter the addressee from performing the said action. By “(un)desirable”, we mean “it is common ground that the outcome is (un)desirable”. In (1), if it is known that the addressee, for some reason, prefers losing over winning, then our judgment about felicity is reversed: (1a) becomes infelicitous while (1b) becomes felicitous.

In this article, we draw attention to another aspect of this puzzle that has seldom been discussed. Not all positive IoDs are equally felicitous: a positive IoD $(!p \vee \triangleright q)$ ² is defective when it is common ground that p yields a less desirable outcome for the addressee than q , as shown by the following pair in an imaginary carjacking scenario:

- (2) a. Hand over your wallet or I’ll take your car. (So give me your wallet!)
- b. #Hand over your car or I’ll take your wallet. (So give me your car!)

When it is common ground that the addressee values her car more than her wallet—losing her car is less desirable than the carjacker’s taking her wallet—(2a) effectively conveys a conditional threat, whereas (2b) sounds bizarre. Furthermore, this empirical generalization, that $(!p \vee \triangleright q)$ is defective when it is common ground that p yields a less desirable outcome for the addressee than q , also covers IoDs like (1). In (1b), the relative desirability that is being compared is between winning and the mere action of betting heads. Assume it is common ground that the mere action of betting heads is less desirable than winning; (1b) is defective. By contrast in (1a), since the mere action of betting heads is preferred over losing, (1a) is unmarked.

¹I would like to thank the semantics and pragmatics group at Heinrich Heine University and the audience at SuB30 for helpful comments and discussion.

²Following Starr (2020), we use $!$ and $\triangleright$ to represent imperative and declarative clauses respectively.

We will show how this generalization can be derived from general pragmatic considerations regarding when imperatives (as well as IoDs) are trivial. In short, an imperative is trivial if it has no effect on the addressee's preferences. Formally, we will couch our analysis in a dynamic system that is inspired by preference semantics (Starr, 2020; Murray and Starr, 2020) but with substantial modifications. The rest of this paper is organized as follows. In section 2, we further elaborate on the puzzle of IoDs and provide an informal sketch of the present proposal. In section 3, we introduce preference semantics and highlight a problem which concerns its definition of logical consequence, more specifically, its failure to derive contextual equivalence in the case of *imperative Hurford disjunctions*. In section 4, we provide our formal analysis and show how it captures both IoDs and imperative Hurford disjunctions. Section 5 addresses some loose ends in regard to Ross's paradox and conditional imperatives.

2. Puzzle of IoDs

2.1. Directive force of the imperative disjunct

Most prior analyses of the contrast in (1) postulate that the imperative disjunct retains its usual directive force, the force that it would possess if the imperative appeared in isolation. This is realized either through reinterpretation of the disjunction as a conjunction of modals (e.g., Franke, 2005; Kaufmann, 2012) or by attributing additional salience to the imperative disjunct via other means (e.g., van Rooij and Franke, 2010; Biezma and Rawlins, 2016). This, in combination with the additional *or-to-if* inference from the disjunction ($!p \vee \triangleright q$) to the conditional ($p \rightarrow q$), results in a pragmatic conflict in the case of infelicitous IoDs. For instance, (1b), repeated below, can be decomposed into two parts:

- (3) #Bet heads or you will win.
 a. Bet heads!
 b. If you don't bet heads then you will win.

The first disjunct retains its usual directive force and commands the addressee to bet heads. At the same time, the conditional in (3b) incentivizes the addressee not to bet heads, which is at odds with the direction issued by the first disjunct.

However, we find the claim that the imperative disjunct in an IoD always carries its usual directive force disputable. For one, as Starr (2020) observes, there are IoDs in which the imperative disjunct exerts a notably weaker force, compared to when it appears in isolation, as exemplified by (4).

- (4) *We're at a used book sale to stock our joint library. We've each found three books, but we only have enough cash for five books. One of us has to put a book back. I say:*
 Put back Waverly or I'll put back Naked Lunch. I don't care which.

Secondly, this hypothesis faces difficulties when it comes to disjunctions that combine an imperative first disjunct with an interrogative second disjunct. For example, consider:

- (5) *We're trying to figure out whom to invite to our wedding. I say:*
 Invite Ann and her husband, or are they no longer married?

As in (4), the first disjunct in (5) does not seem to retain its usual directive force, since this would presuppose that Ann is still married, which would then clash with the interrogative disjunct by rendering the question trivial. Instead, the speaker only appears to suggest inviting Ann and her husband provided that they are still married.

Moreover, it is not immediately obvious how this analysis based on directive force extends to IoDs like (2). To elaborate, (2b), repeated below, conveys two things on this analysis:

- (6) #Hand over your car or I'll take your wallet.
 a. Hand over your car!
 b. If you don't hand over your car then I'll take your wallet.

However, unlike in (3), since losing one's wallet is an undesirable outcome, (6b) does not incentivize the addressee to carry out the action dictated by the conditional's antecedent and *not* hand over her car. Consequently, (6b) should be compatible with (6a), and the following contrast seems to support this. Compare:

- (7) a. #Bet heads! But if you don't bet heads then you will win.
 b. Hand over your car! But if you don't hand over your car then I'll take your wallet (instead).

Conjoining (6a) and (6b) as in (7b) does not give rise to the same clash as observed in (7a). Hence, appealing to the alleged directive force of the first disjunct alone is inadequate.

2.2. Practical triviality: An informal sketch

To explain the contrast in felicity found in (1) and (2), we introduce a notion of practical triviality. To appreciate how triviality comes into play, let us first consider a similar contrast in (8).

- (8) a. Either George is in love, or I'm a monkey's uncle. (Simons, 2001)
 b. #Either George is in love, or I'm a human being.

While (8a) can be used to convey the speaker's high confidence in George's being in love, (8b) is outright odd. The explanation seems straightforward here. Since "I'm a monkey's uncle" is inconsistent with the common ground, (8a) is contextually equivalent to "George is in love", which is informative. By contrast in (8b), since the common ground already entails the second disjunct "I'm a human being", the disjunction as a whole ends up being entailed by the common ground as well. It follows that an assertion of (8b) fails to update the common ground; hence, it is *truth-conditionally trivial*.

We submit that the oddness observed in (1) and (2) is also due to triviality, albeit of a different kind. Loosely speaking, we take it that at the semantics-pragmatics interface, imperatives primarily function to update the addressee's preferences so as to influence their future actions. An imperative is then considered *practically trivial* if it has no effect on the addressee's preference structure. An utterance is infelicitous if it is both truth-conditionally and practically trivial.³

³More comprehensively, we should also include "inquisitive triviality" in the antecedent of this condition to accommodate questions and other utterances that embody inquisitive content (cf. Ciardelli et al., 2018; Zhang, 2025).

In (1b) “bet heads or you will win”, it is common ground that the addressee prefers winning, so when presented with a choice between betting heads and winning, the addressee will choose the latter. However, since “you will win” is a declarative that, when used non-performatively as in this case, does not change the addressee’s preference structure, (1b) as a whole is practically trivial. On top of this, since “bet heads”, being an imperative, does not add any factual information, the whole disjunction is truth-conditionally trivial too. Hence, (1b) is judged infelicitous.

By contrast in (1a) “bet heads or you will lose”, since losing is dispreferred, the addressee will choose the first disjunct, which leads to an update on her preference structure by promoting “betting heads” to become the preferred alternative, thereby making (1a) non-trivial.

The same explanation extends to the contrast in (2). The following background preference can be assumed from the most desirable to the least: *keep one’s car* (c) $>$ *keep one’s wallet* (w) $>$ *lose one’s wallet* ($\neg w$) $>$ *lose one’s car* ($\neg c$). For $!\neg c \vee \triangleright \neg w$ in (2b), since $\neg w$ is preferred to $\neg c$, the addressee will simply opt for $\neg w$. But since “I will take your wallet” is a declarative and has no effect on the addressee’s preference structure when used non-performatively, (2b) is judged infelicitous. Our analysis further predicts that when $\triangleright \neg w$ in the second disjunct is used performatively, the disjunction, which now produces the same discourse effect as $!\neg c \vee !\neg w$, should become felicitous, because the second disjunct which corresponds to the preferred alternative does alter the addressee’s preference structure. This prediction is borne out by (9).

(9) Hand over your car or you will give me your wallet!

