

How can they both be right?: Faultless disagreement and semantic adaptation

Emily Pecsok (egp45@cornell.edu)

Department of Linguistics
Cornell University

Helena Aparicio (haparicio@cornell.edu)

Department of Linguistics
Cornell University

Abstract

Disagreements are speech acts used by interlocutors to challenge previous assertions. When disagreements express subjective views, they can often be perceived as faultless. However, it is unclear whether accepting a disagreement as faultless causes comprehenders to update their own semantic representations of the predicate targeted by the disagreement. Using the vague quantifiers *many* and *few* as a case study, we find in two adaptation studies that participants shifted their meaning representations of the quantifiers after being exposed to disagreements that on average were more likely to be perceived as faultless. The adaptation strengthened the participants' baseline preferences, suggesting that even when a disagreement is judged to be faultless, there exists a perceived asymmetry in the plausibility of the two viewpoints under discussion.

Keywords: vague quantifiers; faultless disagreement; adaptation; subjectivity; experimental semantics and pragmatics

Introduction

Disagreements constitute a common conversational move used by speakers to negotiate what information should be added to the common ground (Stalnaker, 2002). Speakers often disagree on matters that are not objective in nature. Such is the case of disagreements about aesthetic judgments, subjective beliefs, moral values, or vague meanings. It has been observed that disagreements about non-objective matters can be potentially perceived as *faultless*, in the sense that neither speaker is taken to be at fault (Kölbel, 2004; Lasersohn, 2009; Sundell, 2011; Foushee & Srinivasan, 2017). This is exemplified in (1), which contains a disagreement dialogue involving the vague quantifier *many*.

- (1) a. S1: Many people voted by mail this election cycle.
b. S2: No, I don't think many people voted by mail this election cycle.

Even if the speakers in (1) have completely accurate knowledge of the situation, i.e., they know the precise number of mail-in voters, and the total number of possible voters, it is still possible for both speakers to be judged as correct. This is because the speakers in (1) are not necessarily disagreeing about the descriptive content of S1's utterance, but rather about the assertability of the utterance itself (Horn, 1989). Put differently, S2 is not taking issue with the propositional content of S1's assertion, but rather with whether the particular amount of absentee voters should be described as *many*. In this respect, the disagreement in (1) is of a metalinguistic

nature (Barker, 2002): the speakers are disagreeing on how to parametrize the meaning of the vague quantifier *many* in the given context.

The fact that the disagreement in (1) can be perceived as faultless is a direct consequence of the vagueness of the quantifier. A common feature of vague predicates is that their semantic representation underdetermines interpretation. Partee (2004) proposes that the meaning of a sentence like (2-a), containing the vague quantifier *many*, can be modeled as follows.¹

- (2) a. Many [voters in this election cycle]_A [voted by mail]_B
b.

$$\frac{|A \cap B|}{|A|} \geq \theta$$

The equation in (2-b) states that *many* is a relation between two sets of individuals *A* (i.e., the set of *the set of voters in this election cycle*) and *B* (i.e., the set of *mail-in voters*) and a threshold variable θ providing the lower-bound proportion for which the quantifier is still applicable. The meaning in (2-b) returns *true* as long as the proportion resulting from dividing the cardinality corresponding to the intersection of *A* and *B* over the cardinality of *A* is equal to or greater than the threshold θ . Otherwise, the meaning returns *false*. The meaning in (2-b) however, is underspecified with respect to the value of θ , which must be set via pragmatic reasoning (Pepper & Prytulak, 1974; Fernando & Kamp, 1996; Kennedy, 2007; Ramotowska, Haaf, van Maanen, & Szymanik, 2022). For instance, the lower-bound amount of voters that can be felicitously described as *many* in (1) is going to depend on the features of the context (e.g., is the election local or national?), as well as the personal preferences of the speakers (i.e., even after having controlled for the relevant context features, different speakers will have different views on the range of values that constitute a sensible lower-bound amount of voters). Because of the high speaker variability displayed by vague expressions, it has been argued that utterances like (1-a) encode a subjective dimension. For instance, the assertion in (1-a) expresses a viewpoint of what the operative threshold

