

UCLA

Experiments in Linguistic Meaning

Title

Investigating a shared mechanism in the priming of manner and quantity implicature

Permalink

<https://escholarship.org/uc/item/5bf0v924>

Journal

Experiments in Linguistic Meaning, 2(1)

Authors

Cowan, Joe

Katsos, Napoleon

Publication Date

2023-01-27

DOI

10.3765/elm.2.5383

Copyright Information

This work is made available under the terms of a Creative Commons Attribution License, available at <https://creativecommons.org/licenses/by/4.0/>

Peer reviewed

Investigating a shared mechanism in the priming of manner and quantity implicature

Joe Cowan & Napoleon Katsos*

Abstract. In the current paper, we investigate the existence of a shared derivation mechanism between manner and quantity implicature. As per the Gricean-inspired perspective, both manner and quantity implicature are derived in a substantially analogous fashion, relying on the consideration of alternative ways in which the speaker could have spoken, but didn't. In contrast, other accounts (e.g., grammatical accounts) of quantity implicature consider manner implicature and quantity implicature to be distinct in their derivational mechanisms.

Previous studies have found that quantity implicature can prime the derivation of subsequent quantity implicature both within and between quantity implicature subtypes in a structural priming paradigm, suggesting that *ad hoc*, *numeral* and *some* quantity implicature are governed by the same derivational mechanism. We have applied a structural priming paradigm to the case of *manner* implicature to investigate 1) whether *manner* implicature can be primed, 2) whether *manner* implicature can prime *manner* implicature and 3) whether *manner* implicature can be primed by *quantity* implicature. Through *manner-manner* priming, the paper addresses the psycholinguistic reality of *manner*. While *quantity-manner* priming probes the existence of a shared derivational mechanism between the phenomena.

We show that *manner* implicature can prime *manner* implicature under certain experimental circumstances and that *ad hoc* *quantity*, but not *some* *quantity* implicature can also prime *manner* implicature, whereas *some* *quantity* implicature cannot.

Keywords. Manner implicature; quantity implicature; scalar implicature; priming.

1. Introduction. Debate exists concerning the mechanism that drives the derivation of quantity implicature. As per accounts based on Grice's (1975) framework, *quantity* implicature (QI) can be considered a pragmatic phenomenon and a constituent of the wider concept of *conversational implicature*. Below are examples QI:

(1) A: '*Some of the cars are green*' +> Not all the cars are green.

(2) A: *Did you meet John and Bill?*

B: '*I met John*' +> I did not meet Bill.

* Authors: Joe Cowan, University of Cambridge (jc2280@cam.ac.uk) & Napoleon Katsos, University of Cambridge (nk248@cam.ac.uk).

As per a Gricean program, the utterance ‘Some of the cars are green’ is logically compatible with a reality in which all the referenced cars are green and a reality in which some, but not all, of the referenced cars are green. As such, there exists a resulting ambiguity, one that must be elucidated by pragmatic means. To disambiguate the utterance, a listener is said to assume that a more informative, relevant statement must be false, for, if it were true, a rational speaker would have surely used it. For instance, the phrase ‘All of the cars are green’ is more informative than the utterance ‘Some of the cars are green’. If the latter were used to denote a reality in which all referenced cars were green, the speaker would be violating Grice’s *cooperative principle*; the mutually assumed commitment to a set of conversational precepts, which, in part, enjoins a speaker to ensure optimal informativity of expression.

Therefore, upon hearing ‘Some of the cars are green’ a listener is said to assume that the speaker means NOT[all of the cars are green] – a negation of the more informative alternative. Then, NOT[all of the cars are green] is incorporated into the understanding of the articulated utterance ‘Some of the cars are green’; the listener arriving at the understanding that the speaker means to communicate that ‘some but not all of the cars are green’. Such is the derivation of quantity implicature as per Gricean-inspired accounts.

Naturally, there are competing alternatives that posit the derivation of quantity implicature as a phenomenon occurring on a compositional level. Chierchia (2006) and Fox (2007) and Chierchia, Fox & Spector (2012) provide accounts of QI that drastically depart from the Gricean program, explaining quantity implicature to be a grammatical phenomenon, whereby a silent *only* operator is present in the syntax of scalars with an analogous meaning to the overtly expressed ‘only’. Take the following:

- (3) ‘Nancy took some of the toys.’
- (4) ‘Nancy took only some of the toys.’

As per a syntax-oriented explanation, the enriched ‘some but not all’ interpretation of *some* is represented by the covert insertion of a silent *O* (*only*) operator:

- (5) Nancy took *O*[some of the toys]

Turning now to *manner* implicature (MI), per Gricean-inspired accounts, *manner* implicature (MI) can be considered an analogous phenomenon; both QI and MI can be said to arise via violation of Grice’s *cooperative principle* and reasoning about alternative utterances the speaker could have said. The maxim of *manner*, a constituent of Grice’s *cooperative principle*, states that speakers must avoid obscurity and ambiguity of expression and ensure apt utterance brevity, therefore avoid prolix or obtuse utterances. Take the following example:

- (6) A: ‘Can you cook?’
B: ‘I am able to mix ingredients together to form something edible’ +> I cannot cook well.

