



Universiteit
Leiden
The Netherlands

Different scalar terms are affected by face differently

Terkourafi, M.; Weissman, B.; Roy, J.

Citation

Terkourafi, M., Weissman, B., & Roy, J. (2020). Different scalar terms are affected by face differently. *International Review Of Pragmatics*, 12(1), 1-43. doi:10.1163/18773109-01201103

Version: Publisher's Version

License: [Creative Commons CC BY-NC 4.0 license](#)

Downloaded from: <https://hdl.handle.net/1887/87031>

Note: To cite this publication please use the final published version (if applicable).

Different scalar terms are affected by face differently

Marina Terkourafi

Leiden University

Corresponding author: m.terkourafi@hum.leidenuniv.nl

Benjamin Weissman

Rensselaer Polytechnic Institute

Joseph Roy

The American Society for Engineering Education

Abstract

Research on the effect of face-orientation on scalar implicatures has claimed that face-threatening contexts are one type of context in which scalar implicatures are not warranted. However, that research has been based on the two staples of scalar implicature research, *some* and *or*. Given research on scalar diversity has shown that these terms are rather exceptional in inducing high rates of scalar implicatures, we believe it is time for a reassessment. We explored the relationship between scalar implicatures and face concerns by means of an experiment involving eight types of scalar terms in face-boosting and face-threatening contexts. While our results showed that *some* and *or* reliably tended to induce scalar implicatures in both types of contexts, confirming the findings of scalar diversity research in this respect, we failed to replicate previous findings that face-threatening contexts do not induce scalar implicatures. We discuss reasons for these findings and how face concerns should be implemented for future experimentation in this vein.

Keywords

scalar implicatures – face-threat – face-boost – affect – scalar diversity

1 Introduction

The recent experimental turn in linguistic pragmatics (Noveck & Reboul, 2008) has injected new life into debates that had previously been waged almost exclusively on theoretical grounds. One such debate concerns the boundary between semantics and pragmatics; namely, what part of listeners' interpretation of a speaker's utterance is contributed by their knowledge of the language and what part is guided by contextual assumptions based on the surrounding discourse or extra-linguistic factors. Default accounts (e.g., Levinson, 2000) emphasise the importance of the speaker's choice of words, which justifies drawing conclusions listeners may later need to retract if they turn out to be unwarranted in the context at hand. Contextualist accounts (e.g., Sperber & Wilson, 1986), on the other hand, deny the existence of such *a priori* interpretations, arguing instead that listeners begin enriching linguistic representations with contextual information from the start. Finally, constraint-based accounts (e.g., Degen & Tanenhaus, 2015) relativise defaults to contextual conditions and attempt to identify the parameters that constrain their generation. Although the correctness of one of these accounts over the others remains to be proven, the phenomenon of scalar implicature has provided fertile grounds for testing their relative merits. The present article contributes to this literature by experimentally investigating the interpretation of a variety of scalar terms in face-boosting vs. face-threatening contexts and using the results to expand on previous findings in this regard.

Scalar implicatures (henceforth *SIs*) are a type of inference in which the use of an informationally weaker term (e.g., *some*) is taken to mean that an informationally stronger term (e.g., *all*) does not apply.¹ For a real-life example, consider (1) below:

- (1) "Consumers claim some Samsung washing machines explode." (<http://abc13.com/technology/consumers-claim-some-samsung-washing-machines-explode/1058097/>; retrieved 26/7/2017)²

¹ Geurts (2010: 31) defines *SIs* as a subset of a more general category of quantity implicatures, in which the pertinent question is not why the speaker did not say (in a relevant sense of 'say') the stronger proposition but whether the speaker believes the stronger proposition to be true (2010: 3). Since for a politeness motive to be plausible, the "competence assumption" (Geurts, 2010: 29) must hold—the listener thinks the speaker knows whether the stronger proposition is true but is not saying so out of politeness—cases where the speaker is thought not to know which of the two meanings is true are excluded from consideration. In other words, the phenomena discussed in this article belong to Geurts' "scalar" subset, which is why we continue to use the term '*SI*' throughout.

² Following Levinson (1995: 110, fn. 3), we indicate utterances in double quotation marks and

In their response to this complaint, Samsung asserted that the four consumers interviewed “each have different washing machine models [...] [w]hile unfortunate, their experiences are very rare when compared to the number of washing machines we sell each year” (ibid.). This suggests that they interpreted (1) in a two-sided³ way as ‘some but not all Samsung washing machines explode’;⁴ although (1) would be no less true if it turned out that all Samsung washing machines did (if they all explode, then necessarily some of them do). In other words, the inference from “some *x*” to ‘some but not all *x*’ can be cancelled (or does not always arise, depending on the framework one adopts).

This article brings together two lines of work on SIs. The first stems from studies of an expanded range of scalar terms (Doran et al., 2009; Van Tiel et al., 2016, among others), beyond the now classical pairs ⟨*some*, all⟩ and ⟨*or*, and⟩ investigated in much of the related literature. These studies found significant variability in the extent to which different scalar terms are likely to induce a SI, a phenomenon described as “scalar diversity” by Van Tiel et al. (2016). In this line of work, scalar diversity is understood as a property of scalar terms independent of their situational context of utterance. The factors explored for their impact on the likelihood of a two-sided reading are (lexical semantic) properties of the scalar term. This is true also of more recently proposed factors, such as upper-bound excluded local enrichment (UBELE, Sun et al., 2018) and whether the stronger alternative instantiates the critical bound (Simons & Warren, 2018), put forward as additional dimensions along which scalar terms differ.

A second line of work has investigated how situational context affects the likelihood of SI derivation. This work has revealed an interesting relationship between SIs and interpersonal factors such as face (Brown & Levinson, 1987). In a series of experiments, Bonnefon and colleagues found that a scalar term like “some” is significantly less likely to be interpreted as ‘some but not all’ if the listener attributes a polite (i.e., face-saving) intention to the speaker (Bonnefon & Villejoubert, 2006; Bonnefon, Feeney & Villejoubert, 2009; Feeney & Bon-

implicatures, which are not spoken out loud, in single ones. We use the +> symbol to indicate a conversational implicature.

- 3 Other terms used in the literature for what, following Horn (1992), we are calling ‘two-sided’ vs. ‘one-sided’ readings include ‘narrowed’ vs. ‘broadened’ (Bonnefon et al., 2009), ‘upper-bounding’ vs. ‘lower-bounding’ (Breheny et al., 2005) and ‘pragmatic’ vs. ‘semantic’ (Holtgraves & Kraus, 2018).
- 4 This interpretation assures us that the SI was derived in this case (at least by the company’s representatives), despite the fact that it is embedded. For a summary of issues surrounding embedded implicatures, see Chemla & Singh, 2014a, b.

nefon, 2012; Bonnefon, Dahl & Holtgraves, 2015). Based on these results, they concluded that face-threatening contexts are one type of context in which SIs are not warranted.

We address the following research questions: (1) How are scalar terms interpreted when embedded in face-threatening vs. face-boosting contexts? (2) Do different scalar terms behave alike in this respect or is there variation among them? We present an experiment designed to answer these questions and discuss our findings in light of previous work. Our work differs from most previous work on scalar diversity in that we considered utterances embedded in short story contexts rather than in isolation, since we are precisely interested in the impact of situational context on scalar diversity. It also differs from previous work on the impact of face on SIs in that we argue for a different, contextually-motivated, construal of face-boost and face-threat. We also define and theoretically motivate face-boost, which remains under-theorized in previous work. We open with an overview of previous research on these topics (section 2), followed by a discussion of some problematic aspects that serves to clarify how we defined important theoretical notions and why (section 3). In section 4, we describe the experiment and analyse our results. Section 5 discusses their significance and limitations, while section 6 provides some concluding thoughts.

2 Putting scalar diversity in context

2.1 *Previous work on scalar diversity*

Building on Doran et al. (2009), Van Tiel et al. (2016) tested 43 scalar terms representing five types of scales (2 quantifiers, 1 adverb, 2 auxiliary verbs, 6 main verbs, and 32 adjectives) and found significant variability in the extent to which they induced SIs, both in neutral (SI-rates range: 100% to 4%) and non-neutral (SI-rates range: 93% to 4%) sentential contexts.⁵ To explain this variability, they explored a number of factors including association strength between scalar alternatives (implemented as the ease with which a weaker term calls up a stronger alternative), their grammatical class (open vs. closed),

5 In their neutral contexts, they used pronominal subjects (example: "John says: She is intelligent. Would you conclude from this that, according to John, she is not brilliant?"), while in their non-neutral ones the subject was a full NP (example: "John says: This student is intelligent. Would you conclude from this that, according to John, she is not brilliant?"). Of course, the extent to which any sentential context can be considered neutral (e.g., free of stereotypical assumptions about gender and about why 'John' would say that about 'her') is an open question.

relative (between scalar alternatives) or absolute word frequencies, semantic relatedness (the extent to which they share collocates in a corpus), semantic distance (the size of the interval between them on the scale), and boundedness (whether the dimension over which the scalar terms quantify has a specifiable, lexicalized endpoint). Of these, only the last two turned out to explain some of the variance in SI rates; yet, even the strongest predictor, boundedness, only explained 10% of the variance. Another 30% was due to items and participants, while a full 48% remained unexplained. As a tentative explanation for this the authors surmise that statistical tendencies in language may inform participant judgements in ways that cannot be predicted by attributes of the scalars themselves. Thus, while having established the reality of scalar diversity, this study was unable to identify its main source(s).

Following in Van Tiel et al.'s footsteps, Sun et al. (2018) conducted further experiments that confirmed the earlier findings. Using the same materials and two new tests, they were able to identify UBELE as an additional factor explaining some of the variance in Van Tiel et al.'s (2016) results. UBELE is a measure of the propensity of scalar terms to receive a two-sided reading prior to considering the sentential context. Having found that scalar terms differ in this regard, Sun et al. go on to interpret their findings as evidence for a dual route account, in which scalar terms may be locally or globally enriched. Standard Gricean inference based on alternatives corresponds to global enrichment, while local enrichment is the result of Bayesian reasoning based on the prior probability that a term will receive a two-sided or a one-sided reading. Local enrichment is an interesting notion: first, it can go either way (scalar terms may have a preference for two-sided readings—UBELE—or for one-sided readings; that is, it can track genuine underspecification w.r.t. these two 'literal' meanings of a scalar rather than taking the one-sided reading as the encoded one that goes through in case the SI is defeated/not generated); and second, since it does not rely on alternatives, it is not vulnerable to scale availability and how experimental tasks may affect this (cf. McNally 2017). UBELE is of course itself something to be explained: what determines the prior probability that a scalar will be interpreted in a two-sided or one-sided way? Sun et al. do not raise this question; however, one hypothesis is that stylistic and other properties of the scalar (register, dialect, etc.) may interact with societal norms regarding expected uses to generate the prior probabilities reflected—as far as two-sided readings go—in UBELE.

This possibility seems worth exploring in light of recent work on negative strengthening (Benz et al., 2018; Gotzner et al., 2018). Negative strengthening occurs when the negation of a stronger term is taken to imply that a weaker term is also not applicable (e.g., when "not brilliant" is interpreted as 'not intelli-

gent'). This amounts to informational strengthening of "not brilliant" to exclude 'intelligent', a possibility left open by scalar reasoning. The relevant reasoning has been described as blocking, or an application of Horn's (1984) R-principle (inference to the stereotype), which may be attributed to a politeness motivation: a boss who is "not happy" with an employee's performance may well be implicating that she is 'not content' with it while tempering the negativity of her remark. Benz et al. (2018) found that adjectives more likely to be negatively strengthened were less likely to be interpreted in a two-sided way in Van Tiel et al.'s (2016) results; in other words, they found a negative correlation between negative strengthening and SI-rate derivation. Gotzner et al. (2018) investigated the structure of the corresponding adjectival scales and added vagueness (no context-independent standards) and adjectival extremeness to the list of factors affecting the likelihood of a two-sided reading.