As for $!\neg w \vee \triangleright \neg c$ in (2a), since the addressee wants to avoid $\neg c$ at all cost, she will go with the imperative disjunct $!\neg w$. Because this is an imperative, it will command the addressee to adjust her prior preference ordering by promoting $\neg w$ above w . Thus, (2a) is not practically trivial.

The present analysis is also able to explain why the imperative disjunct seems to retain its directive force in certain IoDs such as (1) and (2) but not in (4). In the case of “put back Waverly or I’ll put back Naked Lunch”, since there is no background preference that the addressee prefers one alternative over the other, both alternatives remain as viable options for devising some joint plans (cf. Roberts, 2023). When presented with a choice, the addressee will not automatically choose one over the other. As a result, the imperative disjunct does not immediately win, and its directive force is suppressed. Instead of postulating that the imperative disjunct in an IoD always retains its directive force, our analysis is able to readily derive the directive force in certain restricted circumstances through general pragmatic reasoning.

3. Modeling preference structure

In order to represent the addressee’s preference structure formally, we draw on preference semantics (Starr, 2018, 2020; Murray and Starr, 2020). At the same time, our proposal deviates in a few key aspects from existing preference semantics. First, we will introduce two types of preference: foreground preference which is directly modified by imperatives, and the addressee’s background preference. Second, whereas existing preference semantics defines preference relations over propositions, we define them over possible worlds, which is more standard in dynamic epistemic logic (cf. van Benthem and Liu, 2007; Liu, 2011). We then supply a novel notion of logical consequence. In this section, we will first provide a brief overview of extant preference

semantics and then highlight several problems so as to motivate the aforementioned changes.

3.1. Overview of preference semantics

Preference semantics is a type of dynamic update semantics. It construes meanings of sentences as their context change potentials, namely as functions from input to output contexts. Contexts are understood as preference states $\mathcal{R}$ which are sets of preference relations r . A preference relation is a set of pairs of propositions: $r = \{\langle p, p' \rangle \mid p \text{ is preferred to } p'\}$. Given a set of worlds W representing the total logical space, the initial minimal preference state is represented by $\{r_0\}$ with $r_0 = \{\langle W, \emptyset \rangle\}$; it contains a single preference relation r_0 where the tautological proposition W is preferred to the absurd proposition $\emptyset$.

Declaratives update preference states in a point-wise fashion: as shown in (10), an update on $\mathcal{R}$ with $[\triangleright\varphi]$ intersects every proposition in every $r \in \mathcal{R}$ with the set of worlds where φ is true (i.e., the truth set $|\varphi|$); the empty ordering and orderings whose preferred alternative is the absurd proposition $\emptyset$ are then discarded.

$$(10) \quad \textbf{Declarative update: } \mathcal{R}[\triangleright\varphi] := \{r[\varphi] \mid r \in \mathcal{R} \text{ \& } r[\varphi] \neq \emptyset\},$$

$$\text{where } r[\varphi] := \{\langle p \cap |\varphi|, p' \cap |\varphi| \rangle \mid \langle p, p' \rangle \in r \text{ \& } p \cap |\varphi| \neq \emptyset\}$$

Imperatives update additively: as shown in (11), an update on $\mathcal{R}$ with $[\!|\varphi|]$ adds to every $r \in \mathcal{R}$ a new preference relation $\langle c_r \cap |\varphi|, c_r \cap |\neg\varphi| \rangle$, where c_r is the union of the domain and the range of r , thereby making φ the preferred alternative over $\neg\varphi$ in every $r \in \mathcal{R}$.

$$(11) \quad \textbf{Imperative update: } \mathcal{R}[\!|\varphi|] := \{r \cup \langle c_r \cap |\varphi|, c_r \cap |\neg\varphi| \rangle \mid r \in \mathcal{R}\}$$

Figure 1 shows the effect of updating $\mathcal{R}_0 = \{\{\langle\{w_{pq}, w_{p\bar{q}}, w_{\bar{p}q}, w_{\bar{p}\bar{q}}\}, \emptyset\rangle\}\}$ with $[\triangleright p]$ and $[\!|q|]$.

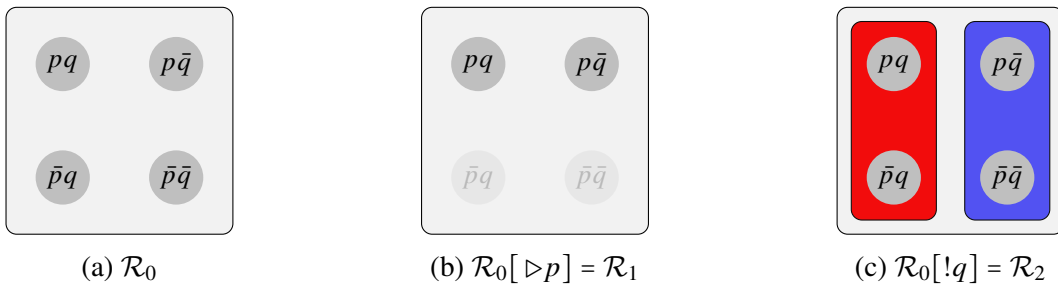


Figure 1: Declarative and imperative updates in preference semantics.

The declarative update $\mathcal{R}_0[\triangleright p]$ eliminates the $\neg p$ worlds and yields $\mathcal{R}_1 = \{\{\langle\{w_{pq}, w_{p\bar{q}}\}, \emptyset\rangle\}\}$ as the output preference state. The imperative update $\mathcal{R}_0[\!|q|]$ adds a new preference relation and yields $\mathcal{R}_2 = \{\{\langle\{w_{pq}, w_{p\bar{q}}, w_{\bar{p}q}, w_{\bar{p}\bar{q}}\}, \emptyset\rangle, \langle\{w_{pq}, w_{\bar{p}q}\}, \{w_{p\bar{q}}, w_{\bar{p}\bar{q}}\}\rangle\}\}$; this update promotes the set of q -worlds above the set of $\neg q$ -worlds.

Now, suppose we further update $\mathcal{R}_2$ with the opposite imperative $\neg q$. This update will add the pair $\langle\{w_{p\bar{q}}, w_{\bar{p}\bar{q}}\}, \{w_{pq}, w_{\bar{p}q}\}\rangle$ to the preference relation contained in $\mathcal{R}_2$, resulting in an output state that contains a preference relation which simultaneously ranks the proposition q above $\neg q$.

and $\neg q$ above q . On Starr's account, the presence of such a *cyclic* preference relation indicates that $!p$ and $!\neg p$ are preferentially inconsistent.

Next for connectives, updating with a conjunction amounts to updating with each conjunct sequentially. Updating with a disjunction amounts to first updating with each disjunct separately and then disjoining the two output preference states via set union.

(12) **Conjunctive and disjunctive updates:**

- a. $R[\varphi \wedge \psi] = R[\varphi][\psi]$
- b. $R[\varphi \vee \psi] = R[\varphi] \cup R[\psi]$

To illustrate, the update $R_0[\triangleright p \wedge !q]$ returns $R_3 = \{\{\{w_{pq}, w_{p\bar{q}}\}, \emptyset\}, \{\{w_{pq}\}, \{w_{p\bar{q}}\}\}\}$ as shown in Figure 2(a). It first eliminates all the $\neg p$ -worlds and then promotes the q -world above the $\neg q$ -world. The disjunctive update $R_0[\triangleright p \vee !q]$ returns the union $R_1 \cup R_2$ from Figure 1 above. Unlike its input state R_0 which consists of a single preference ordering, the output state R_4 comprises two distinct preference orderings: the left disjunct eliminates all the $\neg p$ -worlds and ranks the remaining p -worlds above the absurd proposition $\emptyset$, whereas the right disjunct ranks all the q -worlds above the $\neg q$ -worlds. In a nutshell, disjunction in preference semantics introduces a set of alternative preference orderings.