Here, as with QI, a listener is said to understand that the speaker means NOT the alternative, in this case, ‘I can cook’. However, crucially, unlike QI, this alternative is not more informative, rather, less marked than the articulated utterance. Levinson (2000) repackages Grice’s conversational maxims related utterance form (ambiguity and brevity) to explain this via the *M Heuristic*. The *M Heuristic* comprises of two precepts: one that pertains to the speaker, the other to the listener. According to the *M Heuristic*, the speaker adheres to a maxim whereby to reference an abnormal, unusual, or stereotypical situation, one must use marked language that contrasts with

the unmarked language with which one would typically reference a normal corresponding situation. Simultaneously, a listener is to adhere to the acknowledgment that marked expressions pertain to marked situations (Levinson, 2000). In sum, marked language = marked situation.

Under the Gricean-inspired perspective, QI and MI are derived via the same mechanism; one that relies on consideration, and subsequent negation of, alternative ways in which a speaker could have spoken, but chose not to. QI and MI however diverge regarding the nature of the alternative. In QI, the alternative is more informative while in MI it is less marked. Moreover, whereas in QI the meaning of the alternative is negated, in MI the meaning of the alternative is not negated by the MI, rather the implicated meaning, i.e., the stereotypical connotations of the utterance. Therefore, when choosing to use the prolix ‘I am able to mix ingredients together to form something edible’ a speaker could intend to negate the stereotypical connotations of ‘I can cook’: ‘I can cook edible food’, ‘I can cook visually appealing food’ or ‘I can cook food considered typical or normal of my culture’, etc. Despite the similarities between the nature of the alternative, as per a Gricean-inspired perspective substantial similarities exist between the derivation of QI and MI.

In contrast, the mechanisms governing QI and MI derivation diverge fully if we are to consider QI a grammatical phenomenon. Given that MI concerns not *what* is said, rather *how* it is said, MI is necessarily post-compositional and cannot be explained by the existence of a grammatical operator. QI, on the other hand, is derived by a silent *O*, a grammatical operator dedicated to scalar implicature.

Taking stock, the relation between QI and MI is up for debate; either the derivation of QI and MI implicature are governed by the same pragmatic mechanism, or QI and MI are governed by distinct mechanisms. Moreover, given that a compositional account of QI needn’t preclude any subsequent pragmatic enrichment, it could also be the case that QI and MI diverge in their origin – QI perhaps occurring on a syntactic level – but converge in their appeal to pragmatics to expound their meaning.

2. Background. The nature of the mechanism governing the derivation of QI, and the degree to which this mechanism is shared between subtypes of QI, has been the subject of experimental exploration using structural priming techniques.

Bott & Chemla (2016) investigate the extent to which priming can evidence a shared derivational mechanism between QI subtypes; *some*, *numeral*, *ad hoc*:

- (7) A: ‘*Some of the cats are sleeping*’ +> not all the cats are sleeping. (*Some*)
- (8) A: ‘*I broke two of my fingers*’ +> I broke exactly two of my fingers. (*Numeral*)
- (9) A: ‘*There is a dog in the garden*’ +> There is a dog that isn’t mine in the garden. (*Ad hoc*)

According to Bott & Chemla, the above are all examples of QI.

Bott & Chemla employed a reimagination of Huang, Spelke & Snedeker’s (2013) ‘hidden box’ paradigm to investigate a theorized shared mechanism. The paradigm consisted of priming trials and critical trials in a *prime #1 – prime #2 – critical trial* order. In all trials, the participants were presented with a caption containing a QI, e.g., ‘*There are some arrows*’ and two images and were tasked with matching the caption to one of the displayed images. The images present in the trial were configurations of shapes that either did or didn’t correspond to the QI caption, e.g., a card with 6 arrows and 3 diamonds on it, or a plain card that said ‘*Better Picture?*’.

In the priming trials, participants were presented with weak primes or strong primes. The weak primes consisted of a caption containing a QI implicature, e.g., ‘*There is a triangle*’ (+> There is *only* a triangle) and two images, one image corresponding to the unenriched meaning, e.g., two triangles, and a false image that did not correspond to the caption, e.g., an image of a circle. For lack of a better option, in the weak prime trials, the participants were forced to select an image that matched not with the QI’s enriched meaning, rather the unenriched, non-QI meaning. In contrast, in the strong prime trials, the participants were presented 1) with an image that matched the unenriched QI, e.g., a triangle *and* a star, and 2) an image that matched the enriched QI, e.g., a single triangle. Thus, in the strong prime trials, the participants were given free choice between either deriving or suspending the QI. Given the salience of QI, the participants were exceedingly likely to select the image corresponding to the enriched QI. In the critical trials, the participants were faced with an image matching the unenriched QI and a blank image that read ‘Better Picture?’ – the implication being that the unseen referent of the ‘Better Picture’ card corresponds to the enriched QI.