The importance of adjectival extremeness was confirmed in another study by Simons & Warren (2018), who found that "the critical factor may not simply be the boundedness of the underlying scale, but that the stronger alternative instantiates the critical bound" (2018: 277). Simons & Warren's study is also methodologically interesting, in that, unlike the studies reported so far, which relied on assessments of individual utterances via direct probing, they assessed the interpretation of each scalar adjective in multiple rich contexts in which three scalars were embedded at a time, and asked participants to provide consistency ratings for statements corresponding to different interpretations of each scalar that did not explicitly mention stronger scalemates. They were thus able to show that the scalar diversity effects found in previous studies can also be obtained through more natural, indirect elicitation tasks, eliminating some of the criticisms directed at the earlier studies (e.g., by McNally, 2017) and demonstrating the robustness of the relevant phenomena.

These experimental studies have made significant inroads into understanding the nature of scalar diversity, especially in the adjectival domain. Their results converge in several ways, notably the greater likelihood of two-sided readings for 'logical' terms such as *some* and *or* as opposed to more evaluative and vague ones, and the importance of boundedness as a semantic factor favouring two-sided interpretations. Yet, what has so far remained unexplored is the input of situational context in this process. Van Tiel et al. (2016: 142) explicitly argue against the inclusion of narrative contexts, as these vary in wordiness among others, seeing it as a weakness of Doran et al.'s (2009) study. However, McNally (2017) has pointed out that at least some scalar terms (especially adjectives) may not have fixed meanings outside of specific contexts of use and it is only under such circumstances that they can be candidates for scalar enrichment (when the alternatives are made clear). The upshot of

her discussion is that formal semantics/pragmatics cannot continue to ignore parameters of the context that feed into the interpretation of scalars and therefore also into experimental subjects' considerations when carrying out experimental tasks. In support of her view, Simons & Warren's (2018) study showed that scalar diversity effects also obtain in richer contexts. Going forward, there are some good reasons to expect that situational context may constitute an additional factor interacting with sentence context and the semantics of the scalar to produce the observed rates of SI derivation. The next section provides a critical overview of this research.

2.2 *Previous work on SIs and face*

Studies that have investigated the impact of situational context on SIs have focused on face-threatening vs. face-boosting contexts. Adopting the notion of "face" from Brown and Levinson's work on politeness (1987), where it refers to "the public self-image that [every rational communicator] wants to claim for himself" (1987: 61), researchers have argued that utterances of the type "Some X-ed" are face-boosting if the semantic content of the predicate X is favourable to the listener, but face-threatening when X expresses something unfavourable for the listener. In one set of experiments, Bonnefon et al. (2009) investigated contrasting pairs of utterances such as (2) and (3) below.

(2) Some people loved your poem.

(3) Some people hated your poem.

They presented participants with a scenario in which one of these utterances is confided to the author of a poem that was discussed at a poetry group meeting he was unable to attend by another group member. After reading the scenario followed by one of the two utterances, participants answered the following Yes/No question:

(4) 'From what this fellow member told you, do you think it is possible that everyone loved [hated] your poem?'

They found that participants were inclined to interpret (2) two-sidedly as 'some but not all people loved your poem', but they were significantly less inclined to do the same in the case of (3) (they interpreted (3) as 'some and possibly all people hated your poem').

Bonnefon and his colleagues attributed these results to the different face orientations of the two utterances. Specifically, they argue that in the case of (3),

but not (2), the speaker is likely to want to mitigate the unpleasantness of her message out of consideration for the listener. She may, then, say “Some people hated your poem,” fully knowing that all of them did. Aware of this possibility, the listener may in turn not be fooled by her choice of words and suspect that all of those present hated his poem—all the while being grateful to her for softening the blow. Feeney and Bonnefon (2012) replicated these results with the logical connective *or*, and went on to interpret their findings as indicating that face-threatening contexts favour one-sided interpretations of scalar terms (i.e., interpreting “some” as ‘some and possibly all’).

More recent studies have studied the impact of face-threat on the derivation of SIs using online measures. In one study using ERPs, Holtgraves & Kraus (2018) examined five scalar expressions (*some*, *sometimes*, *like*, *good*, and *probable*) embedded in conversational contexts under face-threatening and non-face-threatening conditions. Participants read short scenarios followed by a target utterance containing the scalar in the first half of the utterance (e.g., *some*) and a two-sided (e.g., *not all*) or one-sided (e.g., *all*) continuation in its second half, as in (5):

- (5) John couldn’t make it to Susan’s party. To make up for it, he made her some cookies and brought them over for the party. After the party was over, John asked Susan if any of his cookies were left over. Susan says:
 There were some left over, specifically, they were not all left over. (two-sided)
 There were some left over, specifically, they were all left over. (one-sided)

In the face-threatening condition, the recipient of Susan’s utterance was John, i.e. the person who had baked the cookies and whose face could therefore be threatened by an assertion that all of his cookies were left over. In the non-face-threatening condition, the recipient of Susan’s utterance was a third party, whose face was assumed to be unaffected by the content of her utterance. Neural responses to the screen containing the scalar term (“some left over”) did not vary between the two conditions. However, neural responses to the one-sided continuation (time-locked to the screen “all left over”) resulted in a larger P300 component compared with the two-sided continuation (time-locked to the screen “not all left over”), and this difference was greater when the situation was face-threatening. Additionally, larger P200 responses⁶ were observed for

⁶ P200 is generally considered an early attention-related ERP component, indexing fast, automatic detection of emotionally salient stimuli. It tends to be more pronounced for negative material (Carretie et al., 2001; Delplancque et al., 2004).

the positively valenced terms (*good* and *like*) in the face-threatening condition (2018: 99), leading the authors surmise that “situations which invoke politeness are more emotionally salient than situations where the probability of offending someone is low” (2018: 102). We return to this finding in section 5 below.

According to Holtgraves & Kraus, this pattern suggests that in conversational contexts in general, it is the two-sided interpretation that is expected. In other words, SIs are generated upon encountering the scalar term, especially when the situation is face-threatening, and that is what is causing the larger P₃₀₀ indicating an updating of the context when the one-sided continuation is subsequently encountered. Conversely, in non-face-threatening contexts, the P₃₀₀ effect on the continuation is smaller, suggesting that the one-sided reading is less unexpected in these contexts.

With respect to these results, one might note that the one-sided continuation (“they were all left over”) in the second part of the utterance in the face-threatening contexts constitutes an (additional) on-record face-threat (rather than its mitigation). This may have led participants to wonder why Susan would say this to John’s face, when she could have stopped after the first part of her utterance and merely implicated it—other than to hurt his feelings. If so, the larger P₃₀₀ effect observed in these contexts could well be indicating participants’ surprise at Susan’s behaviour (their updating the context with information about her face-threatening intention), rather than at the one-sided reading per se (their updating the context with information about the meaning of the scalar). Unfortunately, the experimental setup of this study does not allow us to distinguish between these two possibilities.

At first sight, these results are opposed to those of Bonnefon and his colleagues but another recent study using response times (Mazzarella et al., 2018) suggests how they could be made compatible. Replicating the methodology followed by Bonnefon and his colleagues and positing that the prevalence of one-sided readings in their studies was due to a failure to distinguish between (initial) comprehension of the SI and its (subsequent) acceptance, they divided the experimental task into two parts and elicited subjects’ responses as well as reaction times to two different questions. After presenting participants with speech vignettes defined as face-threatening or face-boosting depending on the predicate used, as in (6):

- (6) a. Imagine you gave a speech at a small political meeting. You are discussing your speech with Denise, who was also there. There were 6 other people in the audience that day. You tell Denise that you are thinking about giving the same speech to another group.

- b. Hearing this, Denise tells you that “Some people hated [loved] your speech.”

they asked them to respond to the following two questions in sequence:

- (6) c. Given what Denise told you, do you think that it is possible that everybody hated [loved] your speech? (the “semantic compatibility” question)
 d. Given what Denise tells you, do you think that she means that you should give the speech again to another group? (the “conversational implicature” question)

The researchers’ goal was to tap into participants’ interpretation of the scalar in the two face conditions (their answers to 6c) separately from their assessment of the speaker’s communicative intention (their answers to 6d). They also measured response times to the reading of the utterance containing the scalar (6b) and, separately, to answering the question in (6c).

What they found was that in face-threatening contexts (“Some people hated your speech”), participants overwhelmingly agreed about the speaker’s communicative intention (93% thought Denise means that you should not give the speech again) and that their interpretation of the scalar (45% YES to the semantic compatibility question and 55% NO) played no evident part in that. Conversely, in face-boosting contexts (“Some people loved your speech”), participants were split regarding the speaker’s communicative intention: 64% thought Denise means you should give the speech again while 36% thought she means that you should not; crucially, these answers were interrelated with their interpretation of the scalar. In particular, of those who interpreted the scalar in a one-sided way (‘some and possibly all loved ...’), 87.5% also thought Denise means you should give the speech again, while for those who interpreted it in a two-sided way (‘some but not all loved ...’) this percentage fell to just above half (53%).

What these results suggest, as the authors themselves are quick to acknowledge, is that “the negative valence of the verb “hate” may be stronger than the positive valence of the verb “love” (2018: 5). In fact, as we discuss in the next section, positively valenced words are not always face-boosting. While for negatively valenced words (“hated”) the negative affect encoded in the predicate may be enough to render their use face-threatening, for positively valenced words (“loved”), the face-boost or -threat may well lie not so much in the predicate itself but in what the speaker is thereby implicating (a positive or negative answer to the conversational implicature question). This inherent asymmetry

makes it likely that both Bonnefon et al.'s and Mazzarella et al.'s results tell us more about the processing of *some* in the scope of negatively and positively valenced predicates than about the effect of face-threat/boost as such.

3 Rationale for the present study

3.1 *How should face-threat be construed?*

In Bonnefon et al.'s (2009) and Mazzarella et al.'s (2018) experiments, an utterance such as (2) was defined as face-boosting based on the fact that what is signified by the positively-valenced predicate "loved" does not threaten the face of the listener, while what is signified by the negatively-valenced predicate "hated" does, making an utterance such as (3) face-threatening. In other words, when the speaker wishes to avoid face-threat (making politeness a relevant consideration for the listener), the context was deemed to be face-threatening, while when there is no reason to suspect the speaker is trying to save the listener's face, the context was deemed to be face-boosting.^{7,8}

However, an utterance such as (2) ("Some people loved your poem") can also be face-threatening, in the sense that it can be used to avoid face-threat. For instance, by indicating that 'not everyone loved your poem' it may be covertly communicating something unfavourable for the addressee; or, it may be used to veil the speaker's own dislike of the listener's poem ('Some people, of whom I wasn't one, loved your poem'). A number of factors could lead to this two-sided interpretation of (2), including the alternatives available in the discourse context (e.g., "everyone" would be a relevant response if it were mutually known that a unanimous vote was needed for the poem to win a prize in a poetry competition, or if (2) were uttered in reply to the question "Did people like my poem?") and the intonational contour of the utterance (whether the quantifier "some" or the predicate "loved" is accented).

The possibility of construing (2) as face-threatening along these lines could be seen as an instance of what is known in psychology as 'negativity bias,'

7 Bonnefon et al.'s (2009) second experiment aimed precisely at disentangling the effects of impact on the listener's face from the effects of simply using a (positively or negatively) evaluative term ("loved" vs. "hated") without impacting the listener's face. However, even in this experiment, the story as well as the rest of the utterance context around the scalar term remain exactly the same across the two conditions (face-boosting vs. face-threatening), meaning that, when impact on face is expected, determining the *direction* of impact (boost or threat) essentially comes down to the lexical semantics of the predicate used.

