

UC Merced

Proceedings of the Annual Meeting of the Cognitive Science Society

Title

A Mentalistic Semantics Explains “Each” and “Every” Quantifier Use

Permalink

<https://escholarship.org/uc/item/5qt582m2>

Journal

Proceedings of the Annual Meeting of the Cognitive Science Society, 44(44)

Authors

Knowlton, Tyler

Trueswell, John

Papafragou, Anna

Publication Date

2022

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

A Mentalistic Semantics Explains “Each” and “Every” Quantifier Use

Tyler Knowlton (tzknowlt@upenn.edu)

MindCORE, University of Pennsylvania, 3740 Hamilton Walk
Philadelphia, PA 19104 USA

John Trueswell (trueswel@psych.upenn.edu)

Department of Psychology, University of Pennsylvania, 425 South University Ave
Philadelphia, PA 19104 USA

Anna Papafragou (anna4@sas.upenn.edu)

Department of Linguistics, University of Pennsylvania, 3401 Walnut Street
Philadelphia, PA 19104 USA

Abstract

“Each” and “every” can be used to express the same truth-conditions but differ in their contexts of use. We adopt a particular psycho-semantic proposal about the meanings of these universal quantifiers: “each” has a meaning that interfaces with the psychological system for representing object-files whereas “every” has a meaning that interfaces with the psychological system for representing ensembles. In five experiments (n=798 total) we demonstrate that this mentalistic account correctly predicts newly-observed constraints on how “each” and “every” are pragmatically used. More generally, these results demonstrate that canonical patterns of language use are affected in predictable ways by fine-grained differences in semantic representations and the cognitive systems to which those representations connect. By treating the output of semantics as mental representations that are more finely articulated than truth-conditions—and by taking seriously the relationship between linguistic meanings and non-linguistic cognitive systems—we can explain otherwise puzzling patterns of language use.

Keywords: Language; psycho-semantics; quantification; pragmatics; quantifier use

1. Introduction

A standard view in linguistic semantics is that expressions are names for things in speakers’ environments (e.g., Davidson, 1967; Montague, 1973; Lewis, 1975; Heim & Kratzer, 1998). For example, the meaning of “frog” is the set of frogs, the meaning of “green” is the set of green things, and the meaning of “every frog is green” is a truth-condition: the set of frogs is related to the set of green things in a particular way. An alternative, mentalistic view holds that meanings are instructions for building thoughts (e.g., Chomsky, 1964; Jackendoff, 1983; Carston, 2008; Pietroski, 2018; Knowlton, Hunter, Odic, Wellwood, Halberda, Pietroski, & Lidz, 2021). On this view, “frog” is a tool for accessing a frog concept, “green” is a tool for accessing a green concept, and putting them both together with “every” yields a complex concept with a particular structure. While this complex concept may

itself be truth-evaluable, it also serves as an instruction to certain cognitive systems (e.g., the system for color processing and the system for representing groups of things).

In this paper, we aim to demonstrate that the mentalistic view about meanings can not only help explain patterns of pragmatic use, but can also help generate novel predictions about which the non-mentalistic view is silent. That is, thinking about meanings as instructions to cognition as opposed to propositions paves the way for connecting semantics and pragmatics to cognitive representations in a psychologically responsible fashion.

As a case study, we consider the English universal quantifiers “each” and “every”. These quantifiers can often be used to describe the very same state of the world. Upon encountering four frogs, all of which are green, it might be appropriate to describe the scene by saying “each frog is green” or by saying “every frog is green”.¹ But at least since Vendler (1962), it has been observed that these very similar quantifiers encourage different contexts of use. An often-reported intuition is that “each” is somehow more individualistic than “every”. This difference can be seen clearly in examples like (1).

- (1) a. Each martini needs an olive.
- b. Every martini needs an olive.

Whereas (1a) calls to mind a scene in which a few particular martinis need garnishes before they’re ready to serve, (1b) is more naturally understood as a general claim or as part of a recipe for making martinis.