Assuming the effect of structural priming, Bott & Chemla hypothesize that, in the strong prime condition, the participant’s selection of the image corresponding to the enriched QI, in the priming trial, is to prime the selection of the ‘Better Picture?’ card (the image tacitly corresponding to the enriched QI) in the succeeding critical trials. Therefore, Bott & Chemla predict an increase in the selection of the ‘Better Picture?’ card in the critical trials in the strong prime condition. Bott & Chemla investigate a potential priming effect in two supra-conditions, firstly, within-category priming, whereby the same subtype of QI, e.g., *some*, is used throughout the *prime #1 – prime #2 – critical trial* sequence and a between-priming condition, whereby priming occurs between QI subtypes, e.g., a *some prime #1 – some prime #2 – ad hoc critical trial* sequence. The between-subtype condition investigates the extent to which there exists a shared derivational mechanism between QI subtypes – the rationale being that if one QI subtype primes another, there exists a shared mechanism.

Bott & Chemla’s investigation indeed reports evidence to suggest that QI can prime QI of the same subtype, suggesting that prior derivation of one QI subtype can prime the subsequent derivation of the same QI subtype. Bott & Chemla also report this affect to hold between QI subtypes *some*, *numeral*, *ad hoc*, suggesting that there exists a shared derivational mechanism between QI subtypes. Bott & Chemla conclude that observed priming effect is likely resultative of the priming of an *O*-operator but is compatible with both Gricean-inspired and Grammatical accounts of QI.

Rees & Bott (2018) present an adaptation of Bott & Chemla’s (2016) ‘Better Picture?’ paradigm with a priming condition that increases the salience of the QI implicature’s alternative (e.g., ‘*some of the letters are B*’s alternative = ‘*all of the letters are B*’). In the alternative condition, participants were presented with a sentence containing a QI’s alternative (e.g., ‘*all of the letters are B*’) and two images, one corresponding to the sentence (e.g., an image with six letter Bs) and one not corresponding to the sentence, but corresponding to the configuration of the enriched QI – the QI to which the sentence is the alternative (e.g., three letter Bs and six letter As). Crucially, the alternative prime does not require the explicit derivation of a QI. These alternative primes were tested against the weak and strong primes as per Bott & Chemla (2016) – the experiment investigating whether the use of QI itself, or the tacit salience of the QIs alternative, is responsible for the demonstrated priming effect in Bott & Chemla (2016). Rees & Bott found that the use of an ‘alternative’ prime is just as effective as eliciting a priming effect as a strong prime in the

case of *ad hoc*, *some* and *numeral* QI, suggesting that the derivation of QI is not responsible for the observed efficacy of the strong prime in Bott & Chemla (2016) and Rees & Bott (2018).

Despite the important conclusions of Bott & Chemla (2016) and Rees & Bott (2018), a critical flaw exists in both studies: no baseline rate of QI implicature judgement is provided in either. To address this concern, Marty, Cowan, Romoli, Sudo and Breheny (2021) presented a rerun of the priming experiments in Bott & Chemla (2016) and Rees & Bott (2018), providing novel baselines with which to assess the purported within-category priming effect. Marty et al. find that while the strong prime is indeed a driver of a positive priming effect in the case of *ad hoc* primes when compared to the established baseline, there is no such increase in *some* QI, concluding that the purported increase of *some* QI after strong *some* primes is rather an effect of reverse priming on the part of the weak *some* prime.

Taken together, important conclusions can be drawn from the three studies: 1) both within- and between-category priming exists for all three types of QI. 2) It is most likely the salience of the alternative driving the observed priming effect 3) the findings are compatible with a broad range of accounts of implicature (Gricean-inspired, grammatical, a.o.)

3. Current study. Given that there exists theoretical debate concerning the extent to which QI and MI are analogous phenomena, it follows that structural priming paradigms may also prove worthwhile means with which to investigate MI and the relationship between QI and MI.

As such, the current paper investigates two research questions. Firstly: can MI be primed? – a question, which is intrinsically motivated, given that, to the extent of our knowledge, the structural priming of MI is novel. Secondly: can QI cross-prime MI? – a question motivated by both its novelty and its ability to illuminate the existence of a shared derivational mechanism governing the derivation of QI and MI and to subsequently inform wider theoretical understanding of the phenomena as either distinct, approximate, or homologous. On a broader level, our endeavors contribute to the emergent body of work concerning the experimental investigation of MI.

4. Methodology. The novel operationalization of MI in a structural priming paradigm represents a challenge; a researcher must create a set of functionally identical trials that reflect a necessarily *ad hoc*, context dependent phenomenon – a non-issue in the case of QI, whereby the manifestations of each reiteration of the QI follow an identical structure (e.g., *some of the x are y*).