8 Mazzarella et al. (2018) adopt this understanding of face-threatening vs. face-boosting, while Holtgraves & Kraus (2018) opt for the more accurate "non-face-threatening" for the latter.

that is, people's tendency to attach greater importance to negatively-valenced events over positively-valenced ones, which can be explained on evolutionary grounds (Baumeister et al. 2001: 358). Closely related to this is the 'severity effect' observed in experimental politeness research, according to which people tend to overestimate the likelihood of negative eventualities (Bonnefon & Villejoubert, 2006; see also Holtgraves & Bonnefon, 2017: 390–391 and the references therein). In other words, given two possible readings of an utterance, people tend to opt for the one with the more negative consequences (see also Holtgraves, 2014).

What is less commonly observed is that this same severity effect can also obtain in the case of positively phrased remarks such as (2) ("Some people loved your poem"). In a now classic article, Boucher & Osgood argued, on the basis of a sample of 13 languages, for "a universal human tendency to use evaluatively positive (E+) words more frequently ... than evaluatively negative words (E-)" (1969: 1), a finding they called the Pollyanna hypothesis. This finding is actually not opposed to the severity effect but may even be explained by it. Faced with the task of letting someone know a not-so-pleasant truth, speakers choose to put a positive twist on things counting on their interlocutor's inferential abilities to figure out the bitter truth. Such uses are not only kinder on the hearer but also easier for the speaker, who might thereby avoid a certain unpleasantness that would otherwise accrue.⁹ Recent experimental results that positive terms like *happy* are more likely to be negatively strengthened (interpreted as excluding weaker alternatives when negated) than their negative counterparts (Ruytenbeek et al., 2017) confirm listeners' awareness of the attested over-use of positively evaluative words (the "sugar coating" effect of Holtgraves & Bonnefon, 2017: 391), providing cognitive grounds for the behavioural asymmetry in the use of positively vs. negatively evaluative terms noted by Boucher & Osgood half a century ago.

To return to the examples in (2) and (3), not only the one-sided interpretation of *some* (= ... possibly all) in "Some people hated your poem" but also its two-sided interpretation (= ... but not all) in "Some people loved your poem" can be explained by the same tendency of interlocutors to interpret incoming utterances in a negative, rather than positive, light. Such an overarching tendency could well override any face-boosting potential encoded into the utterance via the lexical semantics of the positively-valenced predicate "loved." Consequently, the two-sided interpretation of *some* in (2) (= 'Some but not

9 We thank Larry Horn for this timely reminder of the interconnectedness of speaker's and hearer's face concerns.

all people loved your poem') observed by Bonnefon et al. (2009) could be derived on the strength of the same *face-saving* intention claimed to motivate its one-sided interpretation in (3) (= 'Some and possibly all people hated your poem'). This alternative account calls into question the association between face-threatening contexts and one-sided interpretations claimed by Bonnefon et al. (2009). If what were deemed to be face-boosting contexts can be re-analysed as face-threatening ones (a point we take up in section 3.2), then no unique association between face-threatening contexts and one-sided interpretations can be claimed. Rather, face-threatening contexts can lead to *both* two-sided and one-sided interpretations and both types of interpretation incorporate face concerns.

The difficulty with deciding whether (2) should be construed as face-boosting or face-threatening points to a larger problem with the way the notion of face-threat was experimentally implemented in earlier work. Brown and Levinson (1987: 1) construe face-threat as a non-linguistic perlocutionary effect that *would* occur had the speaker not taken linguistic measures (one of their politeness strategies) to pre-empt it, and suggest three sociological variables (Distance, Power, and Ranking of the imposition) which are jointly computed to estimate the risk of face loss (that subsequently drives the choice between strategies). These three sociological variables are *extra*-linguistic variables that depend on the relationship between interlocutors and the culture at hand (1987: 76). Research since Brown and Levinson concurs that whether an utterance constitutes a threat to the speaker's or the hearer's face (or indeed to both; see above and Turner, 1996) and to what extent depends on many factors, including the relationship and prior interactional history between interlocutors and the sequential ordering of the utterance in the flow of events (for a recent overview, see O'Driscoll, 2017). Reading the face-threat engendered by "Some X-ed" off of the lexical semantics of predicate X underestimates the extent to which estimations of face-threat are susceptible to context and vary across situations and cultures. The alternative construal of (2) as face-threatening proposed above suggests that face-boost and face-threat cannot be assessed on the basis of lexical semantics alone but rather emerge out of the interaction between the utterance's encoded content and its situational context of use. More generally, face concerns are omnipresent: they do not enter the picture only when something bad is literally said.

3.2 *How to conceptualise face-boost?*

Another major difficulty with interpreting the results of earlier studies stems from the fact that face-boost is not a possibility in the framework of Brown & Levinson in which these studies are couched. In Brown & Levinson's work,

face consists of two related aspects, a drive for autonomy (“negative face”) and a drive for affiliation (“positive face”) and all speech acts are intrinsically threatening to one of these two aspects of either the speaker’s or the hearer’s face (1987: 65–68). In other words, in the Brown Levinsonian framework, all that language can do is mitigate the degree of face-threat intrinsically present in all speech acts, leading critics to highlight the *remedial* role of politeness in that framework. There is no indication in that framework that speech acts can genuinely *boost* face when there is no need for redress and of the linguistic strategies that can be used to do that.¹⁰ This has been extensively commented on in the literature since the publication of their essay and prompted several revisions.¹¹

The fact that, in the Brown Levinsonian framework, all acts are intrinsically face-threatening but cannot be genuinely face-boosting leaves us at a loss as to how to conceptualise face-boost in that framework. In Bonnefon et al.’s work, the term “face-boosting” (2009: 3) is applied to utterances such as (2), where no face-saving intention is attributed to the speaker (the speaker makes no attempt to avoid face-threat because no face-threat is assumed to be present in the context to begin with). Yet, in section 3.1, we showed that even those instances can be re-analysed as involving a potential threat to face and the wish to disarm it. A more elegant solution to this problem can be obtained if we pay attention to a little-noticed asymmetry between positive and negative face. The next section outlines this asymmetry and builds on it to propose a possible solution.

3.3 *Positive and negative face and the asymmetry between them*

As already mentioned, in Brown & Levinson’s work, face consists of two related aspects, a drive for autonomy (“negative face”) and a drive for affiliation (“positive face”). An important asymmetry exists between these two aspects: negative face can only be constituted by avoiding imposition, i.e. by *not* performing some (verbal) action, while positive face can be constituted *both* by performing some action (e.g., showing approval, face-boosting) and by not perform-

10 Recall that in Brown & Levinson’s framework, positive politeness which does involve face enhancement is also used to redress a threat. Unlike Brown and Levinson, Robin Lakoff’s earlier paper (1973) did allow for acts directly boosting camaraderie under her Rule 3 politeness: “Be friendly”. We thank Larry Horn for reminding us of this.

11 Early examples include Bayraktaroğlu’s (1991) notion of Face-Boosting Acts and Kerbrat-Orecchioni’s (1997) notion of Face-Enhancing Acts, while Leech’s (2014) notion of pos-politeness is a more recent attempt in the same direction. Extensive work on face-boosting acts has been carried out in Spanish-speaking contexts, for an early attempt see Hernández Flores (1999).

ing a different type of action (e.g., withholding criticism, face-saving). That is because on any understanding of negative face that defines it as freedom from imposition, interaction itself is a kind of imposition (minimally, an imposition on the listener's attention and cognitive resources to process what the speaker is saying). On this count, it is not rationally possible to constitute negative face except by avoiding interaction, since interaction is itself a kind of imposition. This also justifies why in their model Brown and Levinson have included silence (avoidance of verbal behavior) as the most polite strategy (to be used when the estimated threat to face is highest). If we take this (strict) Brown & Levinsonian line, then it follows that negative face can only be constituted by avoiding threatening it (i.e. by mitigating behavior which imposes on it or by avoiding interaction altogether), while positive face can be constituted both by avoiding threatening it (e.g., avoiding criticism) and by directly enhancing it (e.g., showing approval).

All of the stories tested in the earlier studies, as well as in our own (see section 4), involve the listener's positive face.¹² This is in line with Bonnefon and colleagues' definition of face as "a sense of *positive* identity and public self-esteem that all humans project and are motivated to support in social interactions" (2009: 250; emphasis added), which corresponds to Brown and Levinson's definition of positive face. Indeed, going back over their experimental materials, one finds that what they termed face-threatening contexts are precisely those in which the speaker works to pre-empt a possible threat to positive face (in this sense, a more appropriate term for these contexts might have been 'face-saving'), whereas their face-boosting contexts are those in which no threat to face is assumed to be present to begin with. Others (e.g., Holtgraves & Kraus, 2018) have opted for the term "non-face-threatening" to describe the latter. However, this alternative term goes against the widely held view that "[p]oliteness, [...] is not a sometime thing" (Fraser 1990: 233)—or, in Scollon and Scollon's words, "there is no face-less communication" (1995: 38).

Adopting Scollon & Scollon's view, we opt for a different solution. Specifically, we regard as 'face-boosting' those contexts in which the speaker's act *enhances* the hearer's positive face (e.g., by expressing affiliation, solidarity, approval or admiration for the hearer), while we reserve the term 'face-threatening' for those contexts in which the speaker's act *threatens* the hearer's positive face (e.g., by not withholding criticism or bad news). While our

12 The following stories from our main study additionally involve avoiding threat to the speaker's or the hearer's negative face: Some-subject story 3, Some-object story 3, Or stories 2 and 3, Often story 4, and Possible stories 2 and 3.

definition of face-boosting contexts has no counterpart in Brown and Levinson's work, our definition of face-threatening contexts is in line with their construal of face-threat as related to Goffman's notion of virtual offence. They write in this regard:

politeness, like formal diplomatic protocol ... *presupposes* that potential for aggression as it seeks to disarm it, and makes possible communication between potentially aggressive parties. But how? Goffman suggests that it is through the diplomatic fiction of the *virtual offence*, or 'worst possible reading' of some action by A that potentially trespasses on B's interests, equanimity or personal preserve. By orienting to the 'virtual offence', an offender can display that he has the other's interests at heart.

BROWN and LEVINSON, 1987: 1; emphasis added

In other words, an act must first engender a potential for face-threat in order for its linguistic clothing to have a role to play in mitigating that threat. That is why the degree of face-threat inherent in an act is estimated by aggregating the values of Distance and Power between interlocutors and Ranking of the imposition—three *extra*-linguistic variables—*before* that act is cast linguistically in terms of one of their politeness strategies able to proportionately redress that threat. The experimental scenarios we term face-threatening/-boosting in our study aimed precisely to set up this potential for face-threat/-boost *before the target utterance is uttered*, to allow us to observe how scalar terms are interpreted in those circumstances.¹³ Our choice to implement face-boost vs. face-threat with respect to the two possibilities available for positive face (*enhancement* and *threat*) achieves greater symmetry between the two types of contexts and additionally acknowledges the omnipresence of face in both types of context (see section 3.1).

13 Since our goal was to set up experimental scenarios in which face-threat is imminent, it might seem as if this is no different from impoliteness. However, note that in impoliteness, the face-threat is linguistically actualized (the offence does not remain virtual), while in our scenarios it is merely expected (see section 3.1). This is because, unlike the linguistic (politeness) strategies discussed by Brown & Levinson, the language used in our stimuli is "neutral" with respect to face-threat and -boost (it does not encode either, as Bonnefon and colleagues' stimuli are claimed to do). This in turn allows us to observe how face concerns influence the interpretation of language without assuming anything about what the language (the predicate used) itself does in these circumstances.

3.4 *Residual considerations*

The fact that face-threat is best construed as a matter of situational context (section 3.1) highlights the importance of a few other situational factors, starting with participant gender. The nexus of im/politeness and gender has been primarily studied by sociolinguists (for a recent overview, see Chalupnik et al., 2017; for a recent experimental study using ERPs, see Jiang & Zhou 2015: 257–259), who found that women tend to use more indirect strategies overall, one explanation for which is that they tend to construe the same acts as more face-threatening than men do.¹⁴ The fact that samples were not balanced for gender in any of the previous experiments (women outnumbered men by as much as 4.5 to 1 in Bonnefon et al., 2009; Feeney & Bonnefon, 2012; Holtgraves, 2014; Holtgraves & Perdeu, 2016; men outnumbered women in Mazzarella et al., 2018) could be yielding a somewhat biased picture in this regard. To alleviate this concern, a gender-balanced sample was used in the study reported below.