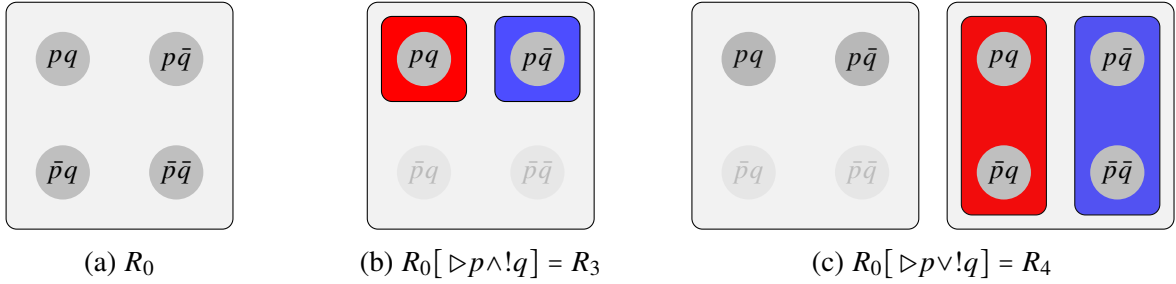


Figure 2: Connectives in preference semantics.

Move on to the definition of logical consequence. A commonly employed consequence notion in update semantics is the *update-to-test* consequence (cf. Veltman, 1996; Groenendijk et al., 1996), which Starr (2020) calls *preferential consequence*, defined as follows:

(13) **Preferential consequence**

$\varphi_1, \dots, \varphi_n \models_p \psi$ iff $R[\varphi_1] \dots [\varphi_n][\psi] = R[\varphi_1] \dots [\varphi_n]$ for all preference states R .

In words, φ entails ψ iff after the update with φ , the update with ψ is always idle. However, preferential consequence is too weak. For example, it fails to vindicate upward monotonicity inferences under imperatives such as $!(p \wedge q) \models !p$. Whereas $R[!(p \wedge q)]$ only promotes the proposition $p \wedge q$ as the preferred alternative over its negation, $R[!(p \wedge q)][!p]$ also promotes p over $\neg p$. Hence, the update with $!p$ after the update with $!(p \wedge q)$ is not idle.

To remedy this, Starr (2020) proposes a different consequence relation named *choice⁺ consequence*. Given a preference relation r (i.e., again, a set of ordered pairs of propositions), a proposition p is considered a good choice in r just in case, roughly, p is preferred to something,

and no other preferred alternatives are preferred to p or to anything entailed by p .⁴ The set of good choices in r , denoted by $ch(r)$, can then be strengthened into a set of good choices⁺, denoted by $ch^+(r)$, by closing it under superset and set intersection. It means that if p is a good choice and $p \subseteq p'$, then p' is a good choice⁺; if p and q are good choices, then $p \cap q$ is a good choice⁺. The Choice⁺ consequence is then defined as follows:

(14) **Choice⁺ consequence**

$\varphi_1, \dots, \varphi_n \models_{ch^+} \psi$ iff $Ch^+(R[\varphi_1] \dots [\varphi_n][\psi]) = Ch^+(R[\varphi_1] \dots [\varphi_n])$ for all preference states R , where $p \in Ch^+(R)$ iff for some $r \in R$: $p \in ch^+(r)$

Choice⁺ consequence vindicates upward monotonicity inferences because they are essentially baked into the definition of $ch^+(r)$. For example, since $p \cap q \subseteq p$, whenever $p \cap q$ is a good choice⁺, p must be a good choice⁺ as well. However, as we shall see, choice⁺ consequence has one notable drawback as it fails to derive the desired contextual equivalence in the case of imperative Hurford disjunctions.⁵

3.2. Difficulties of extant preference semantics

It is well-known that a disjunction cannot be felicitously used when one of its disjuncts contextually entails the other (Hurford, 1974), as exemplified by (15).

(15) John is in Paris or France.

It has been further observed that Hurford phenomena also arise with imperatives.

- (16) a. #Get an American or a Californian to do this job! (Ciardelli and Roelofsen, 2017)
b. #Go visit Paris or go visit France!

Regardless of whether the disjunction takes a narrow scope with respect to the imperative, as in (16a), or a wide scope, as in (16b), such disjunctions are deemed defective. The original preference semantics faces difficulties explaining the oddness of imperative Hurford disjunctions. One popular strategy to explain Hurford phenomena is to appeal to a notion of redundancy. Take the following non-redundancy constraint as an example (cf. Mayr and Romoli, 2016).

(17) **Non-Redundancy (NR):** A sentence φ cannot be used in context c if there is a sentence ψ s.t. ψ is a simplification of φ and is contextually equivalent to φ in c .

- i. ψ is a simplification (à la Katzir 2007) of φ if ψ can be derived from φ by replacing nodes in φ with their sub-constituents;
- ii. φ and ψ are contextually equivalent in c iff $c \cap |\varphi| = c \cap |\psi|$, with $|\varphi|$ and $|\psi|$ denoting the set of worlds where the respective sentence is true.

⁴More precisely, good choices are defined as follows in terms of promoted alternatives which in addition require that the preferred and dispreferred alternatives partition the set of all worlds contained in r .

(i) **Promoted alternatives in r**

$P(r) := \{p \mid \exists p' : \langle p, p' \rangle \in r \text{ \& } p \cup p' = c_r\}$

(ii) **Good choices in r**

$ch(r) := \{p \in P(r) \mid \nexists p' \in P(r), p'' : \langle p', p'' \rangle \in r \text{ \& } p \subseteq p''\}$

⁵With that being said, one advantage of choice⁺ consequence is its ability to avoid Ross's paradox. We will come back to this issue in section 5.

The simple Hurford disjunction in (15) violates NR because “John is in Paris or France” is contextually equivalent to the simplification “John is in France”. In order to adapt the redundancy analysis to imperative Hurford disjunctions, we need a notion of logical consequence that can correctly predict contextual equivalence between the imperative Hurford and one of its structural simplifications. For instance, we should expect the imperative Hurford $!p \vee !f$ in (16b) to be contextually equivalent to its simplification $!f$.

Choice⁺ consequence, however, fails to serve this purpose. More specifically, $!f$ fails to contextually entail $!p \vee !f$. Suppose the context set contains only three worlds w_p, w_m, w_b , representing that the addressee visits Paris, Marseille, or Berlin respectively. Set the input preferential state $R = \{r_0\}$ with $r_0 = \{\{\{w_p, w_m, w_b\}, \emptyset\}\}$. Compute the good choices⁺: $Ch^+(R[!f])$ contains $\{w_p, w_m\}$ along with all propositions entailed by it, whereas $Ch^+(R[!f][!p \vee !f])$ contains $\{w_p\}$ along with all propositions entailed by it. In other words, the proposition that the addressee visits Paris is a good choice⁺ in $Ch^+(R[!f][!p \vee !f])$ but not in $Ch^+(R[!f])$. Because $Ch^+(R[!f]) \neq Ch^+(R[!f][!p \vee !f])$, it follows that $!f \not\#_{ch^+} !p \vee !f$.

To address this inadequacy, we propose a novel notion of logical consequence, couched in an alternative implementation of preference semantics that defines preference orderings over possible worlds rather than propositions. Our world-based preference semantics also facilitates separation of background preference from foreground preference. Recall that in order to apply the practical triviality analysis to IoDs, we need to represent the addressee’s background preference, which can be in conflict with directions issued by imperatives. Suppose that the addressee initially prefers not losing her wallet; the imperative “hand over your wallet”, once accepted, will update the addressee’s preference by making “losing the wallet” the preferred alternative. However, existing preference semantics is unable to represent background preference. If background preference were to be encoded in the same way as foreground preference, then an update with $!\neg p$ on any preference state that has been updated with $!p$ would result in inconsistency.

4. Analysis

4.1. A world-based preference semantics

We first define preference models $\mathcal{M}$. They play more or less the same role as preference relations r found in existing preference semantics.

- (18) **(Preference model)** $\mathcal{M} := \langle c, \geq_b, \geq_f, V \rangle$, where c is a set of worlds representing the context set, $\geq_b$ and $\geq_f$ are two preorders (i.e., reflexive and transitive relations) representing the addressee’s background and foreground preference, and V is a valuation.

For any two worlds w and w' , $w \geq_{b/f} w'$ means that w is no less preferable than w' with respect to the background/foreground preference. Strict preference relations can be defined in terms of non-strict preference relations in the standard way: $w >_{b/f} w'$ iff $w \geq_{b/f} w'$ & $w' \not\geq_{b/f} w$. Two worlds w and w' are preferentially indistinguishable with respect to the background/foreground preference iff $w \geq_{b/f} w'$ and $w' \geq_{b/f} w$, abbreviated as $w \sim_{b/f} w'$. Given a set of worlds W , the minimal preference relation where nothing is strictly preferred is given by the universal relation $W \times W$. Lastly, the foreground preference $\geq_f$ does not have to be connected. That is, there could exist worlds w and w' such that $w \not\geq_f w'$ and $w' \not\geq_f w$, meaning that w and w' are not

comparable. As we shall see in Figure 3 below, incomparable worlds can be used to model updates resulting from conflicting imperatives such as $!p$ and $!\neg p$.