Here, we aim to explain why “each” and “every” differ in their contexts of use, building on recent work at the intersection of linguistics and psychology. In particular, Knowlton, Pietroski, Halberda, and Lidz (2021) and Knowlton (2021) propose that “each” and “every” have formally distinct concepts of universal quantification as their meanings (see Figure 1). On this view, a sentence like “every frog is green” is represented in a way that implicates a single

¹ We leave aside the other main universal, “all”, for three reasons. First, “each” and “every” form a more compelling minimal pair (e.g., “all” takes plural agreement as in “all frogs” whereas “each” and “every” both take singular agreement). Second, some work in

linguistics suggests that “all” is not a genuine quantifier, but an intensifier (e.g., Baker, 1995; Partee, 1995). Third, “all” is orders of magnitude more frequent than “each” and “every”, and its uses more varied.

group (*the frogs are such that they are all green*) and naturally interfaces with the cognitive system for representing ensembles (e.g., Ariely, 2001; Whitney & Yamanashi Leib, 2018). The similar sentence “each frog is green” is represented in a way that implicates only individuals (*any thing that’s a frog is such that it is green*) and naturally interfaces with the cognitive system for representing object-files (e.g., Carey, 2009; Kahneman & Treisman, 1984). These proposed representations differ sharply from the standard semantic treatment (Barwise & Cooper, 1981), on which both “each” and “every” are said to express a relation between two independent groups (*the frogs are a subset of the green things*).

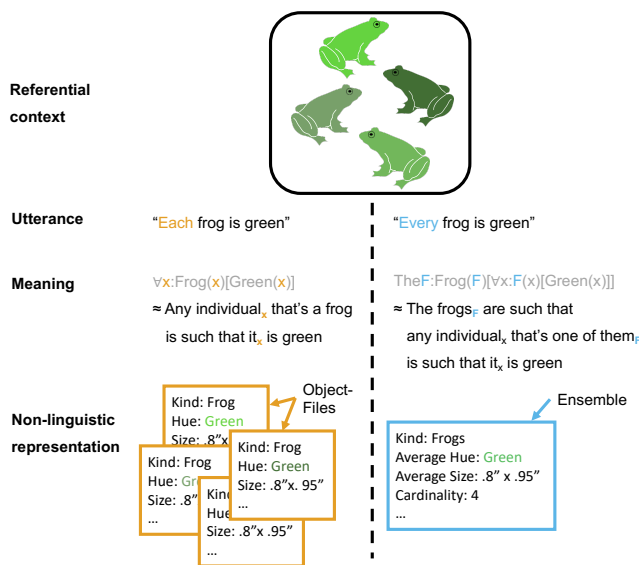


Figure 1: A schematic depiction of the proposed meanings for sentences with “each” / “every” and the resulting non-linguistic representations.

In line with the meaning difference just outlined, in this paper we propose that psychological differences between object-files and ensembles explain many of the pragmatic usage differences between “each” and “every”. Specifically, Section 2 leverages properties of object-files and ensembles to derive three predictions about usage preferences for “each” versus “every”. Section 3 presents the results of five experiments in which these predictions were borne out. Section 4 compares our view to an alternative explanation that retains a more standard semantic picture of quantifier meanings: “each” and “every” have the same meaning, but differ in syntactic position (Beghelli & Stowell, 1997). As we show, such an account cannot capture our findings without additional machinery.

2. Object-Files vs. Ensembles and the “Each”/“Every” Distinction

As noted above, we build off of the proposal that “each” and “every” have different concepts of universal quantification as their meanings (Knowlton et al., 2021b; Knowlton, 2021).

Namely, the pronunciation “each” is connected with a concept that calls for treating the things quantified over as a series of independent individuals (e.g., “each frog is green” roughly means “frog₁ is green & frog₂ is green & ...”). In contrast, “every” serves as an instruction to group the things quantified over (e.g., “the things that are frogs are such that...”). Given that individuals are represented with the object-file system and groups are represented as ensembles, these two universal concepts serve as instructions to two different cognitive systems.

An object-file representation (e.g., Kahneman & Treisman, 1984) is essentially a pointer to an object to which properties are bound (e.g., the object’s size and color). An ensemble representation (e.g., Whitney & Yamanashi Leib, 2018) can also be thought of as a pointer, but one that allows for pointing to many individual objects simultaneously. To accomplish this, ensembles abstract away from individual properties and encode the collection in terms of summary statistics (e.g., the group’s average size, center of mass, and average hue). Both of these cognitive systems are operative in humans as early as infancy and are evolutionarily ancient (for reviews, see Feigenson, Dehaene, & Spelke, 2004; Carey, 2009). Importantly, object-file and ensemble representations need not be representations of visually-presented objects (e.g., both could be formed in response to auditorily-presented tones or could be inferred).