Two additional considerations ought to be kept in mind, although we do not address them in the present study. The first pertains to the type of speech act performed by the utterance in which the scalar is embedded: (2) can be construed as a compliment (or, alternatively, as proposed in 2.2 above, a gentle let-down), (3) as a criticism (or as a warning). The type of speech act participants took the speaker to be performing each time may have affected their interpretation of the scalar above and beyond the face-boosting or face-threatening orientation of the sentential context; or, this face-orientation itself may be emanating not so much from the lexical semantics of the predicate, which researchers controlled for, but from the type of speech act performed, for which they did not control. In short, controlling for the type of speech act in which scalars are embedded is necessary before any claims about how face affects interpretations of scalar terms can be asserted.

The second consideration concerns the range of languages and scalar terms tested. The initial claim that face-threatening contexts favour one-sided interpretations was based on results from two closely related languages (English and French)¹⁵ and three scales (<some, all>, <or, and>, <possibly, probably>). Given the small number of languages and terms studied, as well as the lack of direct correspondence between translational equivalents of quantifiers raised

14 Other explanations of course are possible, from women's greater adherence to standard language norms (Hudson, 1996: 195), of which polite language use can be considered one (Watts, 2002), to their use of indirectness as a means of social filtering and bonding (Morgan, 1991).

15 One study (Pighin & Bonnefon, 2011, study 2) involved Italian pregnant women.

in other studies (Pouscoulous et al., 2007; Banga et al., 2009), expanding the range of languages and scalar terms tested is necessary. When Holtgraves & Kraus (2018) expanded the range of terms to five, they found important differences between the three 'logical' terms (*some, sometimes, probable*) and the two evaluative ones (*like, good*). As a first step toward more comprehensive study of the nexus of face and the interpretation of scalar terms, in this article, we take up the task of testing with a gender-balanced sample and an expanded range of terms, leaving the type of speech act performed and languages other than English to future research.

4 The present study

The present study addresses the following research questions: (1) How are scalar terms interpreted when embedded in face-threatening vs. face-boosting contexts? (2) Do different scalar terms behave alike in this respect or is there variation among them?

We adopt a definition of lexical scales according to which two expressions form a scale if (a) they are equally lexicalized, and (b) they can be ranked according to some metric which orders alternate values as higher or lower and is salient to both speaker and hearer. This definition retains the requirement of lexicalization from previous studies (e.g., Horn, 1972), while relaxing the requirement of strict entailment (along the lines of Hirschberg, 1991). In the present study, the following lexical scales were tested: *⟨some, all⟩* in both Subject and in Object positions,¹⁶ *⟨or, and⟩*, *⟨often, always⟩*, *⟨possible, likely⟩*, *⟨like, love⟩*, *⟨good, excellent⟩*, *⟨unwell, sick⟩*, *⟨misguided, illegal⟩*, *⟨assertive, bossy⟩*, *⟨misleading, lying⟩*. These represent a wider range than previous studies of scalars in face-boosting/-threatening contexts and include items from a variety of scales; the last four represent contextually set up scales identified from press articles and online searches using the heuristic "A but not B".

To minimise the possibility that results will be biased by the linguistic context of the utterance in which a term is appearing, each term was tested using four different utterances, for a total of 32 different utterances (for the full set of stimuli, see Appendix 1).¹⁷ We also included control stimuli, which were

16 We tested Some-subject and Some-object separately because previous research found differences in their acquisition (Armon-Lotem, 2008: 153).

17 The exception here are the last four terms (*unwell, misguided, assertive, misleading*) which were tested as a group, using one utterance per term. Although these are all taken from Non-Entailment Scales, in the results below we discuss them separately given they vary

missing from previous studies. For each utterance, two story versions were constructed. These were intended to provide a face-boosting or face-threatening context for the utterance through the story content. In other words, contrary to the earlier research in which the context was kept stable and the utterances alternated, and more along the lines of Holtgraves & Kraus (2018), we kept the utterance stable and alternated the story versions in which it was embedded. For example, one of the utterances for the adjective *good* was “You have a good sense of rhythm.” This was alternately embedded in one of the story versions below:

- (7) Face boosting: Paul has his first guitar lesson with his new teacher. Paul plays a portion of a song so the teacher can get a sense of his abilities. The teacher, who is eager for new students, tells Paul, “You have a good sense of rhythm.”
- (8) Face threatening: Paul is playing guitar in a competition with a notoriously strict panel of judges. After Paul plays his song, the first judge is silent for a while and then mutters, “You have a good sense of rhythm.”

In these examples, the implicature is that the proposition containing the stronger alternative, ‘You have an excellent sense of rhythm’, is not true. If this implicature is generated, then the speaker would be taken to mean that Paul has a good but not excellent sense of rhythm.

4.1 Norming study

The norming study was designed to ensure that the story versions to be used in the main study were perceived by participants as actually face-boosting and face-threatening in the expected direction. As a reminder, we define ‘face-boosting’ as genuinely enhancing the hearer’s positive face and ‘face-threatening’ as potentially threatening the hearer’s positive face (see section 3.3). In the norming study, participants saw one of the two versions of a story *without the final utterance containing the scalar term* (“You have a good sense of rhythm” in examples (7) and (8) above) in order to test the face orientation of the context itself.

on some dimensions identified as important in recent work on adjectival scales (Benz et al., 2018; Gotzner et al., 2018; Leffel et al., 2019).

4.1.1 Participants

Sixty participants recruited from Amazon Mechanical Turk took part in the norming study. Only individuals with IP addresses from the US were allowed to participate.

4.1.2 Materials

The 64 story versions (2 versions for each of 32 utterances) without the utterance containing the scalar were counterbalanced and sorted across three lists such that each participant saw only one version of each story. Each list contained 21 or 22 story versions; 15 participants saw each list. The order of items was randomized for every participant.

4.1.3 Procedure

After having read one of the story versions (e.g., either (7) or (8) in the examples above), participants were asked to indicate on a five-point Likert scale, only the end-points of which were labelled as 1 = “very unlikely” and 5 = “very likely”, “How likely is it that S will say something nice to H?”¹⁸

4.1.4 Results

To ensure that our story versions were face-boosting/-threatening in the expected directions, mean ratings for Boost and Threat versions were required to be higher or lower than three (the midpoint of the five-point scale used), respectively. Eleven story versions in the initial set failed to meet this criterion; these were revised and presented to 15 new participants in a fourth list. After revisions, all 64 story versions met the criterion. The average rating across all Boost version stories was 4.37 ($sd = 0.41$); the average rating across all Threat version stories was 2.07 ($sd = 0.40$). This norming study ensures that our contexts were indeed face-boosting and face-threatening in the expected direction.

¹⁸ In a few instances, the alternative question “How likely is it that S will say something that H wants to hear?” was used if that was more suitable. Both of these questions manipulate expectations of face-threat/boost rather than face-threat/boost itself. This is in line with Brown & Levinson’s conceptualisation of face-threat in terms of ‘virtual offence’, since linguistic strategies in their model are meant to ensure precisely that face-threat doesn’t ultimately materialise (section 2.3). In other words, it is enough for a situation to engender expectations of face-threat to be characterized as face-threatening (and *mutatis mutandis* for face-boost).

4.2 Main study

4.2.1 Participants

162 participants ($F = 80$, mean age = 35) were recruited from MTurk for the main study. Participants who had taken part in the norming study were not allowed to participate in the main study. Only participants with US IP addresses were allowed to take part. Participants were asked to report what language they speak at home; 161 participants reported speaking English only and one participant reported English and Italian.

4.2.2 Materials

The main study presented the story versions complete with the final utterance containing the scalar, as in the examples in (7) and (8) above. The 64 story versions were counterbalanced over four lists; each list contained 16 critical stories and 7 filler stories for a total of 23 stories. Each participant saw a story in only one of its versions (face-boosting or face-threatening). 40 or 41 participants saw each list. The order of items was randomized for each participant.

4.2.3 Procedure

This time, participants were asked to judge on a five-point scale how likely it is that the speaker means the stronger term in the scale. Again, only the end-points of the scale were labelled, as “very unlikely” (=1) and “very likely” (=5). A full example is given in (9) (*good*, face-threatening version):

- (9) Paul is playing guitar in a competition with a notoriously strict panel of judges. After Paul plays his song, the first judge is silent for a while and then mutters, “You have a good sense of rhythm.”

How likely is it the judge means that Paul has an excellent sense of rhythm?

1 2 3 4 5
Very unlikely Very likely

4.2.4 Results

With ordinal data like the type our experiment generates, rating data typically are sparse between participants and items, and building a maximal random effects model introduces either convergence problems or other issues with the estimation of fixed effects. For this reason, we used Bayesian statistics to analyse our results (Nicenboim & Vasishth, 2016; Kimball et al., 2018). While frequentist approaches aim to address a null hypothesis that a difference between conditions is 0, Bayesian approaches aim to estimate the size of the difference

TABLE 1 Model results for leave one out cross validation (LOO) information criterion (IC)

Model	LOO IC	SE
Model 1: Intercept Only	7120.50	65.04
Model 2: Condition	6912.89	67.94
Model 3: Scalar	6960.45	70.80
Model 4: Condition + Scalar	6663.98	75.03
Model 5: Condition * Scalar	6662.13	73.62

along with a degree of certainty about this estimate (the “credible interval”). They are thus seen as providing a more fine-grained approach and are considered especially helpful in hierarchical mixed-effects models.

Following the methodology described in Bürkner and Vuorre (2018) for ordinal models and using the *r* package *brms* (Bürkner, 2017), fully Bayesian maximal random effects models were fit for the following five models: an independence (null) model (Model 1: Intercept Only), one-factor models for Condition Only (Model 2) and Scalar Only (Model 3), a two-factor model (Condition+Scalar, Model 4) and a two-factor model with an interaction term (Condition * Scalar, Model 5). The independence model assumes that there are no relationships between items, thus providing a baseline from which to compare the relative fit of all the other models. One-factor models assume that all items constitute one underlying construct. These models were used to test the rival hypothesis that one factor (Condition or Scalar) provides the best fit to the data. The two-factor and two-factor-with interaction models test the hypothesis that the items measure two distinct, potentially interacting, constructs. The leave one out cross validation information criterion (LOO IC) gives a comparative measure of model fit and is described in Vehtari, Gelman and Gabry (2017). This method arrives at an Information Criterion value by iteratively training a function based on all but one data point and validating on the excluded point. A comparatively lower LOO IC value indicates that one model captures more information of the data than another. The standard error of each LOO IC gives us a measure of how much noise there is in the data; in comparing two models, if the standard error is larger than the difference in LOO IC, the data lack the information to distinguish between the two models. The results are given in Table 1.

For our purposes there are five models under consideration with the two possible fixed effects, Condition (C) and Scalar (S). Table 2 represents the dif-

TABLE 2 Leave one out cross validation comparison for Models 1–5

Model comparison	Difference in LOO IC	SE	Difference/SE
IO vs C+S	456.53	49.38	9.25
IO vs C	207.61	32.64	6.36
IO vs S	160.05	35.02	4.56
C+S vs C*S	1.85	7.30	.25
S vs C + S	296.47	42.02	7.05
C vs C + S	248.91	38.90	6.39

ferences in LOO between these models, with standard errors. In the Intercept Only (Model 1) versus Condition + Scalar (Model 4) comparison, the positive difference indicates the first model is less informative than the second model—that is, the model with no variables is less informative than the model with Condition and Scalar as fixed effects. The standard error, 49.38, is several times smaller than the difference, 456.53, indicating that the C+S model (Model 4) reliably captures more information than the Intercept Only model (Model 1). Comparing the C+S (Model 4) and C*S (Model 5) models, the difference is 1.85, but the standard error is 7.30, meaning there is no evidence that these models are informationally distinct and, therefore, no evidence that the interaction term should be included. Based on these comparisons, the C+S model, without the interaction term, is the optimal model and will be used for analysis.