Declaratives update c via world-elimination as before: $\mathcal{M}[\triangleright\alpha]$ eliminates all worlds from c that do not satisfy α and then restricts everything else to the remaining α -worlds.

- (19) **(Declarative update)** $\langle c, \geq_b, \geq_f, V \rangle[\triangleright\alpha] := \langle c', \geq'_b, \geq'_f, V' \rangle$, where
- a. $c' := \{w \in c \mid \mathcal{M}, w \models \alpha\}$ (the set of worlds in $\mathcal{M}$ that settle α true)
 - b. $\geq'_b := \geq_b \cap (c' \times c')$
 - c. $\geq'_f := \geq_f \cap (c' \times c')$
 - d. $V'(p) = V(p) \cap c'$

Imperatives only update the foreground preference $\geq_f$ in $\mathcal{M}$ while keeping the rest unchanged.

- (20) **(Imperative update)** $\langle c, \geq_b, \geq_f, V \rangle[!\alpha] := \langle c, \geq_b, \geq_f^*, V \rangle$ where
- a. $\geq_f^* := \geq_f - \{\langle w, v \rangle \mid \mathcal{M}, w \not\models \alpha \text{ \& } \mathcal{M}, v \models \alpha\}$

Intuitively, imperatives can be understood as providing only external reasons for an agent to carry out the corresponding actions. As such, they do not alter the addressee's internal preference.⁶ Hence, imperatives update only foreground preference. The update with $!\alpha$ removes from the input foreground preference $\geq_f$ all pairs $\langle w, v \rangle$ where α is false at w and true at v , thereby making α the preferred alternative over $\neg\alpha$.

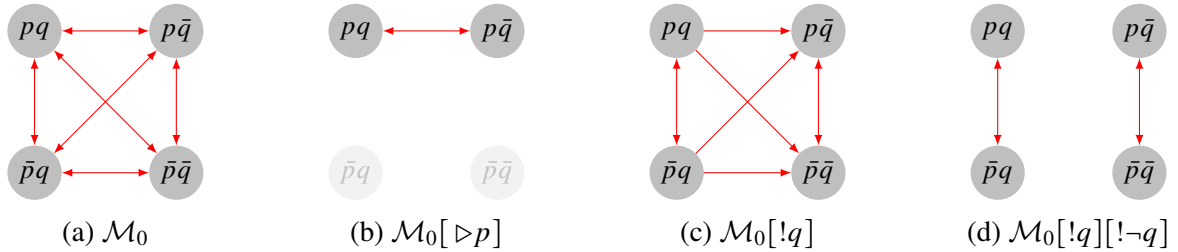


Figure 3: Declarative and imperative updates.

Figure 3 displays various updates on an input model $\mathcal{M}_0$ that contains four worlds. Both the foreground and background preferences of $\mathcal{M}_0$ are the minimal preference given by $W \times W$. For simplicity, only foreground preferences are explicitly depicted by red arrows: an arrow from w to w' represents $w \geq_f w'$; reflexive arrows are also suppressed. As with existing preference semantics, the update with the declarative $\triangleright p$ in (b) eliminates all $\neg p$ -worlds. The update with the imperative $!q$ in (c) eliminates all foreground arrows from $\neg q$ -worlds to q -worlds and returns a model where all q -worlds are now strictly ranked above $\neg q$ -worlds. Further updating it with the opposite imperative $!\neg q$ produces an unconnected foreground preference with incomparable worlds as shown in (d). This means that the resulting preference model is incoherent.

⁶This is not to say the addressee's background preference will never be influenced by conversational exchange. For example, someone who initially prefers iceberg lettuce over romaine lettuce may reverse her background preference after learning about the additional health benefits of romaine lettuce from her interlocutors. Nevertheless, this type of updates is usually associated with learning new factual information, which is typically the result of declarative updates. Therefore, it differs from an imperative update in which a command such as "Don't eat iceberg lettuce. Eat romaine lettuce!" is directly issued to the addressee.

Our next step is to combine the background and foreground preference to derive the combined preference. The intuitive gloss is that an agent will keep as many of her background preferences as possible unless they are overridden by foreground preferences.

$$(21) \quad (\textbf{Combined preference}) \geq := \geq_f - \{ \langle w, v \rangle \mid v >_b w \text{ \& } w \sim_f v \}$$

The combined preference $\geq$ is generated by removing from $\geq_f$ all arrows from w to v such that v is strictly preferred to w in the background, and w and v are preferentially indistinguishable in the foreground. This means that the combined preference will always be a strengthening (i.e., subset) of the foreground preference and will never overturn any foreground preference.

As an example, consider Figure 4 which contains preference models representing the carjacking scenario adduced earlier. The initial model $\mathcal{M}_1$ contains four worlds with c and w representing that the addressee keeps her car and her wallet respectively. Henceforth, we will use blue to depict background preference and red to depict foreground preference. The background preference in $\mathcal{M}_1$ yields the following ordering: $cw >_b c\bar{w} >_b \bar{c}w >_b \bar{c}\bar{w}$. The addressee prefers keeping her car, to keeping her wallet, to losing her wallet, and to losing her car. The foreground preference is given by the universal relation since no imperative has been uttered yet. In (b), the update with the imperative “Hand over your wallet!” promotes all $\bar{w}$ -worlds in the foreground above w -worlds. The combined preference of this output model is shown in (c). It yields the following new preference ordering as expected: $c\bar{w} > \bar{c}\bar{w} > cw > \bar{c}w$.

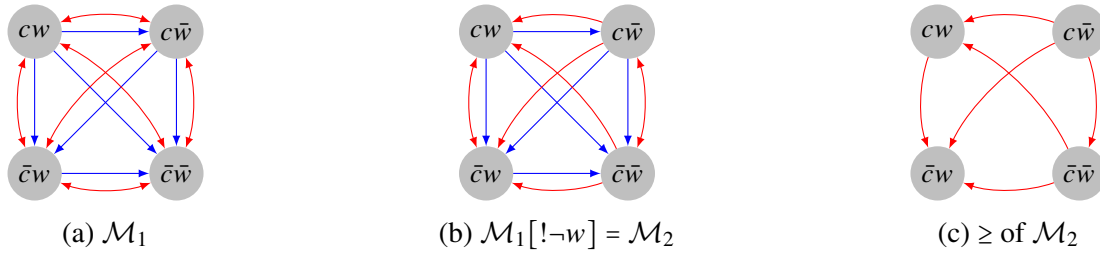


Figure 4: Combined preferences.

Given a combined preference $\geq$, we can calculate the set of *best worlds* in a model relative to $\geq$: w is a best world in $\mathcal{M}$ iff w is no worse than any other worlds in $\mathcal{M}$

$$(22) \quad (\textbf{Best worlds in a preference model}) \text{Best}(\mathcal{M}) := \{w \in c \mid \forall w' \in c : w \geq w'\}$$

$\text{Best}(\mathcal{M})$ returns an equivalence class such that every world in it is equally preferred and is preferred to every world outside it. Figure 5 adduces several examples showing the best worlds, highlighted in yellow, of various models. Notice that the last model in (c) does not contain any best world because the combined preference is not connected. We can then cash out a notion of preferential incoherence in terms of *best worlds*.

$$(23) \quad (\textbf{Incoherent models}) \mathcal{M} \text{ is preferentially incoherent iff } \text{Best}(\mathcal{M}) = \emptyset.$$

It is worth noting that not all models that contain incomparable worlds with respect to the combined preference are preferentially incoherent; only those whose alleged *best worlds* are incomparable are incoherent. Consider the two imperatives in (24) as an example.