In support of the idea that these cognitive systems underlie the meanings of “each” and “every”, Knowlton et al. (2021b) and Knowlton (2021) report that participants recall group summary statistics (cardinality and center of mass) better when evaluating sentences with “every” but recall individual properties (particular hue) better when evaluating sentences with “each”. Of course, this is not to say people *always* represent object-files upon hearing a sentence with “each” or ensembles upon hearing a sentence with “every”. Supposing meanings are instructions doesn’t guarantee that those instructions will always be followed, and other considerations undoubtedly play a role (e.g., how many things there are). The idea is that the meaning representation carries some weight in determining which non-linguistic representational system will be deployed (Lidz, Pietroski, Halberda, & Hunter, 2011).

Taking this proposal as a starting point, we consider how other properties of object-files and ensembles might influence how “each” and “every” are pragmatically used. Three cases are explored in particular.

First, one important property of object-files is that they are subject to more stringent working memory constraints than ensembles. This working memory limit has often been observed in adults in experiments using multiple object tracking (Pylyshyn & Storm, 1988) and change detection (Vogel, Woodman, & Luck, 2001) paradigms. In both cases, performance sharply declines when participants are asked to track five or more objects at once. In infants, the working memory limit for representing multiple object-files has also been well-documented. For example, Feigenson and Carey (2005) show that infants fail to distinguish 1 from 4 objects

when those objects are represented as object-files, despite the fact that infants the same age can reliably distinguish 8 from 16 objects when treating them as an ensemble (Xu & Spelke, 2000). Wood and Spelke (2005) likewise show that 6-month-old infants successfully distinguish 8 versus 4 actions (grouped into two ensembles) but fail to distinguish 4 versus 2 actions (when treated as 6 independent “event-files”).

This working memory difference gives rise to our first prediction. Ensembles are better able to represent large numbers of things, so “every”—whose meaning serves as a call to represent an ensemble—should be preferred over “each” when the domain of quantification is large.

The second relevant difference to consider is that object-file representations treat individuals independently of one another whereas ensemble representations describe many individuals at once. As noted above, this is achieved by abstracting away from individual properties and representing the group’s summary statistics (for review, see Haberman & Whitney, 2012). For example, Haberman and Whitney (2011) show that participants can tell which of two collections of faces are happier on average despite not being able to identify which individual faces changed from one image to the next. In other words, representing some faces as an ensemble amounts to representing properties like their average hue, their average size, and their average happiness at the cost of encoding information about each individual face. This mode of representation naturally licenses a prediction about new members of the group (they will have similar properties to the group average) and it suggests somewhat vague boundaries about which particular things are included in the group.

In short, ensembles more easily support generalization than object-files. This leads to our second prediction: “every” should be preferred when universal quantification is meant to extend beyond the locally-established domain (as in the natural understanding of “every martini needs an olive”).

The third prediction is perhaps the easiest to draw. In the linguistics literature, both “each” and “every” are often said to be bad with collective predicates, which necessarily apply to groups (e.g., Beghelli & Stowell, 1997; Dowty, 1987). For example, sentences like “each/every student gathered in the hall” are said to be infelicitous (cf. “all the students gathered in the hall”). But if “every” does have a meaning that implicates a group representation, then “every” should be better than “each” at combining with collective predicates (even though both are distributive universals).

3. Experiments

The following five experiments test the predictions outlined above. Experiments were designed using PCIBex (Zehr & Schwarz, 2018). All participants were recruited on Prolific and gave informed consent prior to participating.

3.1 Experiment 1a: Domain size (preference)

Experiment 1a tests the prediction that the size of the domain of quantification should impact preferences for using “each” or “every”. In particular, “every” should be preferred given

larger domains, whereas “each” should be preferred given smaller domains (due to the strict working memory limits of object-files).

3.1.1 Methods Participants in this and all subsequent experiments were native English speakers living in the United States. In this experiment, participants (n=100) completed a forced-choice judgment task. They chose between “each” and “every” for 12 sentences in minimally-different pragmatic contexts, manipulated within-subjects. Context either established a small or large domain of quantification. For example, the domain consisted of either “three martinis” or “three thousand martinis”, as in (2).

- (2) a. The bartender at the local tavern has made three martinis. He said that {each/every} martini he made had an olive.
- b. The bartender at the local tavern has made three thousand martinis. He said that {each/every} martini he made had an olive.