Table 3 presents estimates for each parameter of the selected Condition + Scalar model (two parameters for Condition, Threat vs. Boost, and 11 parameters for Scalar).¹⁹ The logit linked was used, so estimates greater than 0 indicate that a higher rating (a one-sided interpretation of the scalar) is more likely compared to the reference category ('Boost' for Condition and 'possible' for Scalar). Bayesian models do not produce a direct analogue to p-values in frequentist models, but the 95% credible intervals, with a 2.5% Lower Credible Interval (LCI) and 97.5% Upper Credible Interval (UCI), for each parameter can be tested to see if there is overlap with 0 (i.e. no difference). If the intervals do not overlap with zero, that means that a null effect is not within 95% of credibility; if the intervals do overlap with zero, then a null effect is within 95% of credibility. The greyed cells in Table 3 indicate parameters for which the 95%

19 For the purposes of statistical analysis, we analysed the four non-entailment scales separately in order for each parameter to correspond to a different lexical item.

TABLE 3 Parameter estimates and credible intervals from C+S model

Variable	Estimate	Est error	2.5% LCI	97.5% UCI
ConditionThreat	-0.5	0.4	-1.3	0.3
ScalarGood	-0.5	0.6	-1.6	0.6
ScalarLike	-0.6	0.6	-1.7	0.6
ScalarNES1—unwell	3.3	1.1	1.2	5.6
ScalarNES2—misguided	-1.5	1.1	-3.6	0.6
ScalarNES3—assertive	1.3	1	-0.7	3.4
ScalarNES4—misleading	1	0.9	-0.8	2.7
ScalarOften	-0.7	0.6	-1.8	0.5
ScalarOr	-3.9	0.7	-5.3	-2.6
ScalarSomeO	-3.8	0.6	-5	-2.6
ScalarSomeS	-4.2	0.6	-5.5	-3

credible intervals do not overlap with zero. As Table 3 shows, *or*, *some*-object, *some*-subject and the non-entailment scalar *unwell* yield non-null effects in the data given the model.

In Figures 1 and 2 below, the 95% credible intervals are plotted from the C+S model for Condition and Scalar, respectively. The y-axis represents the model-estimated probability of each rating (1= two-sided interpretation is most likely; 5 = one-sided interpretation is most likely) where the circle is the estimate and the bars indicate the 95% credible intervals. While Figure 1 shows that there is not sufficient evidence for a consistent overall effect of Condition, Figure 2 shows that three of the scalar terms—*or*, *some*-object, and *some*-subject—had a credible tendency for lower ratings, and one of the non-entailment scales (*unwell*, sick) had a credible tendency for higher ratings.

5 Discussion

To investigate the effect of situational context on the interpretation of scalar terms, we manipulated the face-boosting/-threatening orientation of the context and obtained participant judgements regarding the likelihood that different scalar terms would be interpreted in a one-sided way (including a stronger alternative, i.e. without a SI). For two of these (*or*, and *some*), we found that they tend to induce two-sided readings in all contexts (our prediction that *some* may behave differently in subject and object positions was not confirmed). These

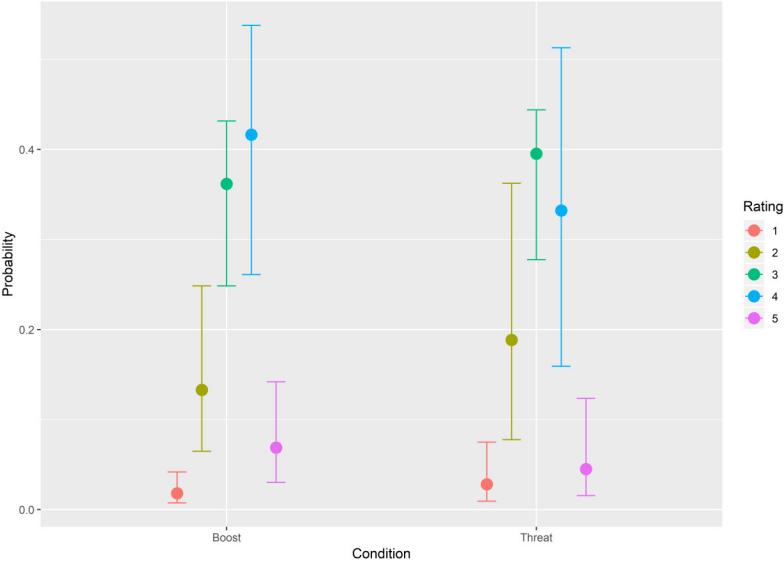


FIGURE 1 C+S model estimates for condition and 95 % credible interval

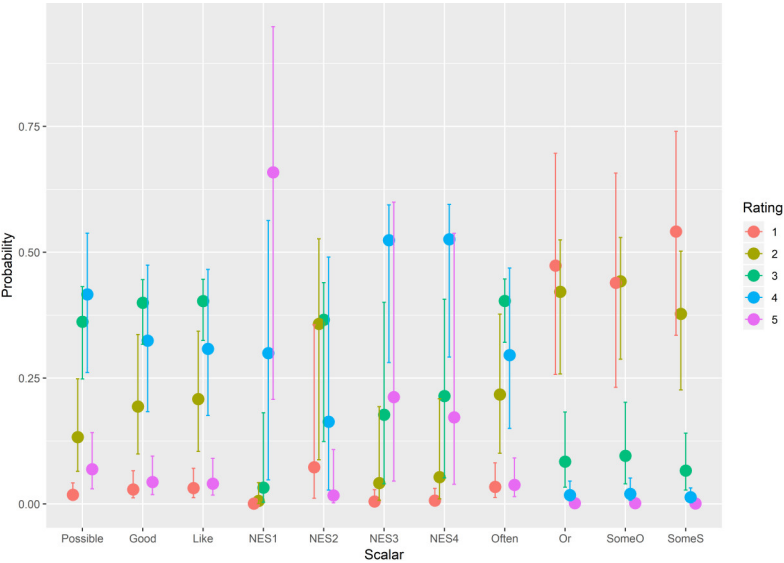


FIGURE 2 C+S model estimates for scalar plus 95 % credible interval

results are in line with those of studies on scalar diversity (section 2.1). These terms seem to generally induce scalar implicatures irrespective of situational context and are less prone to contextual enrichment. However, we were unable to replicate the finding of earlier studies that investigated these terms in face-boosting vs. face-threatening contexts that face-threatening contexts lead to one-sided interpretations (section 2.2). Nor were we able to establish such an effect for the remaining terms. Moreover, for these terms (with the exception of *unwell*),²⁰ our results do not allow establishing a preference for one-sided vs. two-sided readings, suggesting that their nonce context of occurrence plays a major role in their interpretation.

When comparing these results with those of previous studies on scalar diversity and on the effect of face-threat/-boost on scalar interpretation, it is important to highlight how our experimental setup was different. The studies on which the original claim that face-threatening contexts warrant one-sided readings of scalars was based investigated terms that generally tend to induce SIs (*some, or*) in one situational context per term and determined face-orientation (boost vs. threat) based on the lexical semantics of the predicate in whose scope the scalar appeared. This homogeneity in their materials may have inadvertently resulted in the consistency in scalar interpretations that they found, effectively showing how the specific scalar terms investigated behave in the scope of positively- vs. negatively-valenced predicates rather than how scalar terms more generally are interpreted in face-boosting vs. face-threatening contexts. On the contrary, we investigated each term in four different sentential contexts, each of which was alternately embedded in two different story versions. Story versions differed in face-orientation understood as produced in context (for which we controlled by means of norming) but also in other ways (number and identity of interacting characters; type of speech act performed; additional aspects of face threatened, cf. fn. 12), for which we did not control.²¹ Could these additional differences between our story versions be behind the inconclusiveness of our results with respect to face-orientation?

Without denying that these further aspects affected our results, we believe these differences are not wholly responsible for the absence of non-null effects. If that were so, these differences should have equally affected all scalars. How-

20 *Unwell* showed a reliable tendency for one-sided interpretations ('unwell and possibly sick'). Since *unwell* was one of our Non-Entailment Scales terms that were tested in one sentential context only, and given the register difference between the two scalesmates, we do not think this finding warrants generalisation without further testing.

21 A post-hoc classification of our story versions according to seven factors (speaker-ad-

ever, given the robust results obtained for *some* and *or*, we believe that the variable ratings obtained with the other terms are a genuine reflection of scalar diversity and, taking on board the notion of UBELE from Sun et al. (2018), further hypothesise that, compared with *some* and *or* which are subject to local enrichment, the remaining terms we tested are more likely to be globally enriched, calling for fully-fledged Gricean reasoning that can variably produce one-sided or two-sided readings depending on the interlocutors' goals in the specific context of utterance. If this hypothesis is correct, it would explain at least some of the story variance in our results.

Given the multiple contextual parameters that can affect assessments of face-boost/-threat and the fact that such assessments are ultimately subjective to a degree, studies manipulating these notions are bound to be less controlled than experimental pragmatic studies manipulating other parameters affecting interpretation. Yet, we believe experimentation in this field is possible (indeed, necessary!), if only because the underlying face dynamics—despite being ever-present—is all too often ignored. In the rest of this section, we highlight two methodological aspects of our study that follow from our theoretical commitment to the notions of face and face-threat/-boost as understood in the current im/politeness literature and that we believe future studies experimenting with im/politeness should bear in mind. Given the subjectivity in face-related judgements, future studies are also likely to benefit from eliciting demographic and other (e.g., attitudes) information about the participants themselves, to test the possibility that different types of participants interpret different scalar terms differently.

5.1 *The complex interface between face and affect*

As pointed out in section 3.2, face-boosting is not a possible option in the Brown Levinsonian framework. Given this, it is unclear how “face-boosting” should be understood when applied to examples such as “Some people loved your poem” (from Bonnefon et al., 2009).²² To avoid problems with determining what is

dressee pair; one-to-one vs. multi-party interaction; length in words; illocutionary force; interlocutor familiarity; leisure vs. work; private vs. professional) showed that, as regards the two versions of each story, these factors were stable half the time (16/32 stories). Of the remaining 16 stories, 10 were characterized by a change in speaker-hearer pair, which in 3 cases was also a change between multi-party and dyadic conversation. In terms of length in words (range: 16–71, mean 45), the two conditions were more than ten words apart in 2 cases; while another 4 stories were characterized by a difference in illocutionary force between the two conditions.

22 The extremely low ratings for Kindness obtained for face-boosting contexts in their third

face-boosting in our own study, we opted for a construal of face which makes it possible to actively enhance positive face, in line with subsequent research on this topic (Bayraktaroğlu, 1991; Kerbrat-Orecchioni, 1997; Hernandez Flores, 1999; Leech, 2014). Our phrasing of the question used in the norming study (“How likely is it that S will say something nice to H?”) is indicative of this. “Saying something nice” amounts to showing approval, admiration, solidarity, or inclusion—all notions relating to positive face as defined by Brown and Levinson (1987: 61).