¹(2-b) provides the semantic definition for the so called *proportional* reading of the quantifier *many*, where the threshold θ is taken to denote a proportion. Vague quantifiers can also have *cardinal readings* where the threshold variable denoting the lower bound is a cardinal number. In this paper we focus on proportional readings.

for the vague quantifier should be in the context (Kennedy, 2013; Fleisher, 2013; Barker, 2013). This has led some authors to argue that subjective assertions commit the speaker to the truth of the proposition as judged by the speaker herself (Lasersohn, 2005; Rudin & Beltrama, 2019). In other words, the truth of a proposition containing a vague predicate is relativized to a judge parameter that, in most cases, is the speaker.

A parallel body of work on semantic and pragmatic adaptation has shown that listeners are able to track interspeaker variability displayed by vague expressions and that listeners use information about the input statistics associated with a particular speaker to guide utterance interpretation (Yildirim, Degen, Tanenhaus, & Jaeger, 2016; Schuster & Degen, 2019, 2020). More relevant to the present work are recent findings suggesting that comprehenders adapt their meaning representations to reflect the observed input statistics (Heim, Peiseler, & Bekemeier, 2020; Xiang, Kramer, & Kennedy, 2020). For instance, Heim et al. find that comprehenders shift their criteria of applicability of the quantifier *many* after being exposed to extreme uses of the predicate (i.e., uses where *many* describes low percentages, i.e., 20-50%). In particular, their results showed that participants adapted their meaning representations to accommodate the extreme uses observed during the exposure trials.

Here we investigate whether participants update their semantic representations after exposure to metalinguistic disagreements involving the vague quantifiers *many* and *few*. Previous experimental work on the topic of faultless disagreement has focused on the conditions that modulate the acceptability of a disagreement as faultless (Scontras, Degen, & Goodman, 2017; Kaiser & Rudin, 2020, 2021), but little is known about whether disagreements that have the potential of being faultless further modulate the highly malleable meanings that such disagreements target. This question is important because it can help illuminate the strategies deployed by listeners to cope with faultless contradictions. In two adaptation studies, we test the following three hypotheses about how participants might update their representations of vague quantifiers after exposure to metalinguistic disagreement dialogues: 1) we first consider the null hypothesis that metalinguistic disagreements will not give rise to adaptation effects. Such results would be expected if listeners interpret the truth of vague meanings as relativized to the speaker. As such, comprehenders might not be compelled to reconcile the conflicting viewpoints by adapting their own semantic representations; 2) Hypothesis 2 states that comprehenders will try to reconcile the disagreeing views by adopting a middle-ground semantic representation. This hypothesis predicts that comprehenders should adapt towards a representation that is more equidistant to the views expressed by the disagreeing speakers; 3) Hypothesis 3 states that when faced with a metalinguistic disagreement, comprehenders will side with the viewpoint that better represents their initial baseline preferences. This hypothesis predicts selective adaptation effects

such that the updated meaning representation should be more categorical than at baseline.

Our results provide evidence for Hypothesis 3: after exposure to metalinguistic disagreements, participants doubled down on their initial quantifier acceptance views, but only in conditions that had relatively high rates of faultless disagreement. Furthermore, no adaptation effects were detected when the same information was presented as single statements instead of disagreement dialogue. Our results suggest that inconsistent views only alter semantic representations when part of disagreement dialogues that are commonly deemed faultless. Interestingly, the adaptation does not occur in a direction that reconciles the conflicting viewpoints. We discuss the implications of these findings for the study of metalinguistic disagreement.

Experiment 1

The goal of Experiment 1 is to investigate if exposure to faultless disagreements targeting the quantifiers *few* and *many* leads to updated semantic representations of these quantifiers.