For each item, participants received either the large domain or small domain context (e.g., they either saw (2a) or (2b)) and were given a choice between “each” and “every”.

The experiment also included 24 fillers, all of which were cases where the two possible answers differed in form but not in truth-conditional content. For example, the choice might be between “In her favorite book, the main character is a talking dog” and “The main character is a talking dog in her favorite book”.

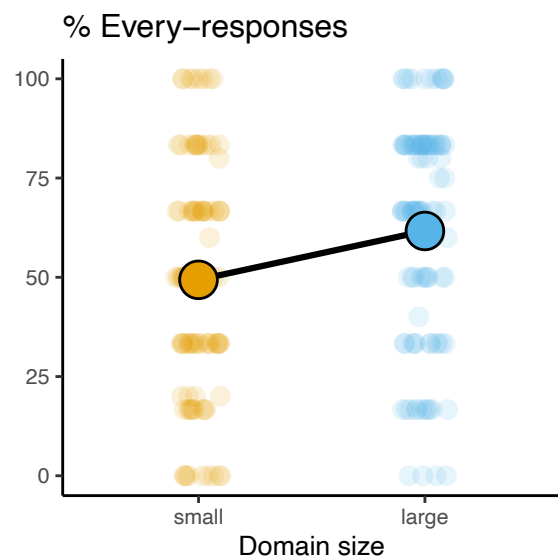


Figure 2: Rate of choosing “every” in the two conditions of Experiment 1a, which manipulated domain size (e.g., “three martinis” vs. “three thousand martinis”).

3.1.2 Results Responses were analyzed by fitting a mixed-effects binomial regression model, specified as follows: answeredEvery ~ context + (1 | subject) + (1 | item). Models

with random intercepts and random slopes for subject and item were also fit, but did not converge (the maximal model that converged across Experiments 1a, 2a, and 3 was one with random intercepts only).

As seen in Figure 2, participants were significantly more likely to pick “every” when given a sentence with a large domain compared to when given a sentence with a small domain ($p < .001$). This is in line with the predictions outlined in Section 2.

It should be noted though, that participants were not more likely than chance to use “each” for small domains (they selected “each” about 51% of the time in this condition). This likely reflects a baseline preference for “every” (perhaps owing to frequency). The important point for our purposes is that despite any baseline preference for “every”, manipulating domain size matters in the predicted way: increasing the domain size encourages using “every”.

3.2 Experiment 1b: Domain size (free response)

Experiment 1b aims to confirm the domain size differences observed in Experiment 1a in a more direct way: simply asking participants how large they think the domain of quantification is upon being given a sentence in which universal quantification is indicated by “each” or “every”.

3.2.1 Methods An independent group of participants ($n=198$) was given the prompt in (3) and asked to type in an answer. The only difference between the two conditions was the quantifier used (manipulated between-subjects).

- (3) If someone said: “{each/every} martini needs an olive”, how many martinis would you guess they have in mind?

3.2.2 Results Participants were more likely to provide an answer within the working memory limit for object-files—three or fewer—in the “each” condition than in the “every” condition ($\chi^2=11.97$, $p < .001$; Table 1).

Table 1: Participants’ responses to being asked the question in (3). “ ∞ ” refers to responses like “infinitely many”. “Exhaustive” responses include answers like “all of them” and “all martinis”. Seven participants did not provide codable responses (e.g., saying “I don’t know”).

	≤ 3	4-5	≥ 6	∞	Exhaustive
<i>Each</i>	62	10	12	0	9
<i>Every</i>	29	13	21	5	30

This corroborates the findings from Experiment 1a: all else equal, “each” is preferred for smaller domains of quantification whereas “every” is preferred for larger domains. This result also accords with how parents prefer to use “each” and “every” in child-directed speech (Knowlton & Gomes, 2022): more often than not, “each” is used to quantify over three or fewer physically-present things.

3.3 Experiment 2a: Generalization (preference)

Experiment 2a tests the prediction that “every” should be preferred in contexts that call for projecting beyond the locally-established domain (i.e., generalizing).

3.3.1 Methods Participants ($n=100$) completed a forced-choice judgement experiment. The method was the same as in Experiment 1a, with one exception: instead of manipulating the size of the domain (e.g., “three” vs. “three thousand”), the size was held constant across both conditions. What differed was whether the quantificational phrase referred back to the domain or explicitly went beyond it. For example, the “local” condition in (4a) refers back to martinis that the bartender made, whereas the “global” condition in (4b) projects beyond this locally-established domain to make a general claim about any martini worth drinking.