In a reimagination of the trials presented in Bott & Chemla (2016), Rees & Bott (2018) and Marty, Cowan, Romoli, Sudo and Breheny (2021), the trials in the present study comprised of a caption involving a MI and two images. The type of MI used between the trials was one of markedness, whereby a marked, prolix description selects are marked, unusual referent. In terms of caption, the MI trials comprised of protracted definitions of quotidian, arbitrarily selected, objects, whereby mention of the common label of the object is circumvented (e.g., for the object ‘church’ the caption ‘*Select the picture with a building used for Christian worship*’ is used, see Fig.1). The prolix definitions were abridged, and sometimes paraphrased, variations of Google’s English Dictionary entries, as provided by Oxford Languages (collected in March-May 2021). The entries were paraphrased to avoid the use of inappropriately formal language and were abridged so that they did not extend beyond a single sentence. In each trial, the caption was paired with two images: one representing a marked referent and the other an unmarked referent; the marked referent representing the enriched version of the MI, the marked-to-marked association.

Select the picture with a building used for Christian worship

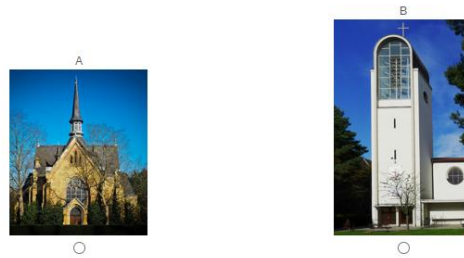


Figure 1: item 'church'

The images were selected from Google Images under the discretion of the researcher, with [object label] and [weird object label] used as search terms. Once a bank of unmarked – marked images were compiled, the items were subject to experimental pre-screenings. Using online survey platform Qualtrics, 45 native speakers of English were asked to provide a description of the 76 images, half were marked items and half were unmarked items in a between-participants presentation mode, ensuring that a participant would not see both the marked and unmarked variations of one object. One image appeared per trial, alongside the phrase ‘*What do you see in this image?*’ and a blank text box in which the participant could submit an answer. The purpose of running the item pre-screening was to ensure that the images selected were indeed appropriately associated with their label, ascertaining each images label rate allowed us exclude images with a < 70% label mention rate.

The rationale for excluding such items that failed this criterion was that, to derive a MI, a participant must be aware that both images pertain to the same label, otherwise the markedness contrast – the crux of the MI – cannot be established. As such, the marked item needs to be a solid referent of its preconceived label; expressing its markedness in a manner which does not compromise its concept belongingness (e.g., in selecting a cat unambiguous in its cat-ness, but nevertheless marked). Resultingly, post item screening, we retained 18 sets of viable unmarked-marked item pairings for use in the subsequent experiments.

5. Baseline. To establish a point of reference from which to interpret the results of the proceeding priming studies, we measured the baseline rates of manner implicature derivation in the 18 selected unmarked-marked pairings. As explained, the experimental items consisted of a caption, a prolix description of the preconceived item label, and two images, both corresponding to the item label, with one of the images a typical example of the item label and the other an atypical example. The participants were asked to select which image best matched the caption with selection of the atypical item taken as indicative of the participant’s derivation of MI, i.e., an atypical caption to atypical image matching. All 18 item pairings appeared in a randomized, within-participants block that also included 10 filler items, whereby each participant provided a judgement on all 18 item pairings.

5.1 PARTICIPANTS. The participant pool comprised of 63 adult monolingual English speakers recruited online via Prolific.co. In all the experiments comprising the current study, participants who had previously participated in any other experiments involving the study were disallowed from participating.

5.2. RESULTS. Fig.2 demonstrates the disparity in MI derivation rates. At the highest point of the spread, ‘*hat*’ has a baseline MI rate of **29.03%**, while at the low end, ‘*hot air balloon*’ shows a

baseline MI deviation rate of **3.23%**, with an average derivation rate of **11.83%** (SD= 8.57%) across all experimental items. The items as they appeared to the experimental participants are shown in Fig.3.