Materials and Procedure

Experiment 1 was modeled after Heim et al. (2020). The experiment contained three blocks. In Block 1, participants were shown a series of images containing 50 circles of different sizes. Each image had blue and/or yellow circles on a gray background. The images ranged from having 20-80% of either color. Each image was paired with a question of the form ‘*Do you think {many, few} of the circles are {yellow, blue}*’ (see Figure 1). Participants were instructed to indicate whether or not they agreed with the statement by pressing F for *yes* or J for *no* on their keyboards. Each participant saw all four quantifier-color pairings twice for each percentage point, resulting in 56 critical trials. Throughout the trials, participants were also given eight attention checks, where they were asked *Do you think {none, all} of the circles are {blue, yellow}*?. These questions were paired with images where all of the circles were only one color.

In Block 2, participants were shown additional images of the circles where the target color made up the same percentages as Block 1 (i.e. 20%, 30%, 40%, 50%, 60%, 70% and 80%). Alongside each image, would be an audio recording of the following dialogue:

- (3) First Speaker: {Many, few} of the circles are {blue, yellow}.
Second Speaker: I disagree,² I don’t think {many, few} of the circles are {blue, yellow}.

There were four speakers in total that all participants heard. Only two speakers, one male and one female, were heard in each trial and the same speakers were paired for every trial.

²The beginning of the denial varied slightly to prevent the dialogues from sounding too repetitive. Beginnings of disagreements included: No; I disagree with you; I don’t agree; and I don’t agree with you.

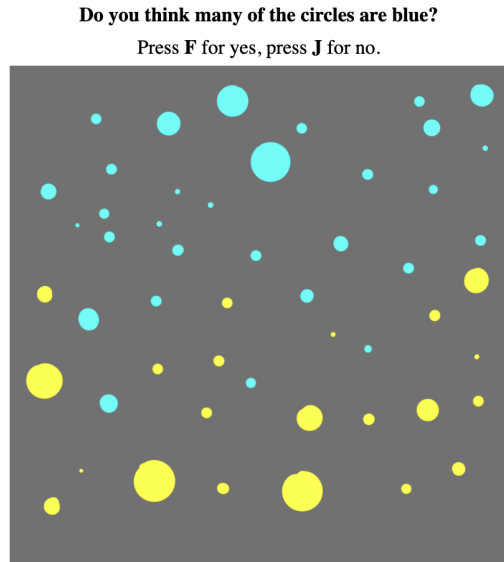


Figure 1: Example trial pertaining to Block 1 and Block 3 (Experiment 1).

Speaker views remained consistent through all percentages, that is, two speakers always used the quantifier *many* and were paired with the two speakers who always used the quantifier *few*. After hearing the dialogue, participants were instructed to select among three options: *Only the first speaker is right*, *Only the second speaker is right* and *Both speakers can be right*. These three options were presented as buttons below the image and participants were told to select the button that best described the dialogue (See Figure 2 for an example). Participants were able to see the prompt and response options while the recordings were playing, but participants were not allowed to respond until the recordings of both speakers were finished. Participants completed four trials for every quantifier-percentage pair, resulting in a total of 56 trials. Before Block 2, participants had an audio check and two practice trials, where two novel speakers disagreed about objective facts regarding whether or not all of the triangles in an image were above all of the squares. Block 3 was an exact repeat of Block 1.

Participants

A total of 60 participants were recruited via Prolific. Participants were born and currently located in the United States, and were self-described native English speakers between the ages of 18 and 35. Participants were paid \$3.00, which amounted to an hourly wage of approximately \$12.00. We excluded participants who failed more than four attention checks and/or failed the audio check presented in Block 2. A total of 4 participants were excluded based on the above criteria. Finally, one participant was excluded due to errors in data collection, leaving a total of 55 participants.

Select the button that, in your opinion, best describes the dialogue.