- (4) a. The bartender at the local tavern made a few martinis. He said that {each/every} martini that he made has an olive.
 b. The bartender at the local tavern made a few martinis. He said that {each/every} martini that’s worth drinking has an olive.

As in Experiment 1a, context (here, “local” vs. “global”) was manipulated within-subjects. But for any given item, participants saw one context and were asked to choose between “each” and “every”.

3.3.2 Results We adopted the same analytical strategy as in Experiment 1a. As seen in Figure 3, participants were more likely to pick “every” when quantification projected beyond the locally-established domain ($p < .001$).

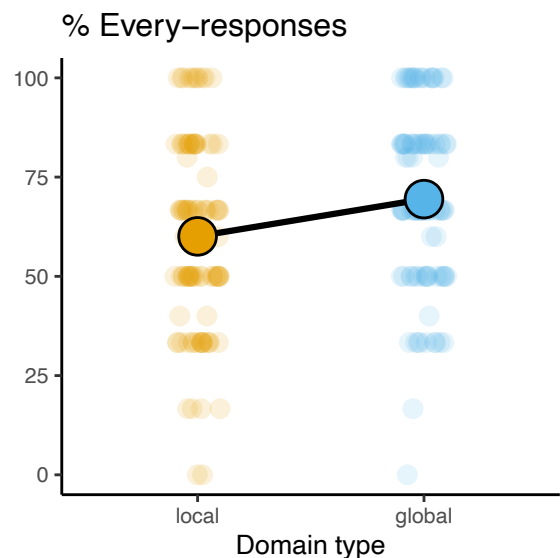


Figure 3: Rate of choosing “every” in the two conditions of Experiment 2a, which manipulated domain type (e.g., “martini that he made” vs. “martini that’s worth drinking”).

These results bear out the prediction that “every” should be preferred in cases calling for generalization. However, there is a potential confound with the stimuli used: the “global” cases implicate larger domains than the “local” cases. For example, the martinis worth drinking likely outnumber the martinis that the bartender just made. For this reason, it might be that the results of Experiment 2a reduce to another case of domain size mattering (as in Experiments 1a-b). Experiment 2b aims to rule out this confound by probing propensity for generalization in a different way.

3.4 Experiment 2b: Generalization (confidence)

Experiment 2b also tests the prediction that “every” should be preferred for generalization beyond the locally-established domain. But instead of asking participants to choose between “each” and “every”, they were either given sentences with “each” or “every” and subsequently asked to generalize. Their propensity to generalize and their confidence in drawing generalizations was compared. The prediction is that “every” should lead to higher rates of generalization and/or higher confidence when generalizing.

3.4.1 Methods Participants ($n=300$) completed a one-trial experiment with the structure in Figure 4. They were shown three novel creatures called “daxes” and were subsequently told that {each/every} dax is green (manipulated between-subjects). Then, they were shown the silhouette of another dax, whose color was hidden by a shadow. They were asked whether they thought that dax was also green (possible responses were “Green” or “Could be another color”). Finally, they were given a slider ranging from 0 to 100 and asked to indicate how confident they were in their response.

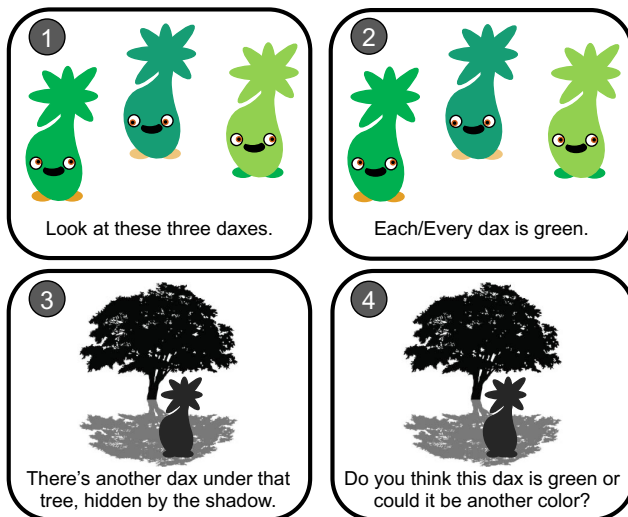


Figure 4: Trial structure of Experiment 2b. After answering the question about the novel dax’s color (4), they were asked to rate their confidence using a slider.