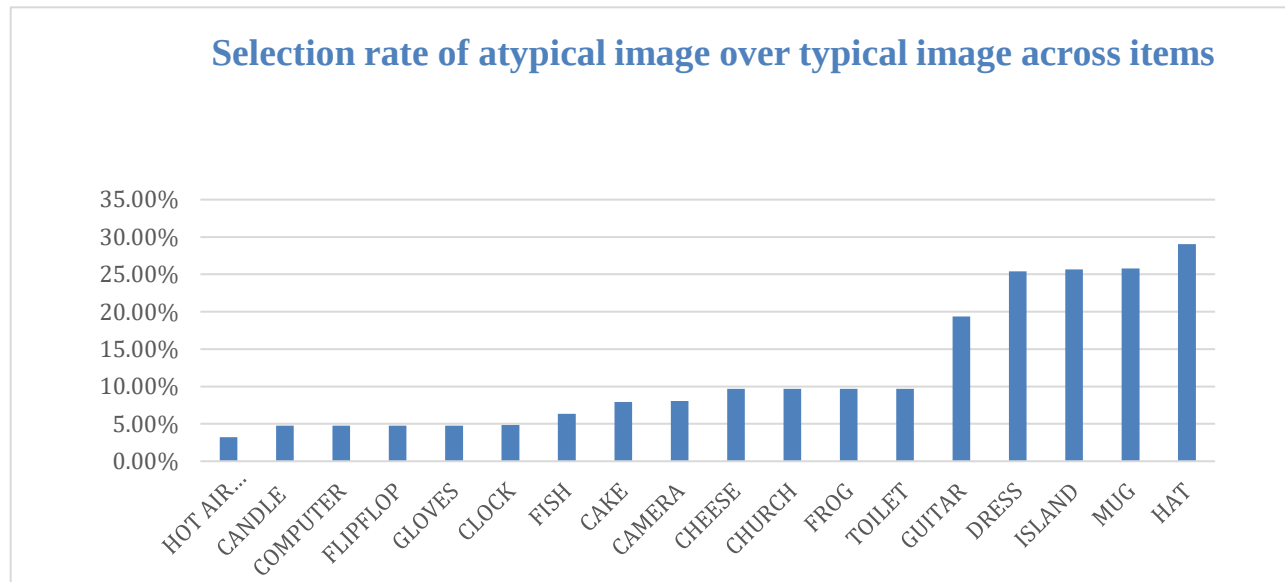


Figure 3: Rate of MI derivation across items

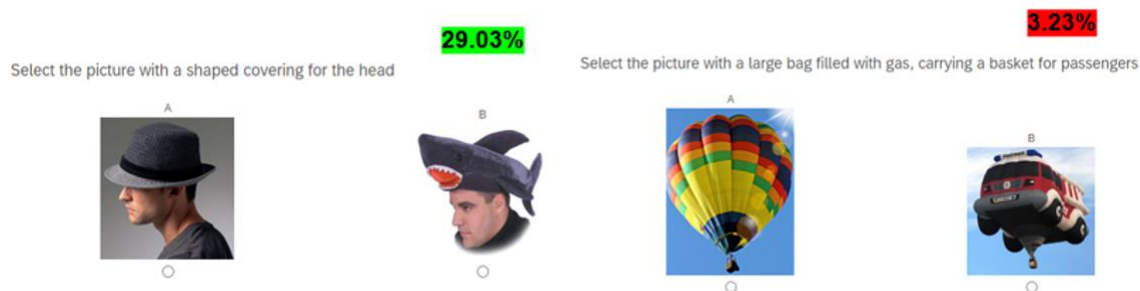


Figure 2: items 'hat' and 'hot air balloon'

5.3. DISCUSSION. MI is, by nature, a heterogenous phenomenon and therefore as the context that guides MI varies, too does MI itself. As such, the baseline rate of MI derivation (taken as the selection rate of the atypical image over the typical image) varies drastically between item pairings. This is expected, as each different MI is essentially a one-off contrast-based implicature influenced by what is said, and how, and what is seen, and is incomparable in uniformity to QI, whereby the implicature relies on more uniform syntax and structure, e.g., 'some of the x are y ' +> 'not all x s are y ' in any imagination of the QI.

The rate of MI derivation is not only heterogeneous but also low compared to the almost ceiling level rates of QI presented in Marty et al. (2021). However, this it to be expected and is consistent with the notion that more inferential and context-dependent reasoning is require in activating the alternative expression in MI leading to reduced derivation (in some items)

6. Experiment 1a. Having established the baseline rate of MI derivation for the 18 prescreened item pairs, we conducted a within-category (MI > MI) priming experiment, addressing the question of MI's ability to be primed.

6.1. STRUCTURE. Experiment 1a followed a *filler – filler – prime #1 – prime #2 – critical trial* order. The critical trials were identical to those presented in the Baseline experiment; they consisted of a free choice between an atypical image and a typical image and a related caption. In contrast, the prime trials, while identical in caption and atypical image to the critical trials, differed in that the typical image was replaced by an unrelated, but equally typical, image of an object e.g., a typical car. This meant that in the prime trials, participants were confronted with a prolix caption and were forced to match it to a marked image, see Fig.4 for an illustration, but were aware of the possibility or existence of unmarked objects within the experiment.

Select the picture with a device for recording visual images



Figure 4: prime 'camera'

The rationale governing this configuration was that the forced selection of the atypical image, and the subsequent derivation of MI, would prime the derivation of MI in the free-choice scenario in the succeeding critical trials à la Bott & Chemla (2016), Rees & Bott (2018) and Marty et al. (2021).

In Experiment 1a, all item pairings were presented both as primes (with an adapted typical image) and as critical trials (as per Baseline) between three trial blocks with participants pooled to one of the three trial blocks. The trial blocks were formulated in a manner whereby each participant ultimately provided judgement on 6 of the overall 18 critical trial pairings while the remaining 12 item pairings functioned as primes.

6.2. PARTICIPANTS. 180 adult monolingual English speakers recruited from Prolific.co.

6.3. RESULTS. In the critical trials of Experiment 1a, the rate of atypical image selection, taken as indicative of MI derivation, was numerically higher at **16.23%** (SD = 12.34%) than our Baseline at **11.86%** (SD= 8.57%), but statistically insignificant as per a Mann-Whitney test $U=126$, $p=0.261$.