We opted for this phrasing for a couple of reasons. First, we wanted to avoid asking participants directly about ‘face,’ as the technical understanding of this term in the im/politeness literature and as a lay term can be different (O’Driscoll, 1996: 8; Terkourafi, 2007: 318–325). Second, this phrasing highlights another aspect of face, namely, its link with affect. Both Goffman (1967: 23) and Brown & Levinson (1987: 28) acknowledged that face has an inextricable affective aspect—suffice it to think how face loss can both result from, and cause, negative affect (e.g., anger). In connection with this, note that affective considerations are not absent from the earlier experiments either. To interpret “Some people hated your poem” ((3) above from Bonnefon et al., 2009) as one-sided, the hearer must presume that the speaker is positively predisposed toward them. The introduction of the target utterance by the phrase “one fellow member confides to you that” in their experiment, with the intimacy signalled by the predicate “confides,” could have facilitated this presumption. If, on the contrary, a negatively affective stance were attributed to the speaker (say, because the speaker is a competitor or known to relish in giving bad news), the speaker’s assertion of the existence of people who hated your poem would not necessarily give rise to the one-sided interpretation. This could be the result of epistemic vigilance (Sperber et al., 2010), implemented through trust in the speaker and in the plausibility of the information in Mazzarella et al.’s (2018) experiments. With a speaker whose motives we find questionable, epistemic vigilance could lead us to accept the implicature (the two-sided reading) rather than reject it. In other words, *default* attributions of affect motivated by the language used in their vignettes may have inadvertently influenced the derivation (or not) of the SI in the earlier experiments.²³

experiment (Bonnefon et al., 2009: 254) could be taken as evidence of participants’ perplexity in this regard.

- 23 Drawing a link between politeness and affect, in a study of *modus ponens* (a non-scalar phenomenon), Demeure et al. (2009) manipulated dis/liking between interlocutors and found that “for the sake of politeness, people use ambiguous statements to correct listeners who dislike them” (2009: 264). Brown & Levinson see such ‘affective distance’ or

The larger P300 responses to the one-sided continuation obtained by Holtgraves & Kraus (2018), with additional larger P200 responses for the positively valenced terms (*good* and *like*) in face-threatening contexts, further suggest that the effects of face-threat are discernible on the emotional/affective plane. While face is not affect *per se*, the (expected/ projected) impact of a situation on the listener's emotions is one way of making tractable this otherwise elusive notion. Given this close link between face and affect, asking about the expected impact on the listener's feelings of what might be said next offers a tangible way of determining the face orientation of the *context* (face-boosting vs. face-threatening) without asking about face directly.

That said, it is important to stress that *face is not affect per se*. Rather, as a mix of self-presentational concerns (our concern for our image in the eyes of others), it is a complex outcome of sociological variables instantiated in the situation. It would therefore be inappropriate to manipulate face-threat/-boost by manipulating affect alone (e.g., by consistently changing a smile to a frown in our story versions to suggest a different emotional state of the speaker). Rather, holistically manipulating the situation is what is required to properly manipulate face-threat/-boost, according to Brown & Levinson's framework. While this may have introduced added complexity to our stimuli, this move is necessary to ensure that what is manipulated is face-threat/-boost and not (just) affect. Our methodological choice to conceptualise face orientation as emergent from the situational context and ascertain its direction through norming further guarantees a modicum of consensus that our face-boosting story versions were indeed those in which enhancement of the listener's positive face was expected, and *mutatis mutandis* for the face-threatening ones.

5.2 *The importance of speaker meaning*

On Grice's definition of implicature, implicatures are a matter of speaker meaning and are generated by virtue of the speaker's intention to convey their contents—or, in the case of Generalized Conversational Implicatures (Grice, 1989: 37), her doing nothing to prevent this. An important point often missed in this connection is whether the speaker actually believes the implicature to be true. This is especially relevant in the case of politeness implicatures: it is the speaker's *intention to have the listener believe* that 'Some but not all people hated your poem' that makes (3) a polite utterance, not whether the speaker actually believes this implicature herself. The listener who is grateful to her for

'liking' as potentially another sociological variable to be added, alongside P, D and R, to their formula estimating the risk of face-threat (1987: 16), rather than as an ingredient of face itself.

softening the blow is grateful not because he takes her to believe that not all people hated his poem but for her considerateness in putting it that way.

Previous research on scalars utilized an inference task asking participants to indicate the likelihood that the stronger term is true, given the speaker's utterance (see (4) above). Both the inference task (Van Tiel et al., 2016) and the verification task also common in SI research (Doran et al., 2009) focus on whether the implicature is true.²⁴ Unlike this research and closer to the approach taken by Mazzarella et al. (2018), we asked participants whether the speaker could have meant the stronger term (see (9) above). Focusing on the speaker's intention means that it is not enough for the inference to be available in the context. What we wanted to know was whether the SI was something that participants thought was meant to be recognized by the listener as speaker-intended. Given Brown and Levinson's explicit adoption of a Gricean framework, in which politeness is precisely supposed to be a matter of intention (1987: 5),²⁵ we believe this is a more appropriate way of finding out whether politeness is a motivating factor for the implicature in these cases.²⁶

6 Concluding thoughts

The possibility that euphemism (a type of politeness) can be relevant to scalar interpretations has been suggested in the context of negative strengthening (Horn, 1992), and more recent work, while focusing on the importance of scale structure, does not reject this possibility (cf. Leffel et al., 2019: 5). Indeed, the lexical semantics of terms used in our study may be part of the reason for the inconclusive results with terms other than *some* and *or*. While our results do not support those of earlier research regarding the direction of inference in face-boosting vs. -threatening contexts, they confirm research on scalar diversity, which has shown that hearers draw SIs less with semantically rich and

24 The questions asked by Doran et al. (2009) and Van Tiel et al. (2016), respectively, were: "Is the speaker's statement true or false?" and "Would you conclude that, according to the speaker, the stronger term is true?"

25 This assumption has been challenged by subsequent research (Arundale, 1999; Terkourafi, 2001, 2003, 2008; Haugh, 2003).

26 Bonnefon et al.'s (2009) third experiment actually aimed to test whether listeners attributed a polite intention to a speaker who used "some" when in fact 'all' was true. However, in that experiment, participants were told *a priori* whether 'all' or 'not all' was true. That is, they were not given the option of deriving the SI (or not) depending on whether they attributed a polite intention to the speaker (or not). As such, this experiment does not establish a direct link between the speaker's polite intention and the *derivation* of the SI.

with vague terms. They thus add to this expanding literature from a different perspective, that of situational context.

Looking ahead, it is possible that, for terms that call for global enrichment, the direction of enrichment may also differ in each situational context, even if these contexts share the same face-orientation. Moreover, the finding that not all scalars are equally sensitive to face concerns raises the possibility that scalars that are more “open to interpretation” do more “politeness work” than those that are not and are therefore more likely to be used to this effect by speakers and interpreted in this light by hearers. Future work in this vein could help us explain which linguistic expressions carry the brunt of doing politeness work, how, and why.

Acknowledgements

Earlier versions of this work were presented at the 15th IPrA conference in Belfast, at departmental seminars at the Universities of Manchester and East Anglia, and at INPRA2018 in Cyprus. We thank the audiences on these occasions for their feedback. We are especially indebted to Larry Horn, Jonathan Culpeper, and Bart Geurts as well as the anonymous reviewers of earlier drafts for their careful reading, which helped us clarify our claims. All remaining errors are our own.

References

- Armon-Lotem, Sharon. 2008. Subject object asymmetry in children's comprehension of sentences containing logical words In P. Guijarro-Fuentes, P. Larrañaga & J. Clibbens (eds.), *First Language Acquisition of Morphology and Syntax: Perspectives across Languages and Learners*, 137–159. Amsterdam: John Benjamins.
- Arundale, Robert. 1999. An alternative model and ideology of communication for an alternative to politeness theory. *Pragmatics* 9: 119–153.
- Banga, Arina, Ingeborg Heutinck, Sanne Berends & Petra Hendriks. 2009. Some implicatures reveal semantic differences. In B. Botma & J. van Kampen (eds.), *Linguistics in the Netherlands*, 1–13. Amsterdam: John Benjamins.
- Baumeister, Roy, Ellen Bratslavsky, Catrin Finkenauer & Kathleen Vohs. 2001. Bad is stronger than good. *Review of General Psychology* 5: 323–370.
- Bayraktaroğlu, Arn. 1991. Politeness and interactional imbalance. *International Journal of the Sociology of Language* 92: 5–34.
- Benz, Anton, Carla Bombi & Nicole Gotzner. 2018. Scalar diversity and negative

- strengthening. In U. Sauerland & S. Solt (eds.), *Proceedings of Sinn und Bedeutung* 22: 191–203. Berlin: ZAS.
- Bonnefon, Jean & Gaelle Villejoubert. 2006. Tactful or doubtful? Expectations of politeness explain the severity bias in the interpretation of probability phrases. *Psychological Science* 17: 747–751.
- Bonnefon, Jean, Aidan Feeney & Gaelle Villejoubert. 2009. When some is actually all: Scalar inferences in face-threatening contexts. *Cognition* 112: 249–258.
- Bonnefon, Jean, Ethan Dahl & Thomas Holtgraves. 2015. Some but not all dispreferred turn markers help to interpret scalar terms in polite contexts. *Thinking and Reasoning* 21: 230–249.
- Boucher, Jerry & Charles Osgood. 1969. The Polyanna hypothesis. *Journal of Verbal Learning and Verbal Behavior* 8: 1–8.
- Breheny, Richard, Napoleon Katsos & John Williams. 2006. Are generalised scalar implicatures generated by default? An on-line investigation into the role of context in generating pragmatic inferences. *Cognition* 100: 434–463.
- Brown, Penelope & Stephen C. Levinson. 1987. *Politeness: Some Universals in Language Usage*. Cambridge: Cambridge University Press.
- Bürkner, Paul. 2017. brms: An R package for Bayesian multilevel models using Stan. *Journal of Statistical Software* 80: 1–28.
- Bürkner, Paul & Matti Vuorre. 2018. Ordinal regression models in psychology: A tutorial. *PsyArXiv*. September 15, 2018.
- Carretié Luis, Francisco Mercado, M. Tapia & J. Hinojosa. 2001. Emotion, attention, and the 'negativity 986 bias', studied through event-related potentials. *International Journal of Psychophysiology* 41:75–85.
- Chalupnik, Małgorzata, Christine Christie & Louise Mullany. 2017. (Im)politeness and gender. In J. Culpeper, M. Haugh & D. Kádár (eds.), *The Palgrave Handbook of Linguistic (Im)politeness*, 517–537. London: Palgrave.
- Chemla, Emmanuel & Raj Singh. 2014a. Remarks on the experimental turn in the study of scalar implicatures, Part 1. *Language and Linguistics Compass* 8: 373–386.
- Chemla, Emmanuel and Raj Singh. 2014b. Remarks on the experimental turn in the study of scalar implicatures, Part 11. *Language and Linguistics Compass* 8: 387–399.
- Degen, Judith & Michael K. Tanenhaus. 2015. Processing scalar implicatures: A constraint-based approach. *Cognitive Science* 39: 667–710.
- Delplanque, Sylvain, Marc Lavoie, Pacal Hot, Laetitia Silvert & Henrique Sequeira. 2004. Modulation of cognitive processing by emotional valence studied through event-related potentials in humans. *Neuroscience Letters* 356:1–4.
- Demeure, Virginie, Jean Bonnefon, & Eric Raufaste. 2009. Politeness and conditional reasoning: Interpersonal cues to the indirect suppression of deductive inferences. *Journal of Experimental Psychology* 35: 260–266.
- Doran, Ryan, Rachel Baker, Yaron McNabb, Meredith Larson & Gregory Ward. 2009.