² Some “purely” collective predicates cannot combine with any quantified subject. For example, “each/every ant in my kitchen is

3.4.2 Results Statistically, rates of generalizing (i.e., answering “green” as opposed to “could be another color”) did not significantly differ (“each”: 61%; “every”: 70%; $p=.09$). But participants’ confidence in generalizing did differ. Those in the “every” condition were significantly more confident in generalizing about the unseen dax’s color than those in the “each” condition (i.e., among “green” responses, confidence was greater for the “every” as compared to the “each” condition, $p<.01$; Figure 5).

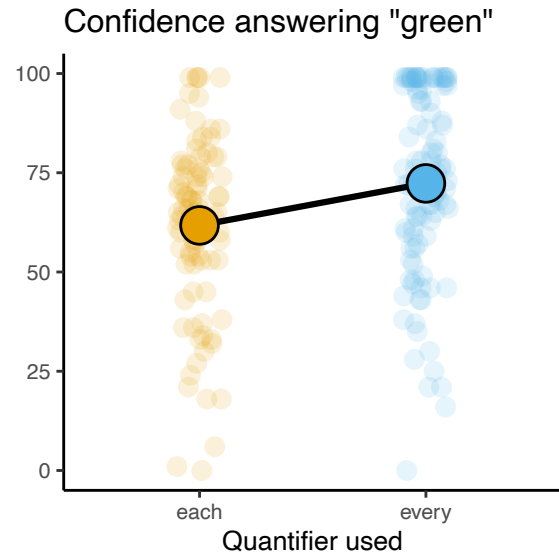


Figure 5: Confidence in answering that the unseen dax was green in Experiment 2b.

This corroborates the findings from Experiment 2a: “every” is preferred for generalizing beyond the locally-established domain (in this case, the initial three daxes). And importantly, Experiment 2b had no domain size confound. In both cases, the domain directly quantified over was the three daxes present on the screen. But introducing those creatures with “every” instead of “each” made participants more comfortable generalizing to a novel instance.

3.5 Experiment 3: Predicate type

Experiment 3 tested the final prediction raised in Section 2: given that only “every” calls for creating a group representation of the things being quantified over, it should be preferred when the predicate in question is collective; that is, when it necessarily applies to a group (e.g., “gather in the hall”; “surround the castle”; “form a line around the block”).²

3.5.1 Methods Participants ($n=100$) completed a forced-choice judgement experiment like Experiments 1a and 2a. What differed between conditions was the predicate used in the second sentence. In the distributive condition, predicates

numerous” is ungrammatical, as is “all/most/some/no ants in my kitchen are numerous” (see Winter, 2002). These were avoided.

applied to individuals, like “went to their locker” in (5a). In the collective condition, predicates necessarily applied to groups, like “gathered in the hall” in (5b). That is, a single student cannot gather. As in Experiments 1a and 2a, participants’ task was to choose between “each” and “every”.

- (5) a. Math class at the local middle school lasts a full hour. After class, {each/every} student went to their locker.
 b. Math class at the local middle school lasts a full hour. After class, {each/every} student gathered in the hall.

3.5.2 Results The analytical strategy from Experiments 1a and 2a was adopted here. As seen in Figure 6, participants were more likely to pick “every” when the predicate was collective than when it was distributive ($p < .001$).

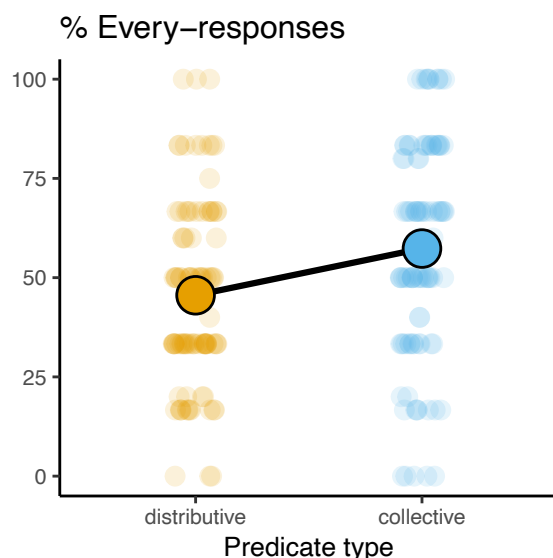


Figure 6: Rate of choosing “every” in the two conditions of Experiment 3, which manipulated predicate type (e.g., “went to their locker” vs. “gathered in the hall”).