6.4. DISCUSSION. The derivation of MI in the prime trials has not significantly increased MI derivation in the critical trials. While it could be the case that MI cannot be primed, considering factors beyond this, the configuration of the prime trials in Experiment 1a could be responsible for the observed lack of priming. There seems to be two reasons the prime trials may be deficient in Experiment 1a, 1) with no certitude can it be said that the participants are engaging in MI in the prime trials; a participant could engage in semantic matching between the caption and the image, irrespective of whether the image is considered marked, circumventing the pragmatic engagement with the caption requisite in triggering MI derivation 2) The matching atypical im-

age may not be considered atypical, or marked, when not appearing next to the less marked, more typical alternative – it seems unlikely that the typicality of, say, a car, can consistently trigger the awareness of the atypicality of, say, a camera.

7. Experiment 1b. Considering the issues with the prime trials of Experiment 1a, we conducted a reiteration with adapted priming trials. In Experiment 1b, the priming trials were almost identical to the critical trials, except a highlighted caption reading ‘THIS IMAGE!’ was placed above the typical image, see Fig.5.



Figure 5: updated item: 'camera'

The experimental instructions were updated so that the participants were told ‘*On some of the images you will be told which image the speaker means. Please select the image they are talking about (it will be labelled 'THIS IMAGE!')*’. The rationale for the adapted prime was that, by including the typical alternative, we contextualize the atypicality of the atypical image, rendering its markedness salient and that the forced selection of the atypical image avoids any purely semantic matching by the participant and instead requires the participant to appeal to pragmatics to rationalize the selection of the atypical image and ultimately arrive at the MI

7.1. PARTICIPANTS. 180 monolingual English speakers recruited from Prolific.ac.uk.

7.2. RESULTS. In the critical trials of Experiment 1b, the rate of atypical image selection was **16.76%** (SD = 8.81%), higher than our Baseline of **11.86%** (SD= 8.57%); a Mann-Whitney test indicating that this finding is statistically significant, $U = 96, p = 0.036$.

7.3. DISCUSSION. With the updated prime trial structure, we can observe a priming effect of the prime trials on MI derivation in the critical trials. While the increase is not dramatic, this could be explained by the pervasive elusivity of MI derivation. As discussed, MI derivation requires the concurrent consideration of various contextual factors, given that this is the case, it could be that, broadly, as derivation of MI is elicited more infrequently than that of QI (as evidenced by the Baseline) even in a primed condition MI derivation remains infrequent. Alternatively, it could be the case that, given the overall low rates of MI derivation, the operationalization of MI in the experimental trials is not robust enough to routinely trigger MI, both in baseline and primed conditions. Given the unnaturalistic context of the trials and the lack of fleshed-out speaker identity, it would be surprising to see lower rates of MI than would be observed ‘in the wild’. Moreover, in either case, it can be said that the increase in atypical image selection is indicative of an increase in MI derivation and, as such, we can conclude that 1) MI can be primed and 2) MI can be primed by MI. Considering evidence of within-category priming of MI, to ad-

dress the questions concerning the relationship between MI and QI, a between-category experimental naturally follows.

8. Experiments 2a & 2b. Considering evidence of within-category priming of MI, to address the questions concerning the relationship between MI and QI, we conducted two between-category priming experiments using *some* and *ad hoc* QI to prime MI.

8.1. STRUCTURE. As in Experiments 1a & 1b, the trials followed a filler – *filler* – *prime #1* – *prime #2* – *critical trial* order, in which the critical trials were identical to those throughout the study at large. Changed were the primes, presented in Fig.5, whereby Experiment 2a concerned *some* QI and 2b *ad hoc* implicature. In Experiment 2a & 2b, the primes were identical in essence to those developed between Bott & Chemla (2016), Rees & Bott (2018) and Marty et al. (2021); *some* primes consisted of two images and a caption. The caption took the form ‘Select the picture in which some of the shapes are [shape]’. Of the images, one matched the enriched QI, e.g., for ‘star’ an image of a card containing three stars and six squares, and the other matching the unenriched alternative e.g., all stars. See Fig. 6.

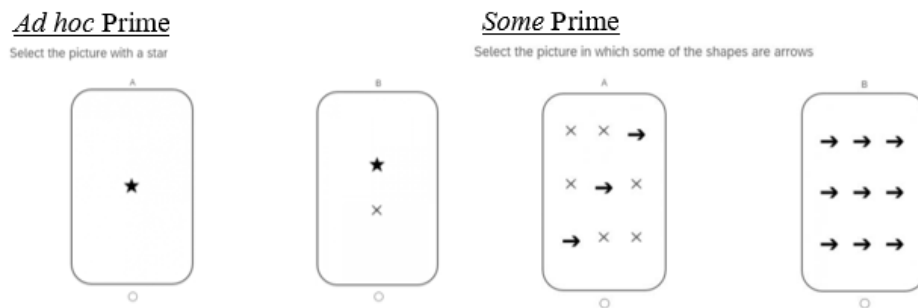


Figure 6: 'ad hoc' and 'some' prime

For Experiment 2b, the *ad hoc* primes had the same caption-image configuration but instead read ‘Select the picture with a [shape]’, with one image matching the enriched QI, i.e., a card containing one instance of a single shape, and the other the unenriched QI, i.e., a card containing two shapes, one pertaining to the caption and the other to a different shape. Given the high baseline rates of QI derivation, the free choice between the enriched or unenriched QI presented on the cards self-selects for the image pertaining to the enriched interpretation and therefore acts as a prime without having to push participants to the selection (e.g., as in 1b with the ‘THIS IMAGE!’ caption). Here, given that MI can be primed, if there is a shared derivational mechanism between QI and MI, we predict that QI derivation in the preceding prime trials will confer greater MI derivation in the succeeding critical trials.