Imperative-or-declarative and practical triviality

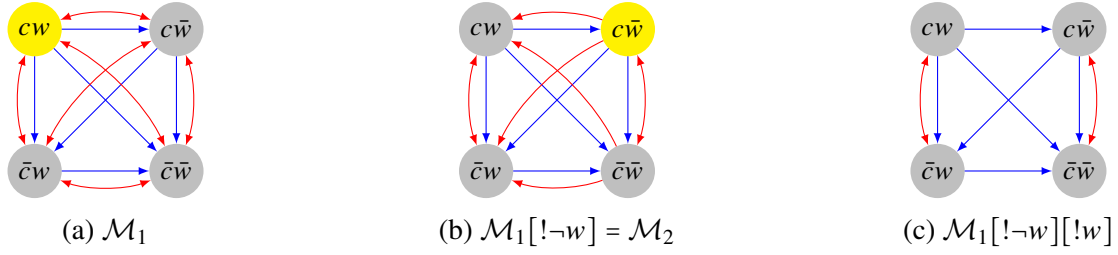


Figure 5: Best worlds and preferentially incoherent model.

- (24) a. Don't invite Paul without inviting Quinn. $!\neg(p \wedge \neg q)$
 b. Don't invite Quinn without inviting Paul. $!\neg(q \wedge \neg p)$

Analyzed as $!\neg(p \wedge \neg q)$ and $!\neg(q \wedge \neg p)$, the above two imperatives (prohibitives) are no doubt compatible. Together, they convey that the addressee should either invite both Paul and Quinn or invite neither of them. This is exactly what we predict: updating a minimal model with $!\neg(p \wedge \neg q)$ and $!\neg(q \wedge \neg p)$ yields a model whose best worlds are w_{pq} and $w_{\bar{p}\bar{q}}$. Yet in this model, the other two worlds $w_{p\bar{q}}$ and $w_{\bar{p}q}$ are incomparable, as the update with $!\neg(p \wedge \neg q)$ eliminates the arrow from $w_{p\bar{q}}$ to $w_{\bar{p}q}$ and the update with $!\neg(q \wedge \neg p)$ eliminates the arrow in the opposite direction. Thus, to allow sentences such as (24) to yield coherent models, our definition of preferential incoherence should zero in on comparability of alleged best worlds.

As with existing preference semantics, declarative and imperative updates are lifted from preference models to preference states. A preference state Σ is a set of preference models. Updating a preference state with a declarative or imperative clause amounts to updating each model in it point-wise, as detailed in (25). A conjunction induces sequential updates of its conjuncts. A disjunction induces parallel updates of its disjuncts and then collects the resulting preference models to form the new preference state.

- (25) **(Updates on preference states)** A preference state Σ is a set of preference models $\mathcal{M}$.
- a. $\Sigma[\triangleright\varphi] = \{\mathcal{M}[\triangleright\varphi] \mid \mathcal{M} \in \Sigma \text{ \& the } c \text{ of } \mathcal{M}[\triangleright\varphi] \text{ is not empty}\}$
 - b. $\Sigma[!\varphi] = \{\mathcal{M}[!\varphi] \mid \mathcal{M} \in \Sigma\}$
 - c. $\Sigma[\varphi \wedge \psi] = \Sigma[\varphi][\psi]$
 - d. $\Sigma[\varphi \vee \psi] = \Sigma[\varphi] \cup \Sigma[\psi]$

To illustrate using the carjacking example, consider the update with (2a) $!\neg w \vee \triangleright\neg c$. Suppose our initial preference state Σ_1 contains a single preference model $\mathcal{M}_1$. Then we have $\Sigma_1[!\neg w \vee \triangleright\neg c] = \Sigma_1[!\neg w] \cup \Sigma_1[\triangleright\neg c] = \{\mathcal{M}_2\} \cup \{\mathcal{M}_3\} = \{\mathcal{M}_2, \mathcal{M}_3\}$, as shown in Figure 6.

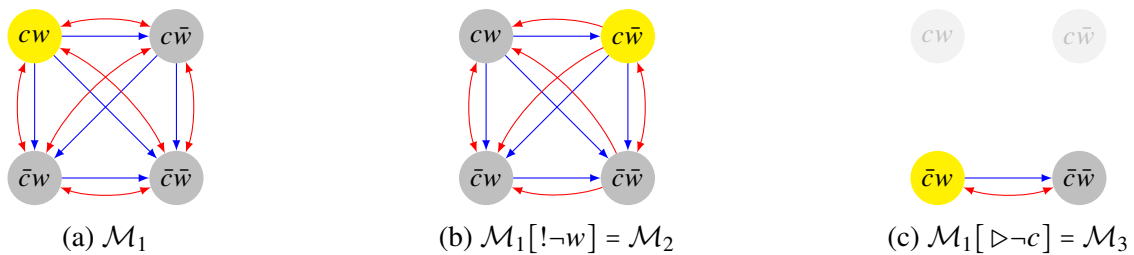


Figure 6: Update on $\Sigma_1 = \{\mathcal{M}_1\}$ with (2a) $!\neg w \vee \triangleright\neg c$.

4.2. Comparative desirability and practical triviality

Recall that in the informal sketch presented back in section 2.2, we have said that when the addressee is offered a choice between losing her wallet and losing her car, she will choose giving up her wallet because of her background preference. We will formally cash out this idea now via a notion of comparative desirability, which enables us to compare the desirability of different models contained in a preference state.

- (26) **(Comparative desirability)** For any two models in Σ sharing the same background preference, $\mathcal{M}$ is more desirable than $\mathcal{M}'$ (written $\mathcal{M} \gg \mathcal{M}'$) iff $Best(\mathcal{M}) \neq \emptyset$ and $\forall w \in Best(\mathcal{M}) \exists w' \in Best(\mathcal{M}') : w >_b w'$.⁷

We evaluate comparative desirability against the addressee's background preference. Since we have assumed that background preference is not affected by updates, we could assume that all models in the same state will share the same background preference. Applying this definition to the preference state $\{\mathcal{M}_2, \mathcal{M}_3\}$ above in Figure 6, we have $\mathcal{M}_2 \gg \mathcal{M}_3$ given that the addressee prefers $c\bar{w}$ to $\bar{c}w$ in the background. We then use $MD(\Sigma)$ to denote the set of most desirable models of Σ .

- (27) **(Most desirable models)** $MD(\Sigma) := \{\mathcal{M} \in \Sigma \mid \nexists \mathcal{M}' \in \Sigma : \mathcal{M}' \gg \mathcal{M}\}$

Whenever $MD(\Sigma)$ contains a single element, this is the model that the addressee will instinctively choose among its alternatives. For instance in Figure 6, $MD(\{\mathcal{M}_2, \mathcal{M}_3\}) = \{\mathcal{M}_2\}$; hence, the addressee will choose the model that corresponds to the disjunct $\neg w$ rather than $\triangleright \neg c$ in (2a). By contrast, Figure 7 details the update with the infelicitous IoD in (2b) $\neg c \vee \triangleright \neg w$.

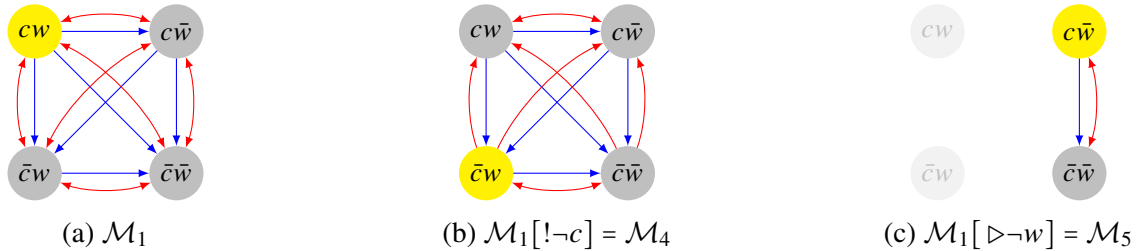


Figure 7: Update on $\Sigma_1 = \{\mathcal{M}_1\}$ with (2b) $\neg c \vee \triangleright \neg w$.

We have $\Sigma_1[\neg c \vee \triangleright \neg w] = \{\mathcal{M}_4, \mathcal{M}_5\}$. The best worlds are: $Best(\mathcal{M}_4) = \{\bar{c}w\}$ and $Best(\mathcal{M}_5) = \{c\bar{w}\}$. Since $c\bar{w} >_b \bar{c}w$, it follows that $\mathcal{M}_5 \gg \mathcal{M}_4$. Therefore, $MD(\Sigma_1[\neg c \vee \triangleright \neg w]) = \{\mathcal{M}_5\}$.