These data alone do not militate against the standard assumption that “every” is a distributive universal (like “each”). But, as predicated, they suggest that “every” is better able than “each” to combine with collective predicates.

4. General Discussion

Across five experiments, we observe that “every” is preferred over “each” (i) when the domain of quantification is large as opposed to small, (ii) when quantification is meant to generalize beyond the locally-established domain, and (iii) when the quantificational phrase combines with a collective predicate as opposed to a distributive one. These differences in the pragmatic use of “each” and “every” are predicted by the psycho-semantic view outlined in Section 2. That is, they can be derived from the properties of object-files and

ensembles, the two cognitive systems that, by hypothesis, interface with the meanings of “each” and “every”.

4.1 A syntactic alternative

The proposal adopted throughout this paper posits different meanings for “each” and “every”. As noted at the outset, it is not standard in semantics to assume that these two universal quantifiers correspond to distinct concepts. To take a classic example, Beghelli and Stowell (1997) offer a view on which “each” and “every” share a common meaning and only differ in a syntactic feature. It is worth asking whether a purely syntactic view might accommodate the above results.

The gist of Beghelli and Stowell’s proposal is this: “each” is marked with a diacritic, causing it to undergo syntactic movement to associate with a distributivity operator, which is situated relatively high in the syntactic tree. The distributivity operator is responsible for ensuring that predicates apply to individuals (i.e., it essentially has the meaning we propose for “each”). The distributivity operator also happens to be located higher in the syntactic tree than the generic operator, which is responsible for giving sentences generic meanings (i.e., meanings that project beyond the locally-established domain). In contrast, “every” can remain lower in the syntactic tree, beneath these operators.

Because “each” always associates with the distributivity operator whereas “every” only sometimes does, this purely syntactic view can potentially capture the finding that “each” is worse with collective predicates. And assuming the generic operator gives a generic interpretation only to things within its scope (below it), this view can also potentially capture the finding that “every” is preferred for generalizing beyond locally-established domains. But it is not clear how a purely syntactic view could explain the observed difference with respect to domain size. Being lower in the tree than “each” does not obviously imply that “every” would be preferred for larger domains. And it is likewise non-obvious that the propensity of “each” to associate with the distributivity operator would help explain the domain size difference.

Of course, the scoped-based and psycho-semantic views are not mutually exclusive. The right approach to differences between “each” and “every” might be to preserve a role for purely linguistic machinery while still pushing some of the explanatory burden to non-linguistic cognition (e.g., the properties of object-file and ensemble representations). Future work will consider this possibility in earnest.

4.2 Conclusion

The present results provide a case study in linking the psycho-semantics and pragmatics of quantifiers. Importantly, we do not endorse a reductionistic view that obviates the role of linguistic meaning. In contrast, we hope to have shown that thinking of linguistic meanings as finely-articulated mental representations and taking seriously the cognitive systems with which they interface can help explain otherwise puzzling patterns of linguistic use. Moreover, integrating semantics, pragmatics, and cognition in this way can allow for making and testing novel predictions.

Acknowledgments

For helpful discussion throughout this project, we thank Jeffrey Lidz, Paul Pietroski, Justin Halberda, Zoe Ovans, Victor Gomes, and Sandy LaTourrette.