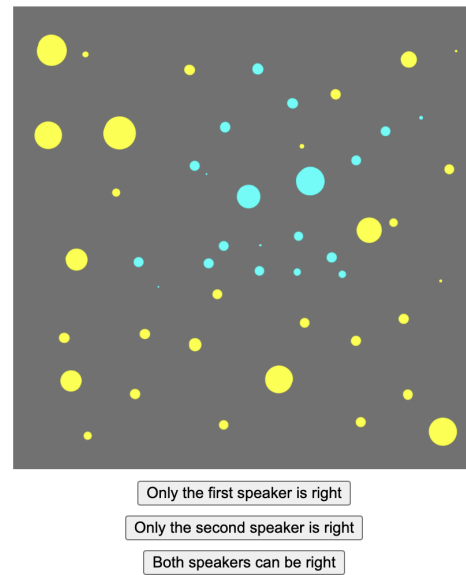


Figure 2: Example trial pertaining to Block 2 (Experiment 1).

Results

Figure 3 contains the proportions of faultless disagreement acceptance (i.e., selection of *Both speakers can be right* responses) in Block 2. As expected, participants displayed more categorical behavior when the target color made up 20% or 80% of all the circles, where there was a strong preference for responses indicating that only one speaker could be right for both quantifiers. Since we expect to find potential adaptation effects at percentages with high acceptance of faultless disagreement, we used the two extreme percentages—20% and 80%—to determine which percentages had significantly different acceptance rates of faultless disagreement. We fit a logistic mixed effects regression model comparing the faultless disagreement acceptance rates at 30%, 40%, and 50% to the baseline 20%. All three percentages presented significantly higher rates of faultless disagreement compared to 20% (30%: $\beta = 1.93$; 40%: $\beta = 4.63$; 50%: $\beta = 4.17$, all p 's < 0.05). The same model was used to compare the faultless disagreement acceptance at 50%, 60%, and 70% to the baseline 80%. Results showed that the rates of acceptance for 50% and 60% were significantly different (50%: $\beta = 4.2$; 60%: $\beta = 3.2$, , p 's < 0.05) but acceptance at 70% was not for (70%: $\beta = 0.5$, $p > 0.05$). The percentages that demonstrated statistically significant differences—30%, 40%, 50%, and 60%—were further analyzed for adaptation effects and will be discussed below.

Figure 4 contains the results comparing quantifier applicability judgments in Block 1 and Block 3. To determine whether the quantifier acceptability patterns shifted after exposure to disagreements, we fit a logistic mixed effects regression model to the data for each of the four target color percentages where participants perceived the disagreements

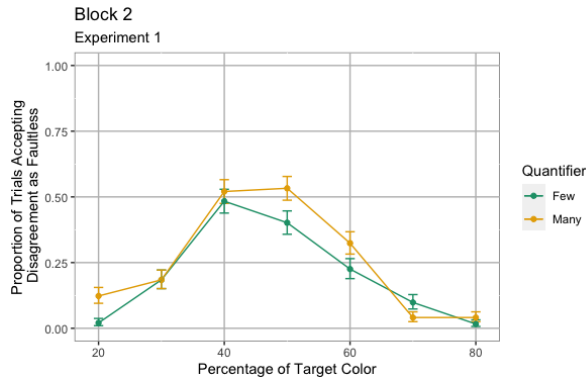


Figure 3: Proportion of acceptance of faultless disagreement (i.e., selection of *Both speakers can be right* response) at the different percentages tested. Error bars correspond to 95% confidence intervals.