To take stock, we have seen that the most desirable model after the update with the felicitous IoD $\neg w \vee \triangleright \neg c$ is $\mathcal{M}_2$ in Figure 6, whereas the most desirable model after the update with the infelicitous IoD $\neg c \vee \triangleright \neg w$ is $\mathcal{M}_5$ in Figure 7. Compare $\mathcal{M}_2$ and $\mathcal{M}_5$ to their input model $\mathcal{M}_1$. We can make the following observation: $\mathcal{M}_5$ and $\mathcal{M}_1$ share the same foreground preference if we restrict attention to the set of worlds contained in $\mathcal{M}_5$, namely $\{c\bar{w}, \bar{c}w\}$; by contrast, $\mathcal{M}_2$ and $\mathcal{M}_1$ differ substantively with respect to their foreground preference.

⁷Because $Best(\mathcal{M})$ and $Best(\mathcal{M}')$ are sets of worlds and thus can be understood as representing propositions, we may conceive the definition above as a definition for what it means for a proposition to be more preferable than another. Effectively, this $\forall \exists$ definition is a standard way to retrieve a preference relation over propositions from a preference relation over possible worlds (see, e.g., Liu, 2011).

We now formulate practical triviality to capture this observation. For simplicity, we focus on cases where the MDs of the input and the output preference state both contain a single model.

- (28) **(Practical triviality)** φ is practically trivial with respect to Σ if the $\geq_f$ of $MD(\Sigma)$ is identical to the $\geq_f$ of $MD(\Sigma[\varphi])$ when restricted to the c of $MD(\Sigma[\varphi])$.

According to (28), with respect to the preference state $\Sigma_1 = \{\mathcal{M}_1\}$, (2b) $!\neg c \vee \triangleright \neg w$ is practically trivial, whereas (2a) $!\neg w \vee \triangleright \neg c$ is not. Additionally, given that Σ_1 , $\Sigma_1[!\neg c \vee \triangleright \neg w]$ and $\Sigma_1[!\neg w \vee \triangleright \neg c]$ all contain the same set of worlds c , (2b) and (2a) are both truth-conditionally trivial. Consequently, as (2b) “Hand over your car or I will take your wallet” is both truth-conditionally and practically trivial, it cannot be felicitously used in the context represented by Σ_1 .

4.3. Logical consequence and imperative Hurford disjunction

Just like the extant preference semantics, the Veltman-style update-to-test consequence given in (13) is too weak on the present analysis as well. For instance, it fails to vindicate the upward monotonicity inference from $!(p \wedge q)$ to $!p$. Since $\Sigma[!(p \wedge q)][!p]$ will also promote p -worlds above $\neg p$ -worlds in addition to promoting $(p \wedge q)$ -worlds above $\neg(p \wedge q)$ -worlds, we do not have $\Sigma[!(p \wedge q)][!p] = \Sigma[!(p \wedge q)]$ for all Σ .

Fortunately, a suitable consequence notion is not hard to come by. In fact, we have already introduced most of the necessary formal components needed to define such a notion. The only remaining step is to lift the set of best worlds from preference models to preference states.

- (29) **(Best worlds in a preference state)** $Best(\Sigma) := \bigcup_{\mathcal{M} \in \Sigma} Best(\mathcal{M})$

The set of best worlds in a preference state consists of all best worlds from every model in that state. We can similarly lift the definition of incoherent models from (23) to incoherent states: a preference state Σ is incoherent just in case $Best(\Sigma) = \emptyset$. Next, we can define a new consequence, which we call *optimality consequence*, as follows:

- (30) **(Optimality consequence)** $\varphi_1, \dots, \varphi_n \models_o \psi$ iff $Best(MD(\Sigma[\varphi_1] \dots [\varphi_n][\psi])) = Best(MD(\Sigma[\varphi_1] \dots [\varphi_n]))$ for all Σ .

Underlying optimality consequence is the idea that after the updates with all the premises, the additional update with the conclusion should not have any effect on the set of worlds considered *optimal* (i.e., the set of best worlds in the most desirable models).

Let us demonstrate how optimality consequence improves upon Starr’s (2020) choice⁺ consequence by establishing contextual equivalence in the case of imperative Hurford disjunction.

- (31) a. Go visit France or go visit Paris! $!f \vee !p$
 b. Go visit France! $!f$

To start with a simpler case, let us consider an initial state $\Sigma_0 = \{\mathcal{M}_0\}$ as shown in Figure 8(a). Assume that there are three worlds: the addressee visits Paris, Marseille, or Berlin. The background preference is completely neutral and thus suppressed in the figure. Computing the outputs of the relevant updates returns: $\Sigma_0[!f] = \{\mathcal{M}_1\}$ and $\Sigma_0[!f][!f \vee !p] =$

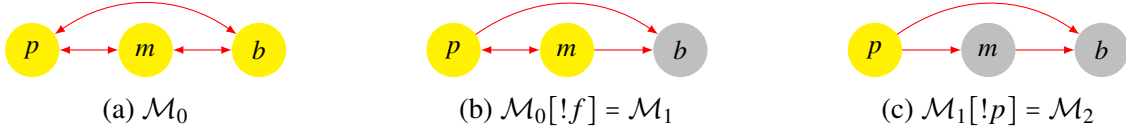


Figure 8: Imperative Hurford disjunction with a neutral background preference.

$\{\mathcal{M}_1\}[!f \vee !p] = \{\mathcal{M}_1, \mathcal{M}_2\}$. Because the two worlds w_p and w_m are preferentially indistinguishable in the background, $\mathcal{M}_1$ and $\mathcal{M}_2$ are equally desirable. So $MD(\Sigma_0[!f]) = \{\mathcal{M}_1\}$ and $MD(\Sigma_0[!f][!f \vee !p]) = \{\mathcal{M}_1, \mathcal{M}_2\}$. The sets of best worlds in the most desirable models are: $Best(MD(\Sigma_0[!f])) = \{w_p, w_m\}$ and $Best(MD(\Sigma_0[!f][!f \vee !p])) = \{w_p, w_m\}$. Hence, it follows that $Best(MD(\Sigma_0[!f])) = Best(MD(\Sigma_0[!f][!f \vee !p]))$, as desired.

The same result obtains even against a non-neutral background preference. Figure 9 depicts a scenario in which the addressee has an internal preference for visiting Marseille over the other two cities, with $\Sigma'_0 = \{\mathcal{M}'_0\}$. Running the calculation, we have $\Sigma'_0[!f] = \{\mathcal{M}'_1\}$ and $\Sigma'_0[!f][!f \vee !p] = \{\mathcal{M}'_1, \mathcal{M}'_2\}$. Because $w_m >_b w_p$, we have $\mathcal{M}'_1 \gg \mathcal{M}'_2$. Thus, $MD(\Sigma'_0[!f]) = MD(\Sigma'_0[!f][!f \vee !p]) = \{\mathcal{M}'_1\}$, meaning $Best(MD(\Sigma'_0[!f])) = Best(MD(\Sigma'_0[!f][!f \vee !p]))$.

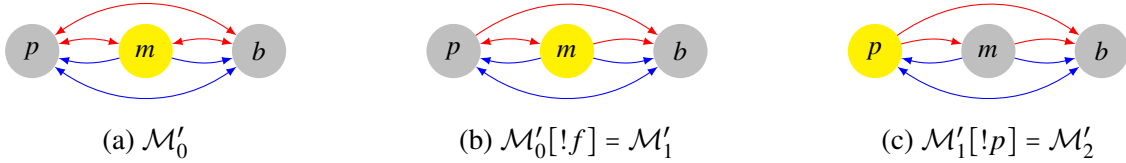


Figure 9: Imperative Hurford disjunction with a non-neutral background preference.

Let us supply one more example to further demonstrate the fruitfulness of our optimality consequence. Consider the following or-to-if inference from (32a) to (32b).