- On the non-unified nature of SI: An empirical investigation. *International Review of Pragmatics* 1: 1–38.
- Feeney, Aidan and Jean Bonnefon. 2012. Politeness and honesty contribute additively to the interpretation of scalar expressions. *Journal of Language & Social Psychology* 32: 181–190.
- Fraser, Bruce. 1990. Perspectives on politeness. *Journal of Pragmatics* 14: 219–236.
- Geurts, Bart. 2010. *Quantity Implicatures*. Cambridge: Cambridge University Press.
- Goffman, Erving. 1967. *Interaction Ritual: Essays in Face-to-Face Behavior*. Chicago: Aldine.
- Gotzner, Nicole, Stephanie Solt & Anton Benz. 2018. Scalar diversity, negative strengthening, and adjectival semantics. *Frontiers in Psychology* 9: 16–59.
- Grice, H. Paul. 1989. *Studies in the Way of Words*. Cambridge, MA: MIT Press.
- Haugh, Michael. 2003. Anticipated versus inferred politeness. *Multilingua* 22: 397–413.
- Hernandez-Flores, Nieves. 1999. Politeness ideology in Spanish colloquial conversation: The case of advice. *Pragmatics* 9: 37–49.
- Hirschberg, Julia. 1991. *A Theory of Scalar Implicature*. New York: Garland Press.
- Holtgraves, Thomas. 2014. Interpreting uncertainty terms. *Journal of Personality and Social Psychology* 107: 219–228.
- Holtgraves, Thomas & Jean Bonnefon. 2017. Experimental approaches to linguistic (im)politeness. In J. Culpeper, M. Haugh & D. Kádár (eds.), *The Palgrave Handbook of Linguistic (Im)politeness*, 381–401. London: Palgrave.
- Holtgraves, Thomas & Brian Kraus. 2018. Processing scalar implicatures in conversational contexts: An ERP study. *Journal of Neurolinguistics* 46: 93–108.
- Holtgraves, Thomas & Audrey Perdeu. 2016. Politeness and the communication of uncertainty. *Cognition* 154: 1–10.
- Horn, Laurence R. 1972. *On the Semantic Properties of Logical Operators in English*. Ph.D. thesis, UCLA.
- Horn, Laurence R. 1984. A new taxonomy for pragmatic inference: Q-based and R-based implicature. In D. Schiffrin (ed.), *Meaning, Form and Use in Context (GURT '84)*, 11–42. Washington: Georgetown University Press.
- Horn, Laurence R. 1992. The said and the unsaid. In C. Barker & D. Dowty (eds.), *Proceedings of SALT II*, 163–192. Department of Linguistics, Ohio State University.
- Hudson, Richard. 1996. *Sociolinguistics*. Cambridge: Cambridge University Press.
- Jiang, Xiaoming & Xiaolin Zhou. 2015. Impoliteness electrified: ERPs reveal the real-time processing of disrespectful reference in Mandarin utterance comprehension. In M. Terkourafi (ed.), *Interdisciplinary Perspectives on Im/politeness*, 239–266. Amsterdam: John Benjamins.
- Kerbrat-Orecchioni, Catherine. 1997. A multilevel approach in the study of talk in interaction. *Pragmatics* 7: 1–20.
- Kimball, Amelia, Kailen Shantz, Christopher Eager & Joseph Roy. 2018. Confronting

- quasi-separation in logistic mixed effects for linguistic data: A Bayesian approach. *Journal of Quantitative Linguistics* 1–25.
- Leffel, Timothy, Alexandre Cremers, Jacopo Romoli & Nicole Gotzner. 2019. Vagueness in implicature: The case of modified adjectives. *Journal of Semantics* (in press).
- Length, Russell V. 2016. Least-squares means: The R package lsmeans. *Journal of Statistical Software* 69: 1–33.
- Levinson, Stephen C. 1995. Three levels of meaning. In F. Palmer (ed.), *Grammar and Meaning: Essays in honour of Sir John Lyons*, 90–115. Cambridge: Cambridge University Press.
- Levinson, Stephen C. 2000. *Presumptive Meanings: The Theory of Generalized Conversational Implicature*. Cambridge, MA: MIT Press.
- Liddell, Torrin & John K. Kruschke. 2018. Analyzing ordinal data with metric models: What could possibly go wrong? *Journal of Experimental Social Psychology* 79: 328–348.
- McNally, Louise. 2017. Scalar alternatives and scalar inference involving adjectives: A comment on van Tiel, et al. 2016. In J. Ostrove, R. Kramer & J. Sabbagh (eds.), *Asking the Right Questions: Essays in Honor of Sandra Chung*, 17–28.
- Mazzarella, Diana, Emmanuel Trouche, Hugo Mercier & Ira Noveck. 2018. Believing what you're told: Politeness and scalar inferences. *Frontiers in Psychology* 9: 12–33.
- Morgan, Marcyliena. 1991. Indirectness and interpretation in African American women's discourse. *Pragmatics* 1: 421–451.
- Nicenboim, Bruno & Shravan Vasishth. 2016. Statistical methods for linguistic research: Foundational Ideas—Part 11. *Language and Linguistics Compass* 10: 591–613.
- Noveck, Ira & Anne Reboul. 2008. Experimental pragmatics: A Gricean turn in the study of language. *Trends in Cognitive Science* 12, 425–431.
- O'Driscoll, Jim. 1996. About face: A defence and elaboration of universal dualism. *Journal of Pragmatics* 25: 1–32.
- O'Driscoll, Jim. 2017. Face and (im)politeness. In J. Culpeper, M. Haugh & D.K. Kádár (eds.), *The Palgrave Handbook of Linguistic (Im)politeness*, 89–118. London: Palgrave.
- Pouscoulous, Nausicaa, Ira Noveck, Guy Politzer & Anne Bastide. 2007. A developmental investigation of processing costs in implicature production. *Language Acquisition* 14: 347–375.
- Ruytenbeek, Nicolas, Steven Verheyen & Benjamin Spector. 2017. Asymmetric inference towards the antonym: Experiments into the polarity and morphology of negated adjectives. *Glossa: A Journal of General Linguistics* 2: 1–27.
- Scollon, Ron & Susanne W. Scollon 1995. *Intercultural Communication: A Discourse Approach*. Oxford: Blackwell.
- Simons, Mandy & Tessa Warren. 2018. A closer look at strengthened readings of scalars. *Quarterly Journal of Experimental Psychology* 71: 272–279.

- Sperber, Dan, Fabrice Clément, Christophe Heintz, Olivier Mascaro, Hugo Mercier, Gloria Origgi & Deirdre Wilson. 2010. Epistemic vigilance. *Mind & Language* 25: 359–393.
- Sperber, Dan & Deirdre Wilson. 1986. *Relevance: Communication and Cognition*. Oxford: Blackwell.
- Sun, Chao, Ye Tian & Richard Breheny. 2018. A link between local enrichment and scalar diversity. *Frontiers in Psychology* 9: 20–92.
- Terkourafi, Marina. 2001. *Politeness in Cypriot Greek: A frame-based approach*. Unpublished PhD Thesis. University of Cambridge, Department of Linguistics.
- Terkourafi, Marina. 2003. Generalised and particularised implicatures of linguistic politeness. In P. Kühnlein, H. Rieser & H. Zeevat (eds.), *Perspectives on Dialogue in the New Millennium*, 149–164. Amsterdam: John Benjamins.
- Terkourafi, Marina. 2007. Toward a universal notion of face for a universal notion of co-operation. In I. Kecskes & L. Horn (eds.), *Explorations in Pragmatics: Linguistic, Cognitive and Intercultural Aspects*, 313–344. Berlin: Mouton de Gruyter.
- Terkourafi, Marina. 2008. Toward a unified theory of politeness, impoliteness, and rudeness. In D. Bousfield & M. Locher (eds.), *Impoliteness in Language: Studies on its Interplay with Power in Theory and Practice*, 45–74. Berlin: Mouton de Gruyter.
- Turner, Ken. 1996. The Principal principles of pragmatic inference: Politeness. *Language Teaching* 29: 1–13.
- van Rooij, Robert & Katrin Schulz. 2004. Exhaustive interpretation of complex sentences. *Journal of Logic, Language and Information* 13: 491–519.
- van Tiel, Bon, Emiel van Miltenburg, Natalia Zevakhina & Bart Geurts. 2016. Scalar diversity. *Journal of Semantics* 33: 137–175.
- Vehtari, Aki, Andrew Gelman & Jonah Gabry. 2017. Practical Bayesian model evaluation using leave-one-out cross-validation and WAIC. *Statistics and Computing* 27: 1413–1432.
- Watts, Richard. 2002. From polite language to educated language. In R. Watts & P. Trudgill (eds.), *Alternative Histories of English*, 155–172. London: Routledge.

Appendix 1—Main study stimuli

⟨some, all⟩—Subject 1

Face boosting: John is a contestant in a local poetry competition. On the big night, he gets up on stage and reads his poem. When he returns to his seat, his girlfriend Mary smiles and tells him: “Some people loved your poem.”

Face threatening: John has submitted a poem to a venerated poetry club hoping to become a member. This is John’s third attempt to join the club, and he is in a fierce competition with George for the chance to join. While they are waiting to hear the results, George goes up to John and tells him: “Some people loved your poem.”

⟨some, all⟩—Subject 2

Face boosting: Nicole is at gymnastics practice and performs her new routine in front of her parents for the first time. After the routine, she goes over to her parents and her mom tells her, “Some parts were perfect.”

Face threatening: Nicole is at gymnastics practice and performs her new routine in its entirety for the first time in front of her tough and demanding coach. After the routine, she goes over to her coach, who tells her, “Some parts were perfect.”

⟨some, all⟩—Subject 3

Face boosting: Margaret and William are dropping their daughter off at her first dance practice. They are greeted warmly by the instructor, and are very impressed by how welcoming she is. At the end of their short conversation with her, the instructor tells them: “Some parents stay to watch the rehearsal.”

Face threatening: Margaret and William are dropping their daughter off at her first dance practice. As practice starts, they are hovering around the side of the room, watching the rehearsal. The instructor feels uncomfortable and their presence is bothering her, so she comes over to talk to them. She says: “Some parents stay to watch the rehearsal.”

⟨some, all⟩—Subject 4

Face boosting: Patrick, has spent all day in the kitchen preparing a Thanksgiving feast with many dishes. After the feast, his loving wife who has always encouraged him to get more involved in the kitchen tells him: “Some dishes were wonderful.”

Face threatening: Patrick has spent all day in the kitchen preparing a Thanksgiving feast with many dishes. After the feast, his tough-to-please mother-in-law remarks to him: “Some dishes were wonderful.”

⟨*some*, *all*⟩—Object 1

Face boosting: Barbara's two sons have been entertaining themselves all day, playing safely and happily outside while allowing Barbara to complete her work. She wants to reward them for being good boys, so she puts out a plate of cookies and calls them inside. She says to them, "You can have some cookies."

Face threatening: Barbara's two sons have been a nuisance all day, continuously bothering her and asking her for cookies while she is trying to complete her work. Frustrated, she finally gives in and opens the box, saying to them, "You can have some cookies."

⟨*some*, *all*⟩—Object 2

Face boosting: Kevin and his professor have been doing some exciting research and are now looking at a list of conferences where Kevin can travel and present their work. His professor tells him, "Looks like you'll go to some conferences."

Face threatening: Kevin has been begging his professor to let him travel to conferences to present his work, but his professor is a bit hesitant. Kevin shows his professor a list of conferences, to which his professor responds, "Looks like you'll go to some conferences."

⟨*some*, *all*⟩—Object 3

Face boosting: Oliver and his mother are walking down the street when they see a store with a sign in the window that says "Free Books!" They head in and are greeted by the store owner, who shows Oliver a row of children's books. Oliver shyly proceeds to the row and looks through the books with his mother. The store owner walks over to them and says, "Feel free to take some books."

Face threatening: Ron is an avid reader of second-hand books. One day he walks past a new second-hand bookstore and sees a sign in the window that says "Free Books!" He reaches into his backpack and pulls out a large grocery bag, then enters the store and starts loading books into it. The store owner goes up to him and says, "Feel free to take some books."

⟨*some*, *all*⟩—Object 4

Face boosting: Jennifer, a Democrat, is chatting with her Democrat friends. She says: "I disagree with some Republicans."

Face threatening: Jennifer, a Democrat, is arguing with her Republican friends. She says: "I disagree with some Republicans."

<or, and> 1

Face boosting: Nick and Amanda went on a date last week and both had a great time. Nick eagerly asks Amanda when they can hang out again and she responds, “Wednesday or Thursday.”