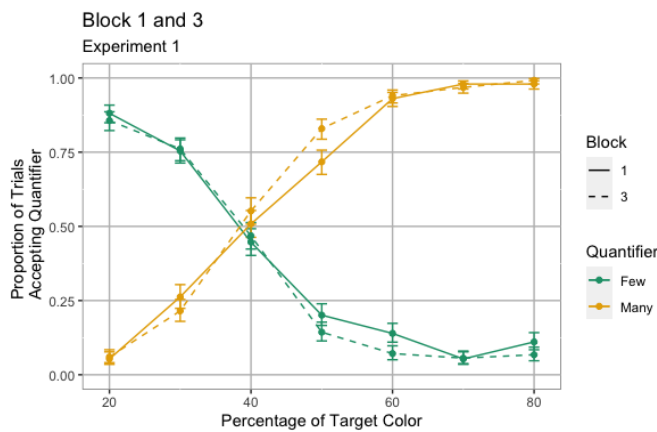


Figure 4: Proportion of acceptance of *many* and *few* at the different percentages tested in Block 1 and Block 3. Error bars correspond to 95% confidence intervals.

to be faultless significantly more than baseline (percentages 30%-60%). The model predicted the response data from the fixed effects of BLOCK (Block 1 vs. Block 3), QUANTIFIER (*few* vs. *many*) and their interaction. The model also included random intercepts by-participant and by-condition random slopes. Predictors were contrast coded. Results revealed that the interaction was significant at 50% ($\beta = 1.45$, $z = 4.46$, $p < 0.001$) and 60% ($\beta = 5.81$, $z = 5.65$, $p < 0.001$). Simple effect analyses revealed that at 50% the interaction was uniquely driven by differences in *many*, with participants providing significantly higher ratings in Block 3 compared to Block 1 ($\beta = 2.19$, $z = 2.28$, $p < 0.05$). No differences were detected for *few*; conversely, at 60% the interaction was mainly driven by *few*, which showed significantly lower ratings at Block 3 compared to Block 1 ($\beta = -3.08$, $z = -2.1$, $p < 0.05$). No differences were detected for *many* at this percentage. Finally, no significant differences were detected at 30% and 40% ($p > 0.05$ for both percentages).

Discussion

The current results show evidence that participants indeed adapted their representation of the quantifiers *few* and *many* in some of the percentages where disagreement dialogues in Block 2 had been perceived to be faultless at higher rates than baseline, i.e. percentages 50% and 60%. Conversely, no differences were found at the baseline percentages 20%, 70%, and 80%, for which the rates of faultless disagreement in Block 2 were the lowest. The detected adaptation effects showed participants displaying less uncertainty about the applicability of the quantifier in Block 3 compared to Block 1. Specifically, participants showed higher acceptability rates for *many* at 50% and lower acceptability rates for *few* at 60%. This entails that participants had less meaning uncertainty as the percentages got further away from 40%, the point of maximal uncertainty about whether the predicate was applicable or not. The current results therefore rule out the null hypothesis (Hypothesis 1) that no adaptation effects would occur as a result of exposure to potentially faultless disagreement. Hypothesis 2 also cannot account for the observed results, as participants did not reconcile the disagreements by becoming more uncertain about the applicability of the quantifiers. The results are best explained by Hypothesis 3, which states that participants will strengthen their baseline assessment of the applicability of the quantifier after exposure to disagreements deemed to be faultless.

Adaptation effects were however not found in all the percentages where acceptance of faultless disagreement in Block 2 was higher than baseline, namely at 30% and 40%. For 30%, it is possible that the rates of faultless disagreement were still too low to result in detectable semantic updates, as acceptance at 30% was lower than any of the other analyzed percentages. The lack of adaptation effects at 40% is more interesting. The baseline proportion of quantifier acceptance for both *many* and *few* at 40% was 0.5. The acceptance rate of faultless disagreement in Block 2 was also 0.5 at 40%. We argue that this lack of effect can be accounted for by Hypothesis 3; when, on aggregate, participants had maximal uncertainty about the predicate, the opinions expressed in the disagreement dialogue were equally plausible to them, hence the lack of unidirectional reinforcement. We further elaborate on this point in the General Discussion.

A potential confound in the interpretation of the current results is that the adaptation effects could have been the result of the listeners hearing opinions expressed by the speakers rather than the result of hearing such opinions expressed as part of a disagreement dialogue. Experiment 2 has the goal of teasing these two explanations apart.