8.2. PARTICIPANTS. 360 adult monolingual English speakers recruited from Prolific.co., split equally between studies.

8.3. RESULTS. In the *ad hoc* prime condition, the MI derivation rate increased to **18.04%** (SD = 12.05%), a significant increase compared to a baseline of (11.86%) as per a Mann-Whitney test, $U=93, p=0.028$, in contrast, while we observed an increase from the baseline to **15.67%** (SD = 9.95%) is observed in *some* QI prime condition, this is not found to reach statistical significance, $U=110, p= 0.102$.

8.4. DISCUSSION. While we gathered evidence of cross-category priming exists, this can only be said to be the case with *ad hoc* QI and MI, not *some* QI and MI. The lack of a priming effect of *some* trials on MI could have been resultative of a more integral issue with the experimental structure, for instance, the lack of a robust conversational context in both the prime and MI critical trials when paired with the conventionalization of *some* QI, and a lack of contextual boundness not found in *ad hoc* QI, could have disengaged the participants from more profound consideration of speaker intention and thus failed to elicit the salience of MI in the eventual critical trials. While this could be the case, the difference between the ability of *some* and *ad hoc* QI to prime MI could be due to a more integral difference between the two types of QI and their relation to MI.

If we are to consider the nature of the alternative in *some* and *ad hoc* QI and MI, it could be the case that the disparate levels of contextual consideration requisite in the search for the alternative are responsible for the disparate potency of *some* and *ad hoc* QI as primes. For instance, the derivation of *some* QI relies on the search for conventional alternatives, i.e., ‘most’ or ‘all’ which are then then negated. In contrast, as is naturally the case with MI, the alternative in the case of *ad hoc* QI is less conventionalized. For instance, take the following sentence and its potential *ad hoc* implicatures:

(10) A: ‘I took Bernie out for a walk’ $\rightarrow$ I didn’t go to the shops/I didn’t walk Rover.

Here, to construct the *ad hoc* scale with which to interpret the utterance, context is needed to determine whether the scale is a scale of tasks completed, or dogs walked. In (10) the alternative is not uniform between the potential *ad hoc* implicatures; ‘I took Bernie out for a walk and did the shopping’ and ‘I took Bernie and Rover out for a walk’ being two alternatives. As such, the nature of *ad hoc* QI’s alternative approximates that of MI; both are reliant on context for their generation. Therefore, it could be the case that the contextual boundness of *ad hoc* QI is what allows for the priming relationship between *ad hoc* QI and MI to occur and the lack of requisite contextual awareness in the generation of the alternative is where a priming effect between *some* QI and MI is thwarted.

9. Overall discussion & conclusion. In general, in both baseline and prime conditions, we see a far lower rate of MI than can be observed for *ad hoc* and *manner* QI. As explained, this aligns with the understanding of MI as a context-dependent, one-off phenomenon; given that MI is derived from context and is not lexically or structurally triggered, it follows that in order for MI to be derived, a robust context, i.e., one in which MI derivation is motivated, is needed. Despite this, in Exp 1b, we did find evidence that MI derivation rates can be augmented by the existence of a MI prime in the context of our experimental setting. The existence of a within-category MI priming effect gives credence to the psycholinguistic reality of markedness contrasts and to MI generally.

In terms of cross-category priming, we observe a small cross-priming effect between *ad hoc* QI and MI implicature but not between *some* QI and MI, Exp. 2a and 2b, see Table 1 and Fig.8, which is compatible with our understanding of the importance of the role of context in the derivation of these three types of implicature. The demonstration of a cross-priming effect is novel across pragmatic maxims, having previously been demonstrated between sub-types of QI (Bott & Chemla, 2016; Rees & Bott, 2018; Marty, et al., 2021). While the existence of this cross-priming effect is compatible with both Gricean-inspired and Grammatical accounts of QI, the existence of a priming effect between *ad hoc* QI and MI suggests that *ad hoc* QI is, at least in

part, a pragmatic phenomenon and that a theoretical dichotomy between *ad hoc* QI and MI as truly distinct phenomena is unjustified.