- (32) a. Hand over your wallet or I will take your car. $!\neg w \vee \triangleright \neg c$
 b. If you keep your wallet then I will take your car. $\triangleright(w \rightarrow \neg c)$

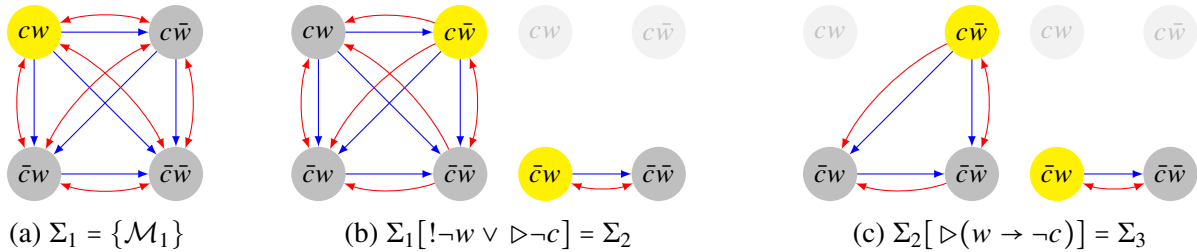


Figure 10: Or-to-if inference.

Assume $\Sigma_1 = \{\mathcal{M}_1\}$ as the initial preference state again. The relevant updates are depicted in Figure 10. In regard to the update with $\triangleright(w \rightarrow \neg c)$, although we have yet to supply a semantic clause for conditionals, we may for now simply treat this update as removing all worlds where the conditional's antecedent is true but the consequent is false—that is, removing all $(c \wedge w)$ -worlds. Given that the prior update with $!\neg w \vee \triangleright \neg c$ ensures that no $(c \wedge w)$ -world is part of the best worlds in its output state, the further update with $\triangleright(w \rightarrow \neg c)$ will always be idle as it has no effect on the set of best worlds.

5. Loose ends

While our rendition of preference semantics outperforms Starr’s (2020) analysis in a number of respects, such as accounting for imperative Hurford disjunction and resolving the puzzle of IoDs, it leaves several other problems that Starr’s framework is able to tackle unaddressed. We set out to tighten up these loose ends in this section.

5.1. Ross’s paradox

The first problem pertains to Ross’s paradox (cf. Ross, 1941).

- (33) a. Post the letter $!p$
 b. Post the letter or burn the letter $!p \vee !b$

Classical logic validates the unwelcome inference from (33a) to (33b). Starr’s choice⁺ consequence successfully blocks this inference. However, this is achieved only by globally prohibiting the rule of disjunction introduction from applying to imperatives. As a result, it fails to vindicate the inference from $!f$ to $!f \vee !p$ in the case of imperative Hurford disjunctions.

By contrast, the optimality consequence we employ is closer to classical consequence in this respect: it validates both the desirable inference from $!f$ to $!f \vee !p$ and the undesirable inference from $!p$ to $!p \vee !b$. To block Ross’s paradox, we will appeal to general pragmatic considerations rather than baking the restriction into the definition of logical consequence. A promising strategy is to appeal to interlocutors’ tendency to neglect “zero models”. For instance, the empty set, which canonically represents the absurd context, can be conceived of as a “zero-model”. Analyses based on this so-called “neglect zero” effect have been proposed to tackle free choice inferences and ignorance inferences (e.g., Aloni, 2022a, b). In fact, it is instructive to compare the odd inference in (33) with the equally odd inference from (34a) to (34b).

- (34) a. The letter has been posted.
 b. The letter has been posted or the letter has been burned.

This inference fails to preserve assertability. One can assert (34a) in a context without being able to assert (34b), as the latter carries the ignorance inference that the speaker still considers the second disjunct as a live possibility. This ignorance inference can be derived via the “neglect zero” principle. Suppose our initial context set c is $\{w_{p\bar{b}}, w_{\bar{p}b}, w_{\bar{p}\bar{b}}\}$ —that is, it is common ground that the letter cannot be both posted and burned. In a dynamic setting, the “neglect zero” principle translates to a requirement that prevents the two constituent updates $c[\varphi]$ and $c[\psi]$ belonging to the disjunctive update $c[\varphi \vee \psi]$ from yielding the empty set; if either $c[\varphi] = \emptyset$ or $c[\psi] = \emptyset$, then the whole update is deemed undefined (Aloni, 2022b). In the case of the illicit inference from (34a) to (34b), since $c[p][b] = \emptyset$, it follows that $c[p][p \vee b]$ is undefined. But since $c[p] = \{w_{p\bar{b}}\}$ and is not undefined, it follows that $c[p][p \vee b] \neq c[p]$, which means, under the standard update-to-test consequence, we have $p \not\models p \vee b$.

One way to extend this analysis to Ross’s paradox is to establish $!p \not\models_o !p \vee !b$ under the optimality consequence. We can accomplish this by enriching the “neglect zero” principle and adapting it to suit the current framework. In addition to requiring that a constituent update not yield any

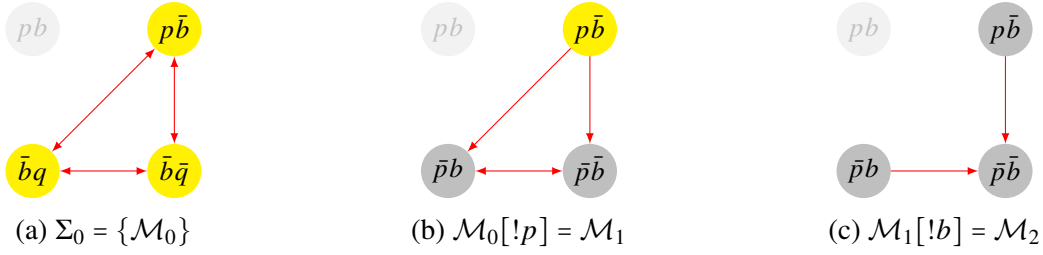


Figure 11: Ross's paradox.

preference state that contains the empty model (i.e., the model whose c is $\emptyset$), the principle also requires that the constituent update not yield any state that contains preferentially incoherent models (i.e., models whose set of best worlds is $\emptyset$).

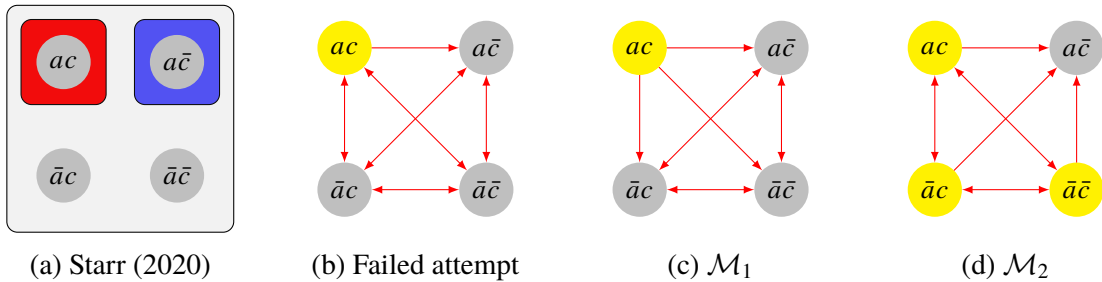
For example, consider the initial state $\Sigma_0 = \{\mathcal{M}_0\}$ as shown in Figure 11(a). It contains no background preference between the three possible alternatives. The update $\Sigma_0[!p]$ then yields $\Sigma_1 = \{\mathcal{M}_1\}$ as shown in 11(b). The subsequent update $\Sigma_1[!p \vee !b]$ yields $\Sigma_2 = \{\mathcal{M}_1, \mathcal{M}_2\}$. However, $\mathcal{M}_2$, as shown in 11(c), is preferentially incoherent as it contains no best worlds because of its unconnected preference structure. Hence, we may treat the update $\Sigma_1[!p \vee !b]$ as undefined according to “neglect zero”, which thereby blocks the inference from $!p$ to $!p \vee !b$.

5.2. Conditional imperatives

In section 4.3, we assumed that conditionals are material implications whose update potential comes down to eliminating all worlds where the conditional's antecedent is true and the consequent is false from all preference models. However, this simple treatment of conditionals fails to generalize to conditional imperatives such as (35).