Experiment 2

Experiment 2 seeks to address a potential confound in Experiment 1. As discussed in the previous section, it is possible that the adaptation effects observed in Experiment 1 were the result of participants simply tracking the production statistics associated with each speaker, independently of the fact that

these different viewpoints were expressed as part of (potentially faultless) disagreement dialogues. To discard this possible interpretation, Experiment 2 tested the same opinions expressed by speakers in Experiment 1 as independent statements rather than as part of a disagreement dialogue.

Materials and Procedure

Experiment 2 followed the same blocked design as in Experiment 1. Block 1 and Block 3 used the same exact stimuli tested in Experiment 1. Block 2 differed from Experiment 1 in that the visual displays were not accompanied by a disagreement dialogue but rather by recordings containing a single speaker uttering assertions such as (4-a) or (4-b):

- (4) a. {Many, few} of the circles are {blue, yellow}.
b. I don't think {many, few} of the circles are {blue, yellow}.

The four speakers in Experiment 1 were used again in Experiment 2. Speaker views remained consistent through all percentages, that is, two speakers always described the image as having *many* or *not few* circles of the target color while the other two speakers always described the image as having *few* or *not many* circles of the target color. Participants were instructed to indicate whether or not they agreed with the speaker by pressing F for *yes* or J for *no* on their keyboards. There were four trials for every sentence-percentage pair, resulting in a total of 112 trials for Block 2. The procedure followed for Blocks 1 and 3 was identical to that of Experiment 1.

Participants

Sixty participants were recruited via Prolific. We required participants to be born and be currently located in the United States. Participants were born and currently located in the United States, and were self-described native English speakers between the ages of 18 and 35. Participants were paid \$3.80, which amounted to an hourly wage of approximately \$12.00.

As in Experiment 1, Experiment 2 included a total of sixteen attention checks in Block 1 and Block 3. The same exclusion criteria used in Experiment 1 were followed in Experiment 2. Two participants were excluded due to failure to meet these criteria. Data from the 58 remaining participants was submitted to the analyses reported below.

Results

Figure 5 contains the proportion of acceptance for *many* and *few* in Block 1 and Block 3. The same analyses performed in Experiment 1 were used to analyze the data for Experiment 2. No significant interactions of QUANTIFIER and BLOCK were detected in any of the analyzed percentages (30%: $\beta = 0.80$; 40%: $\beta = 0.14$; 50%: $\beta = 0.41$; 60%: $\beta = 0.23$, all p 's > 0.05).

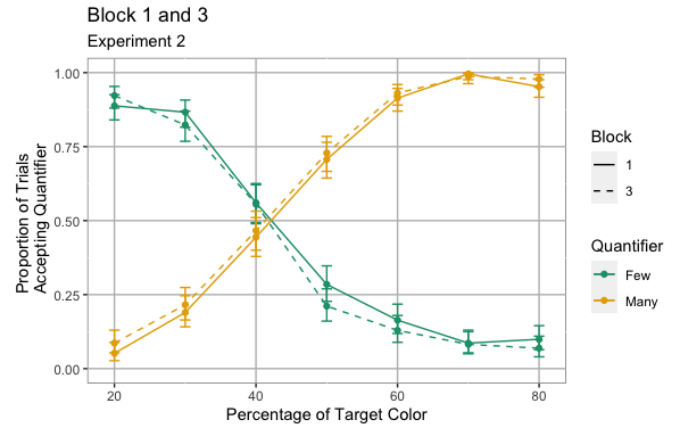


Figure 5: Proportion of acceptance of *many* and *few* at the different percentages tested in Block 1 and Block 3. Error bars correspond to 95% confidence intervals.

Discussion

In light of Experiment 1 results, the lack of adaptation effects for Experiment 2 suggests that the sole fact of exposing participants to inconsistent views expressed in isolation, not as part of a discourse, is not sufficient to induce adaptation effects.