While the novelties of the presented study have provided evidence suggesting that MI can be primed by MI and by *ad hoc* QI, something that has theoretical ramifications, the same novelties have limited the presented study. Given the rudimentary nature of the operationalization of MI in the trials, it is likely that the availability of MI derivation is being suppressed and that, with preceding iterations of priming studies involving MI, MI can be experimentally constructed in a way in which augments the baseline rate of MI, the effectiveness of MI as a prime and the susceptibility of MI as a subject of priming. A further question raised by the current study is the relationship between some QI and MI; is it the case that some QI cannot prime MI? Or is it the case that, given the lesser degree of similarity between some QI and MI when compared to *ad hoc* QI and MI, that the experimental structure is inhibiting a potential some QI to MI priming effect?

Experiment:	Rate of MI Derivation	Significance in difference from Baseline
Baseline: MI	11.83%	N/A
Experiment 1a: MI – MI prime	16.23%	$P=0.261$
Experiment 1b: MI – MI prime	16.76%	$p= 0.036$
Experiment 2a: <i>some</i> QI – MI prime	15.67%	$p= 0.102$
Experiment 2a: <i>ad hoc</i> QI – MI prime	18.04%	$p=0.028$

Table 1: Collated results

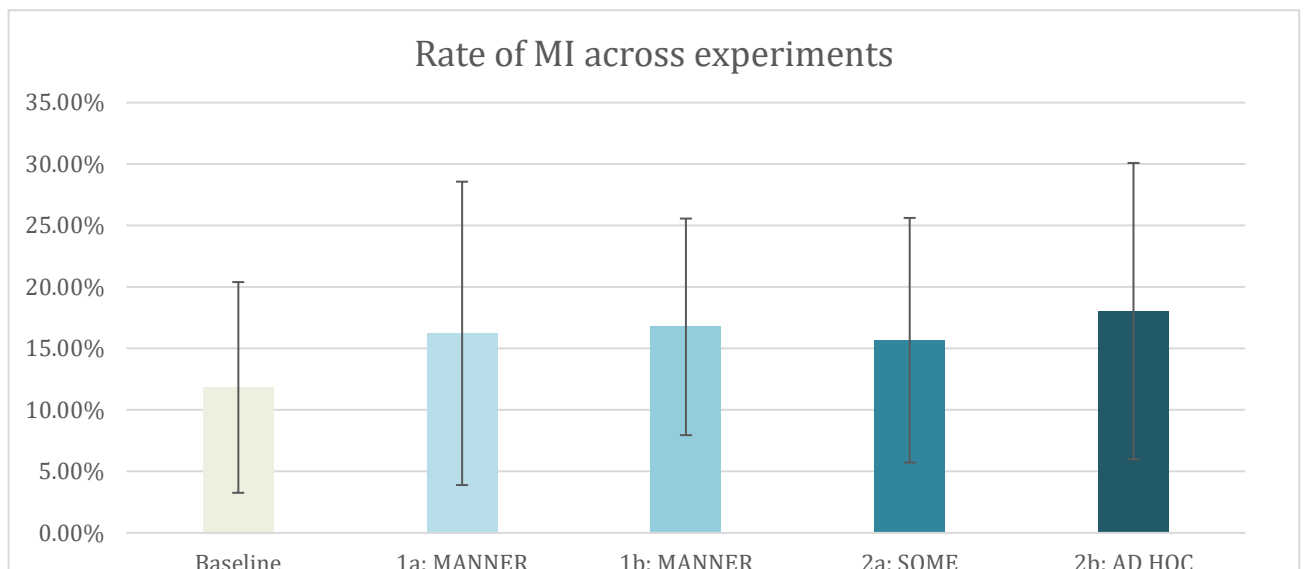


Figure 7: Collated Results

Bibliography:

- Bott, L., & Chemla, E. 2016. Shared and distinct mechanisms in deriving linguistic enrichment. *Journal of Memory and Language*, 91, 117–140. <https://doi.org/10.1016/j.jml.2016.04.00>
- Chierchia, G. 2004. Scalar implicatures, polarity phenomena, and the syntax/pragmatics interface. In *Structures and beyond* (pp. 39–103). Oxford, United Kingdom: Oxford University Press.
- Chierchia, G., Fox, D., & Spector, B. 2012. Scalar implicature as a grammatical phenomenon. In P. Porter, C. Maienborn, & K. von Stechow (Eds.), *An international handbook of natural language meaning* (Vol. 3, pp. 2297–2331). Berlin, Germany: De Gruyter
- Grice, H. P. 1975. Logic and conversation. *Syntax and Semantics 3: Speech Acts*, 26-40
- Huang, Y. T., Spelke, E., & Snedeker, J. 2013. What Exactly do Numbers Mean? *Language Learning and Development*, 9(2), 105–129. <https://doi.org/10.1080/15475441.2012.658731>
- Levinson, S. C. 1983. *Pragmatics* (Cambridge Textbooks in Linguistics). Cambridge, United Kingdom: Cambridge University Press
- Levinson, S. C. 2000. *Presumptive meanings: The theory of generalized conversational implicature*. The MIT Press.
- Marty, P., Cowan, J., Breheny, R., Romoli, J., & Sudo, Y. 2021. What primes what – an experimental framework to explore alternatives for SIs. Paper presentation presented at the 34th Annual CUNY Conference on Human Sentence Processing, New York , United States of America.