References

- Ariely, D. (2001). Seeing sets: Representation by statistical properties. *Psychological Science*, 12(2):157–162.
- Baker, M. C. (1995). On the absence of certain quantifiers in Mohawk. In E. Bach, E. Jelinek, A. Kratzer, & B. H. Partee (Eds.), *Quantification in natural languages* (pp. 21–58). Dordrecht: Kluwer.
- Barwise, J. and Cooper, R. (1981). Generalized quantifiers and natural language. In *Philosophy, Language, and Artificial Intelligence*, (pp. 241–301). Springer.
- Beghelli, F. and Stowell, T. (1997). Distributivity and negation: The syntax of each and every. In Szabolcsi, A. (Ed.), *Ways of Scope Taking*, (pp. 71–107). Springer.
- Carey, S. (2009). *The Origin of Concepts*. Oxford Series in Cognitive Development. Oxford University Press.
- Carston, R. (2008). Linguistic communication and the semantics/pragmatics distinction. *Synthese*, 165(3):321–345.
- Chomsky, N. (1964). *Current Issues in Linguistic Theory*. De Gruyter Mouton.
- Davidson, D. (1967). Truth and meaning. In *Philosophy, Language, and Artificial Intelligence*, (pp. 93–111). Springer.
- Dowty, D. (1987). Collective predicates, distributive predicates, and all. In *Proceedings of the 3rd Eastern States Conference on Linguistics* (pp. 97–115). Ohio State University.
- Feigenson, L., Dehaene, S., and Spelke, E. (2004). Core systems of number. *Trends in Cognitive Sciences*, 8(7):307–314.
- Feigenson, L., and Carey, S. (2005). On the limits of infants' quantification of small object arrays. *Cognition*, 97(3), 295–313.
- Haberman, J. and Whitney, D. (2011). Efficient summary statistical representation when change localization fails. *Psychonomic Bulletin & Review*, 18(5):855–859.
- Haberman, J. and Whitney, D. (2012). Ensemble perception: Summarizing the scene and broadening the limits of visual processing. In Wolfe, J. and Robertson, L. (Eds.), *From Perception to Consciousness: Searching with Anne Treisman*, (pp. 339–349). Oxford University Press.
- Heim, I. and Kratzer, A. (1998). *Semantics in Generative Grammar*. Blackwell Oxford.
- Jackendoff, R. (1983) *Semantics and cognition* (Vol. 8). MIT press.
- Kahneman, D. and Treisman, A. (1984). Changing views of attention and automaticity in varieties of attention. In Parasuraman, R. and Davies, D. R. (Eds.), *Varieties of Attention* (pp. 29–61). Academic Press.
- Knowlton, T. (2021). *The psycho-logic of universal quantifiers*. Doctoral dissertation, Linguistics Department, University of Maryland, College Park.
- Knowlton, T. and Gomes, V. (2022). Linguistic and non-linguistic cues to acquiring the strong distributivity of “each”. *Proceedings of the Linguistic Society of America (PLSA)*, 7(1), 5236.
- Knowlton, T., Hunter, T., Odic, D., Wellwood, A., Halberda, J., Pietroski, P., and Lidz, J. (2021a). Linguistic meanings as cognitive instructions. *Annals of the New York Academy of Sciences*.
- Knowlton, Pietroski, Halberda, and Lidz (2021b). The mental representation of universal quantifiers. *Linguistics and Philosophy*.
- Lewis, D. (1975). Languages and language. In Gunderson, K. (Ed.), *Minnesota Studies in the Philosophy of Science*, volume 7. University of Minnesota Press.
- Lidz, J., Pietroski, P., Halberda, J., and Hunter, T. (2011). Interface transparency and the psychosemantics of most. *Natural Language Semantics*, 19(3):227–256.
- Montague, R. (1973). The proper treatment of quantification in ordinary English. In *Approaches to Natural Language*, pages 221–242. Springer.
- Partee, B. (1995). Quantificational Structures and Compositionality, In E. Bach, E. Jelinek, A. Kratzer, & B. H. Partee (Eds.), *Quantification in natural languages* (pp. 541–601).
- Pietroski, P. M. (2018). *Conjoining Meanings: Semantics Without Truth Values*. Oxford University Press.
- Pylyshyn, Z. W. and Storm, R. W. (1988). Tracking multiple independent targets: Evidence for a parallel tracking mechanism. *Spatial Vision*, 3(3):179–197.
- Vendler, Z. (1962). Each and every, any and all. *Mind*, 71(282):145–160.
- Vogel, E. K., Woodman, G. F., and Luck, S. J. (2001). Storage of features, conjunctions, and objects in visual working memory. *Journal of Experimental Psychology: Human Perception and Performance*, 27(1):92.
- Whitney, D. and Yamanashi Leib, A. (2018). Ensemble perception. *Annual Review of Psychology*, 69:105–129.
- Winter, Y. (2002). Atoms and sets: A characterization of semantic number. *Linguistic Inquiry*, 33(3), 493–505.
- Wood, J. N., and Spelke, E. S. (2005). Infants' enumeration of actions: Numerical discrimination and its signature limits. *Developmental science*, 8(2), 173–181.
- Xu, F., and Spelke, E. S. (2000). Large number discrimination in 6-month-old infants. *Cognition*, 74(1), B1–B11.
- Zehr, J., and Schwarz, F. (2018) PennController for Internet Based Experiments (IBEX).