General Discussion

Decades of research in the Gricean pragmatics tradition has consistently shown that speakers and listeners strive to coordinate during conversation in order to converge on the intended utterance interpretation. When it comes to vague language, with its high interspeaker variability, such coordinative behavior is critical for successful communication (Lassiter & Goodman, 2017; Tessler, Tsvilodub, Snedeker, & Levy, 2020). Recent work on semantic and pragmatic adaptation has argued that coordination in the presence of highly variable input is in part possible because listeners are able to track speaker-specific input statistics, and even update their own semantic representations to approximate that of the speaker (Schuster & Degen, 2020; Xiang et al., 2020; Yildirim et al., 2016). But coordination also has its limits. Disagreements are examples of conversational moves that signal precisely the listener's inability or unwillingness to coordinate with the speaker. Remarkably, faultless disagreements do not seem to cause coordination breakdowns. The current paper has investigated potential strategies that listeners might deploy in order to cope with faultless contradictions. Using the case study of vague quantifiers, we examined to what extent listeners seek to minimize contradictions by updating their own semantic representation of the disputed predicates.

Our results show that listeners do indeed adapt their semantic representations in conditions where disagreements are more likely to be perceived as faultless. In particular, all the adaptation effects were such that participants reinforced their

preferences displayed at baseline instead of converging towards 50% acceptability rates. Our results are therefore incompatible with the hypothesis that listeners resolve potentially faultless metalinguistic disagreements by adopting a vaguer threshold—i.e., one that ranges over a wider set of proportions—that would allow them to accommodate the viewpoints expressed in the disagreement dialogue. Instead, our results are more consistent with Hypothesis 3: exposure to metalinguistic disagreements that were on average more likely to be judged as faultless caused participants to further shift their semantic representation towards the viewpoint for which they had initially displayed a preference.

The question remains of how the adaptive behavior displayed by participants at 50% and 60% can be reconciled with the high rates of faultless disagreement observed in those same percentages. We argue that a disagreement can be perceived as faultless while still being asymmetrical, i.e., one view could be perceived as more plausible than the other. If participants identified with the view that they took to be more plausible, which arguably should correspond to baseline preferences in Block 1, the selective adaptation effects observed in Experiment 1 would be expected. Above and beyond the direction of the adaptation effects detected in Experiment 1, a second data point seems to be consistent with this account. We note that at 40% the initial acceptability rates of both quantifiers were essentially at chance (Block 1), indicating peak uncertainty about the applicability of the quantifier. Interestingly, the perceived faultlessness of both speakers was judged to be the highest (about 0.5 acceptance) at this percentage for both quantifiers (Block 2). It is therefore likely that participants perceived the opinions displayed in the dialogue as being equally plausible, thus explaining the lack of adaptation effects observed in this percentage (Block 3).

A second important finding pertains to the lack of adaptation effects in Experiment 2. The goal of Experiment 2 was to rule out the possibility that the selective adaptation patterns displayed by participants in Experiment 1 were unrelated to the fact that the different viewpoints were part of a disagreement dialogue. Taken as a whole, the findings from Experiments 1 and 2 indicate that the observed adaptation effects were indeed the result of exposure to potentially faultless disagreements.

Conclusion

Accepting a disagreement as faultless requires accommodating two contradictory statements. The present study investigated the ways in which comprehenders might circumvent this contradiction. In two adaptation studies, we find that exposure to potentially faultless disagreements led to reduced meaning uncertainty about the applicability of the quantifier in the context. No parallel effects were attested when participants heard speakers express opinions in isolation, i.e., not as part of a disagreement dialogue, implying that the observed adaptation effects were due to the disagreements.