Face threatening: Nick has been pressuring Amanda to go on a date with him for weeks. She decides, begrudgingly, to give him a chance. Nick eagerly asks her when they can go out and she much-less-eagerly responds, “Wednesday or Thursday.”

<or, and> 2

Face boosting: Susan brings her 8-year-old son, Ben, to the bank because the nice teller always offers Ben candy. The teller sees Ben eagerly approaching the counter and says to him, “You can have Skittles or Starbursts.”

Face threatening: Susan brings her annoying 8-year-old son, Ben, to the bank because the bank tellers usually have candy to hand out. The cranky old teller, annoyed at Ben’s over-enthusiasm to get candy, tells him “You can have Skittles or Starbursts.”

<or, and> 3

Face boosting: Mark is helping his friend Greg move a bunch of boxes in Greg’s garage. They head into the house for a break, and Greg offers Mark a cold drink. Mark asks Greg what he has, and Greg responds, “There’s water or beer.”

Face threatening: Mark is helping his friend Greg move a bunch of boxes in Greg’s garage. Mark keeps complaining about being hot and tired, and Greg is frustrated and just wants to hurry the process. Greg eventually tries to shut Mark up and says: “There’s water or beer.”

<or, and> 4

Face boosting: Quincy, the beloved superstar athlete on the track team, is chatting with his coach after practice and asks him what he’ll be participating in in the upcoming meet. His coach responds, “You’ll be doing high jump or hurdles.”

Face threatening: Quincy, one of the worst athletes on the track team, is annoying his coach by asking what he’ll be participating in in the upcoming meet. His coach responds, “You’ll be doing high jump or hurdles.”

<often, always> 1

Face boosting: Thomas is thanking his therapist at the end of a session. Thomas has been feeling much better recently, and thinks the therapy is going really well. He tells his therapist: “Your suggestions are often helpful.”

Face threatening: Thomas is going to fire his therapist at the end of this session.

Thomas doesn't get along with the therapist on a personal level, which he feels impacts his ability to make progress. He tells his therapist: "Your suggestions are often helpful."

⟨often, always⟩ 2

Face boosting: Derek and Peter have been best friends since kindergarten. Derek has been hoping to ask Courtney out on a date for a while, but has noticed recently that she spends a lot of time with Chris. Derek asks Peter, who he always goes to for support, if he knows anything about Courtney's relationship status. Peter responds: "I don't know, but they're often together."

Face threatening: Derek has been hoping to ask Courtney out on a date for a while, but has noticed recently that she spends a lot of time with Chris. Derek asks Courtney's best friend, who knows Courtney doesn't like Derek, about Courtney's relationship status with Chris. Her friend responds: "I don't know, but they're often together."

⟨often, always⟩ 3

Face boosting: Simon just moved to a new city for work. He doesn't have any friends in the area, but his co-workers have been very welcoming and nice to him. Simon asks one of these co-workers, Jacob, what he's doing for Thanksgiving, which is coming up soon. Jacob responds, "I often host a Thanksgiving dinner for my friends."

Face threatening: Simon just moved to a new city for work. He doesn't have any friends in the area, but he tries to be friendly to his new co-workers. Unfortunately, his new co-workers, especially Jacob, find him annoying. Simon asks Jacob what he's doing for Thanksgiving, which is coming up soon. Jacob responds, "I often host a Thanksgiving dinner for my friends."

⟨often, always⟩ 4

Face boosting: Joanne is the executive chef in a prestigious restaurant. She just hired Trent to work in the kitchen. Paige has been working there for a year, so she is tasked with training Trent. She is very friendly and helpful while training him, and the working environment is very comfortable. Paige notes, "Joanne is often pleased when the plate presentation is perfect."

Face threatening: Joanne is the executive chef in a prestigious restaurant. She just hired Trent to work in the kitchen. Paige has been working there for a year, so she is tasked with training Trent. She is clearly stressed out from working under the strict guidance of Joanne, and is not very friendly or helpful while training Trent. Paige notes, "Joanne is often pleased when the plate presentation is perfect."

⟨possible, likely⟩ 1

Face boosting: Michael is talking to his girlfriend, who lives in Chicago, about where he may wind up getting placed for his new job. He's going over the options and enthusiastically concludes: "It's possible that I'll end up in Chicago."

Face threatening: Michael is talking to his girlfriend, who lives in Los Angeles, about where he may wind up getting placed for his new job. She keeps asking him to go over the options and he concludes, somewhat frustrated, "It's possible that I'll end up in Chicago."

⟨possible, likely⟩ 2

Face boosting: Riley is hoping to buy a new car and has been saving up her money for a while. Her parents are rich, but she usually doesn't like to rely on them for financial help. When she decides to ask for their support, her mother responds: "It's possible we could help you out."

Face threatening: Riley is hoping to buy a new car but hasn't been responsible with her money. She's asked her parents for several loans, but never seems to save up any money towards the car. When she goes to ask them for a loan for the 5th time, her mother responds: "It's possible we could help you out."

⟨possible, likely⟩ 3

Face boosting: Victor's car is in the shop for the afternoon, so he asks his roommate, Reggie, who just bought a new car and is excited to drive it at all times, if he could get a ride home from work later. Reggie responds, "It's possible that'll work, it depends on the timing."

Face threatening: Victor sold his car and has been repeatedly asking his roommates if they'd pick him up at work and give him a ride home. Today, he asks his roommate Reggie, who's getting frustrated with the situation, for a ride back in the afternoon. Reggie responds, "It's possible that'll work, it depends on the timing."

⟨possible, likely⟩ 4

Face boosting: Alex and his best friend Marco are out at a bar, and Alex spends a lot of the night chatting with a girl. At the end of the night, they exchange numbers, and Alex is telling Marco about it on the way home. Alex asks Marco if he thinks she'll ever text him, to which Marco responds, "It's possible."

Face threatening: Alex is over at Marco's house for a barbecue. Marco can't help but be bothered that Alex spends most of the time hitting on Marco's sister. Alex gives her his number at the end of the night, and later asks Marco if he thinks she'll ever text him. Marco responds, "It's possible."

⟨like, love⟩ 1

Face boosting: Henrietta and Brandon are brainstorming where to go for dinner. Henrietta suggests her favourite Indian food restaurant, and Brandon enthusiastically responds, “Sure, I like Indian food.”

Face threatening: Henrietta and Brandon are arguing yet again about where to go for dinner. Henrietta suggests her favourite Indian food restaurant, and Brandon, tired of all the arguing, responds, “Sure, I like Indian food.”

⟨like, love⟩ 2

Face boosting: Natalie is getting ready for a night out. She shows the clothes she's thinking about wearing to Sofia, her friendly and supporting roommate, who tells Natalie, “I like your outfit.”

Face threatening: Natalie is getting ready for a night out. She shows the clothes she's thinking about wearing to Sofia, her harsh and critical roommate, who tells Natalie, “I like your outfit.”

⟨like, love⟩ 3

Face boosting: Mara is showing some of her paintings at an art exhibition in hopes of selling them. She notices her favourite uncle standing in front of one particular painting for a while, so she walks up to him and asks him what he thinks. He responds, “I like this painting.”

Face threatening: Mara is showing some of her paintings at an art exhibition in hopes of selling them. She notices a famous art critic in attendance, walking around and shaking his head in disapproval. While he's paused in front of one particular painting, she walks over to him and asks him what he thinks. He responds, “I like this painting.”

⟨like, love⟩ 4

Face boosting: Emily wants to ask Brett out on a date. Brett's best friend, Adam, knows that Brett would happily say yes if Emily asked him. Emily consults Adam about the situation, and asks if Brett would like going to the movies. Adam tells her: “Yeah, Brett likes going to the movies.”

Face threatening: Emily wants to ask Brett out on a date. Brett's best friend, Adam, knows that Brett would be horrified at the proposition and would not want to go out with Emily. Emily consults Adam, and asks if Brett would like going to the movies. Adam tells her: “Yeah, Brett likes going to the movies.”

⟨good, excellent⟩ 1

Face boosting: David has been hoping for and expecting a promotion at work recently. David's boss calls him into the office, smiles, and says, "You are a good employee."

Face threatening: David's company has been going through some rough times financially and thus has had to make several layoffs around the office. David is one of the least productive employees at the company, and doesn't have a good relationship with the boss. David's boss calls him into the office and says, "You are a good employee."

⟨good, excellent⟩ 2

Face boosting: Lucy and her boss, Karen, are in a meeting to discuss Lucy's recent idea to try and save the company money. Karen, who likes Lucy and has been talking to the CEO about promoting her, responds to Lucy's suggestion: "You're presenting a good idea."

Face threatening: Lucy and her boss, Karen, are in a meeting to discuss Lucy's recent idea to try and save the company money. Karen, who was very close to firing Lucy just last month and still doesn't trust her to handle the responsibilities of the job, responds to Lucy's suggestion: "You're presenting a good idea."

⟨good, excellent⟩ 3

Face boosting: Robin goes to meet with Professor Wilson, her favorite professor, in hopes of getting her to write a letter of recommendation. They get along extremely well, and Professor Wilson has told Robin that she'd love to write her a letter. Professor Wilson responds to Robin's request: "I'll be sure to write you a good recommendation."

Face threatening: Robin goes to meet with Professor Wilson in hopes of getting her to write a letter of recommendation. They don't get along very well at all, but Robin is out of options and desperately needs a letter from someone. Professor Wilson responds to Robin's request: "I'll be sure to write you a good recommendation."

⟨good, excellent⟩ 4

Face boosting: Paul has his first guitar lesson with his new teacher. Paul plays a portion of a song so the teacher can get a sense of his abilities. The teacher, who is eager for new students, tells Paul, "You have a good sense of rhythm."

Face threatening: Paul is playing guitar in a competition with a notoriously strict panel of judges. After Paul plays his song, the first judge is silent for a while and then mutters, "You have a good sense of rhythm."

Non-entailment scales 1 *(unwell, sick)*

Face boosting: Julian knows that his friends usually go for drinks after class and has been hoping they will ask him to join them. One day after class they finally extend the invite to him. Julian responds: "Sorry I can't, I've been feeling unwell."

Face threatening: Julian's classmates want him to come out to a bar with them after class, but Julian doesn't like them and doesn't want to go with them. When class is over, they extend the invite. Julian responds: "Sorry I can't, I've been feeling unwell."

Non-entailment scales 2 *(misguided, illegal)*

Face boosting: Cameron and Victoria are con artists, partners in crime. Victoria is describing to Cameron her newest plot, which will involve swindling a top government official out of a lot of money. Cameron responds: "Did you consider your plan may be misguided?"

Face threatening: Cameron and Victoria are con artists who have been caught by the police for attempting to swindle a top government official out of a lot of money. At the opening of their trial, the judge asks them: "Did you consider your plan may be misguided?"

Non-entailment scales 3 *(assertive, bossy)*

Face-boosting: Sarah is telling her friend Gina about an incident with a guy who kept hitting on Sarah at a party and wouldn't take no for an answer. Gina is consoling Sarah and responds: "In cases like this, it's OK to be assertive."

Face-threatening: Sarah and Gina have been working on a class project. Sarah doesn't feel like it's been going very well, and Gina's personality rubs her the wrong way. She tells Gina this, and Gina responds: "In cases like this, it's OK to be assertive."

Non-entailment scales 4 *(misleading, lying)*

Face boosting: Natasha is a community organizer in her town and is often out on the streets with her friend Nina gathering signatures for one cause or another. Recently, she has been frustrated at how the local city officials treat the public and is thinking of starting a petition to hold them accountable for their behaviour. She tells Nina about it and she says "I can see why you think they have been misleading you."

Face threatening: Natasha is a community organizer in her town and is often out on the streets gathering signatures for one cause or another. Some people, like her neighbor, Louis, consider her a bit of a trouble-maker. Recently, she has been frustrated at how the local city officials treat the public and is thinking of starting a petition to hold them accountable for their behaviour. She tells Louis about it and he says "I can see why you think they have been misleading you."