(35) If the air conditioner is on then close the window! $o \rightarrow !c$

Intuitively, (35) does not rule out any world. On Starr's (2020) account, conditional imperatives have the update potential of promoting the proposition that satisfies both the antecedent and the consequent over the proposition that satisfies the antecedent but not the consequent. For example, the update with (35) on the initial preference state $\{\{\{w_{ac}, w_{a\bar{c}}, w_{\bar{a}c}, w_{\bar{a}\bar{c}}\}, \emptyset\}\}$ returns $\{\{\{w_{ac}, w_{a\bar{c}}, w_{\bar{a}c}, w_{\bar{a}\bar{c}}\}, \emptyset\}, \{\{w_{ac}\}, \{w_{a\bar{c}}\}\}\}$ which, as depicted in Figure 12(a), ranks the proposition $\{w_{ac}\}$ in red above the proposition $\{w_{a\bar{c}}\}$ in blue.

Figure 12: Conditional imperative: $a \rightarrow !c$.

On our account, which defines preferences directly over worlds rather than over propositions, interpretation of conditional imperatives turns out to be more difficult. To start, we cannot simply interpret $a \rightarrow !c$ as an update that removes the arrow from the world $w_{a\bar{c}}$ to the world w_{ac} , thereby promoting w_{ac} over $w_{a\bar{c}}$. As Figure 12(b) shows (assuming that the background preference is neutral), the resulting preference model will contain a foreground preference that is non-transitive and thus illicit. In order to ensure that $\geq_f$ is still transitive, we must remove, at minimum, either the arrows $ac \geq_f \bar{a}c$ and $ac \geq_f \bar{a}c$ together as shown in 12(c), or the arrows $\bar{a}c \geq_f a\bar{c}$ and $\bar{a}\bar{c} \geq_f a\bar{c}$ together as shown in 12(d).

However, both options face difficulties. If we remove $ac \geq_f \bar{a}c$ and $ac \geq_f \bar{a}c$, the addressee's preference structure now prefers w_{ac} to all the other worlds, thereby collapsing $a \rightarrow !c$ to $!(a \wedge c)$ (*Turn on the air conditioning and close the window*). Alternatively, if we remove $\bar{a}c \geq_f a\bar{c}$ and $\bar{a}\bar{c} \geq_f a\bar{c}$, then $w_{a\bar{c}}$ becomes the least preferred world, thereby collapsing $a \rightarrow !c$ to $!\neg(a \wedge \neg c)$. In other words, the conditional imperative becomes an imperative to see to it that a material implication holds. While this may seem like a better option than collapsing $a \rightarrow !c$ to $!(a \wedge c)$, it is still less than ideal. For one, on this interpretation $a \rightarrow !c$ would be equivalent to $\neg c \rightarrow !\neg a$. This prediction is not borne out as (35) above and (36) below seem to issue different commands.

(36) If the window is open then keep the air conditioner off!

Suppose the air conditioner is on and the window is open. (35) commands the addressee to close the window, whereas (36) commands the addressee to turn off the air conditioner. In terms of their intended update effect, whereas (35) aims to promote w_{ac} over $w_{a\bar{c}}$, (36) aims to promote $w_{a\bar{c}}$ over w_{ac} . Interpreting $a \rightarrow !c$ as $!\neg(a \wedge \neg c)$ fails to respect this difference.

That being said, there is a third possibility to model conditional imperatives available to us on the present account. We could treat $a \rightarrow !c$ as generating a preference state that comprises both $\mathcal{M}_1$ in Figure 12(c) and $\mathcal{M}_2$ in 12(d). Such a preference state can be derived from taking the update $\Sigma[p \rightarrow !q]$ to first promote w_{pq} over $w_{p\bar{q}}$ in every $\mathcal{M} \in \Sigma$ —as demonstrated in Figure 12(b)—and then generate a set of *minimal contractions* based on the results from the first step so that transitivity is reestablished for every model in the final output state. Both $\mathcal{M}_1$ and $\mathcal{M}_2$ are minimal contractions of the illicit structure in 12(b). This solution overcomes the aforementioned problem of collapse. For instance, $a \rightarrow !c$ and $\neg c \rightarrow !\neg a$ now induce different updates: the model $\mathcal{M}_1$ in 12(c) will be a member of $\{\mathcal{M}_0\}[a \rightarrow !c]$ but not of $\{\mathcal{M}_0\}[\neg c \rightarrow !\neg a]$.⁸

6. Conclusion

We have developed a world-based preference semantics capable of explaining the puzzles surrounding imperative-or-declaratives through a notion of practical triviality. In addition, the optimality consequence newly defined in this framework yields a number of desirable results. In future work, we plan to further enrich the framework so as to integrate interrogatives.

⁸A difficulty remains. Unlike Starr (2020), we fail to render $a \rightarrow !c$ and $a \rightarrow !\neg c$ incompatible. For example, if we further update $\{\mathcal{M}_0\}[a \rightarrow !c] = \{\mathcal{M}_1, \mathcal{M}_2\}$ from Figure 12 with the conflicting imperative $a \rightarrow !\neg c$, the output state will not be incoherent. This is so because the model that contains $w_{a\bar{c}}$ and $w_{a\bar{c}}$ as its best worlds will belong to the final output state; thus, not every model in the output is preferentially incoherent. One way out of this problem is to redefine incoherence so that a preference state is considered incoherent so long as it contains at least one preferentially incoherent model. We will leave further investigation of this proposal to future work.

References

- Aloni, M. (2022a). Logic and conversation: the case of free choice. *Semantics and Pragmatics* 15, 5–1.
- Aloni, M. (2022b). Neglect-zero effects in dynamic semantics. In *Tsinghua Interdisciplinary Workshop on Logic, Language, and Meaning*, pp. 1–24.
- Biezma, M. and K. Rawlins (2016). Or what?: Challenging the speaker. *Proceedings of NELS 46th 1*, 93–106.
- Ciardelli, I., J. Groenendijk, and F. Roelofsen (2018). *Inquisitive Semantics*. Oxford University Press.
- Ciardelli, I. and F. Roelofsen (2017). Hurford’s constraint, the semantics of disjunction, and the nature of alternatives. *Natural Language Semantics* 25(3), 199–222.
- Clark, B. (1993). Relevance and “pseudo-imperatives”. *Linguistics and Philosophy*, 79–121.
- Franke, M. (2005). Pseudo-imperatives. Master’s thesis, ILLC, Amsterdam.
- Groenendijk, J., M. Stokhof, F. Veltman, et al. (1996). Coreference and modality in the context of multi-speaker discourse. *Context Dependence in the Analysis of Linguistic Meaning*, 195–216.
- Hurford, J. R. (1974). Exclusive or inclusive disjunction. *Foundations of Language* 11(3), 409–411.
- Katzir, R. (2007). Structurally-defined alternatives. *Linguistics and Philosophy* 30, 669–690.
- Kaufmann, M. (2012). *Interpreting Imperatives*, Volume 88. Springer Science & Business Media.
- Liu, F. (2011). *Reasoning about Preference Dynamics*, Volume 354 of *Synthese Library*. Dordrecht: Springer.
- Mayr, C. and J. Romoli (2016). A puzzle for theories of redundancy: Exhaustification, incrementality, and the notion of local context. *Semantics and Pragmatics* 9(7), 1–48.
- Murray, S. E. and W. B. Starr (2020). The structure of communicative acts. *Linguistics and Philosophy* 44(2), 425–474.
- Roberts, C. (2023). Imperatives in a dynamic pragmatics. *Semantics and Pragmatics* 16, 7–EA.
- Ross, A. (1941). Imperatives and logic. *Theoria* 7(1), 53–71. In “Discussions”.
- Schwager, M. (2004). Don’t be late or you’ll miss the first slot. *NASSLI, UCLA*.
- Simons, M. (2001). Disjunction and alternativeness. *Linguistics and Philosophy* 24(5), 597–619.
- Starr, W. B. (2018). Conjoining imperatives and declaratives. In *Proceedings of Sinn und Bedeutung*, Volume 21, pp. 1159–1176.
- Starr, W. B. (2020). A preference semantics for imperatives. *Semantics and Pragmatics* 13, 6–1.
- van Benthem, J. and F. Liu (2007). Dynamic logic of preference upgrade. *Journal of Applied Non-Classical Logics* 17(2), 157–182.
- van Rooij, R. and M. Franke (2010). Promises and threats with conditionals and disjunctions. *Ms., University of Amsterdam*.
- Veltman, F. (1996). Defaults in update semantics. *Journal of Philosophical Logic* 25(3), 221–261.
- Zhang, Y. R. (2025). QUD-mediated redundancy. In *Proceedings of Sinn und Bedeutung*, Volume 29, pp. 1784–1801.