References

- Barker, C. (2002). The dynamics of vagueness. *Linguistics and philosophy*, 1–36.
- Barker, C. (2013). Negotiating taste. *Inquiry*, 56(2-3), 240–257.
- Fernando, T., & Kamp, H. (1996). Expecting many. In *Semantics and linguistic theory* (Vol. 6, pp. 53–68).
- Fleisher, N. (2013). The dynamics of subjectivity. In *Semantics and linguistic theory* (Vol. 23, pp. 276–294).
- Foushee, R., & Srinivasan, M. (2017). Could both be right? children's and adults' sensitivity to subjectivity in language. In *Cogsci*.
- Heim, S., Peiseler, N., & Bekemeier, N. (2020). “few” or “many”? an adaptation level theory account for flexibility in quantifier processing. *Frontiers in Psychology*, 11.
- Horn, L. R. (1989). A natural history of negation. *University of Chicago Press*.
- Kaiser, E., & Rudin, D. (2020). When faultless disagreement is not so faultless: What widely-held opinions can tell us about subjective adjectives. *Proceedings of the Linguistic Society of America*, 5(1), 698–707.
- Kaiser, E., & Rudin, D. (2021). Arguing with experts: Subjective disagreements on matters of taste. In *Proceedings of the annual meeting of the cognitive science society* (Vol. 43).
- Kennedy, C. (2007). Vagueness and grammar: the semantics of relative and absolute gradable adjectives. *Linguistics and Philosophy*, 30(1), 1–45.
- Kennedy, C. (2013). Two sources of subjectivity: Qualitative assessment and dimensional uncertainty. *Inquiry*, 56(2-3), 258–277.
- Kölbel, M. (2004). Faultless disagreement. *Proceedings of the Aristotelian Society*, 104, 53–73.
- Lasnik, P. (2005). Context dependence, disagreement, and predicates of personal taste. *Linguistics and philosophy*, 28, 643–686.
- Lasnik, P. (2009). Relative truth, speaker commitment, and control of implicit arguments. *Synthese*, 166(2), 359–374.
- Lassiter, D., & Goodman, N. D. (2017). Adjectival vagueness in a bayesian model of interpretation. *Synthese*, 194, 3801–3836.
- Partee, B. H. (2004). Many quantifiers. *Compositionality in Formal Semantics*, 241–258.
- Pepper, S., & Prytulak, L. S. (1974). Sometimes frequently means seldom: Context effects in the interpretation of quantitative expressions. *Journal of Research in Personality*, 8(1), 95–101.
- Ramotowska, S., Haaf, J. M., van Maanen, L., & Szymanik, J. (2022). Most quantifiers have many meanings..
- Rudin, D., & Beltrama, A. (2019). Default agreement with subjective assertions. In *Semantics and linguistic theory* (Vol. 29, pp. 82–102).
- Schuster, S., & Degen, J. (2019). Speaker-specific adaptation to variable use of uncertainty expressions. In *Annual*

- meeting of the cognitive science society.*
- Schuster, S., & Degen, J. (2020). I know what you're probably going to say: Listener adaptation to variable use of uncertainty expressions. *Cognition*, 203, 104285.
- Scontras, G., Degen, J., & Goodman, N. D. (2017). Subjectivity Predicts Adjective Ordering Preferences. *Open Mind*, 1(1), 53-66.
- Stalnaker, R. (2002). Common ground. *Linguistics and philosophy*, 25(5/6), 701–721.
- Sundell, T. (2011). Disagreements about taste. *Philosophical Studies*, 155(2), 267–288.
- Tessler, M. H., Tsvilodub, P., Snedeker, J., & Levy, R. P. (2020). Informational goals, sentence structure, and comparison class inference. In *Proceedings of the annual conference of the cognitive science society*.
- Xiang, M., Kramer, A., & Kennedy, C. (2020). Semantic adaptation in gradable adjective interpretation..
- Yildirim, I., Degen, J., Tanenhaus, M. K., & Jaeger, T. F. (2016). Talker-specificity and adaptation in quantifier interpretation. *Journal of Memory and Language*, 87, 128-143.