

Young Children's Developing Expectations about the Language of Events

Kristen Syrett and Sudha Arunachalam

1. Introduction

Speakers routinely provide cues in their utterances about their intended meaning that a listener must retrieve using inferential processes that go beyond pure lexical semantics and semantic composition. A classic example of this speaker-hearer meaning negotiation comes from conversational implicatures (Grice 1975). To take two well-known cases, if a speaker delivers the utterance in (1a) with the existential quantifier *some*, the hearer might infer that the toddler did *not* eat *all* of the broccoli. Similarly, if a speaker delivers the utterance in (1b) with disjunction (*or*), a hearer might infer that the toddler ate *either* the broccoli *or* the peas, but *not both*.

- (1) a. The toddler ate some of the broccoli.
- b. The toddler ate the broccoli or the peas.

The listener who calculates these so-called *scalar implicatures* (Grice 1975; Horn 2006) is claimed to go through roughly the following reasoning process. First, the listener assumes that the speaker is being cooperative (i.e., adhering to the Cooperative Principle) and delivering a maximally informative utterance. Second, let us assume that lexical items such as *some* and *or* participate in a scale, where alternatives are ordered by lexical entailment ($\langle e_1 \dots e_n \rangle$). If a speaker has a choice of selecting any of the scalar alternatives, including a stronger lexical alternative (e_1), but delivers an utterance with a weaker alternative (e_2), this must be either because s/he knows that the stronger alternative does not hold, or does not know if it holds. Thus, while the speaker did not explicitly indicate the quantity by stating *some but not all* or signal exclusive disjunction by stating *one or the other and not both*, the listener can confidently infer that this is the intended meaning (Geurts 2010; Grice 1975; Horn 1984; Levinson 1983, 2000).

A wide range of experimental evidence, which we do not have the space to review here, has provided evidence that both adults and older preschool age

* Kristen Syrett, Rutgers University–New Brunswick (kristen.syrett@rutgers.edu); Sudha Arunachalam, Boston University (sarunach@bu.edu). This research was supported by start-up funds to KS and NIH K01DC013306 to SA. We are grateful to our research assistants, the participating preschools and families, and the audience at BUCLD 40 for their insightful comments and questions.

children compute such *generalized conversational implicatures* (GCIs), although there is significant variability in their ability to do so and the rapidity with which they do so, depending on experimental methodology, the target lexical items, and the context in which the utterance is delivered.

Until now, the vast majority of work on implicature calculation has focused on entailment-based scales (and more specifically on implicatures associated with certain lexical items, *some* in particular). There is, however, a second class of conversational implicatures—*particularized conversational implicatures* (PCIs)—which are context dependent: they are not associated with specific lexical items, and the accessibility of the implicature depends on specific aspects of the context at hand (Levinson 2000). Take, for example, the utterance in (2).

(2) Ashley is crying.

A speaker who utters (2) in the context of a bustling preschool classroom might merely wish to assert that Ashley is crying. However, the speaker might also intend for the listener to retrieve the meaning that *Ashley and no other preschooler in the room* is crying. In both cases, it is true that Ashley is crying, but with the second interpretation, there is an extra layer of pragmatic meaning that is not encoded in the semantics. From a semantic perspective, the assertion is true as long as $\llbracket \text{Ashley} \rrbracket$ is in the extension of $\llbracket \text{crying} \rrbracket$ in that context.

How would one calculate this second meaning, given that the speaker has not explicitly indicated exhaustivity of the subject? The same sort of reasoning process that is inherent to GCIs is relevant here: the listener compares the target utterance with possible alternatives. Given that the speaker *only* mentioned Ashley, and no one else, in the assertion, the listener may infer that the predicate applies to no one else in the context. (Of course, since this is an implicature, it is entirely cancelable, as demonstrated by the fact that the speaker could continue with the follow-up, “In fact, she’s not the only one. What happened?”) If the context is stripped down so that there are only two children in the domain (Ashley and Donovan) and the speaker utters (2), the listener can be much more confident that the speaker intends to indicate the *only* Ashley is crying.

What’s more, although such PCIs are not tied to specific lexical items, reasoning about such implicatures may be deployed at a local level. For example, once the listener has evidence that the speaker has finished making reference to the grammatical subject—i.e., when the speaker goes on to pronounce the auxiliary verb *is*—s/he may deduce that only that individual is being predicated of. (Again, of course, this deduction may or may not prove correct, and is also dependent upon other knowledge, such as the range of other constructions the speaker uses, elements of the discourse context, and so forth.)

A challenge for the young language learner, then, is to determine what a speaker’s intended meaning is in a *particular* discourse context, recruiting their syntactic, semantic, and pragmatic skills in tandem. In addition, as a child becomes more adult-like, s/he should become more rapid in making this determination as the utterance unfolds. While the field of language acquisition

has witnessed a growth in studies investigating when and how rapidly children calculate GCIs, very little is known about their calculation of PCIs, and how this ability is deployed to pick out a speaker's intended reference in real time.

Our goals in this study were thus twofold. First, we sought to investigate how children reason about a speaker's event descriptions in a discourse context, and what this reasoning process—which depends on the calculation of a PCI—reveals about their pragmatic expectations about the intended referent of event descriptions. Second, we sought to gather evidence about how they deploy these expectations as the description unfolds. Here, we focus on the intransitive frame with a singular subject as in (2), as compared with a conjoined plural subject (e.g., *Ashley and Donovan*).

2. Previous Research on Particularized Conversational Implicatures

The type of PCI described above has recently been investigated experimentally. Here, we summarize the findings of three main studies demonstrating that adults and young children alike are able to calculate an “and nothing else” implicature in a given context. However, the findings leave open some important questions, which we highlight at the end of this section.

Breheny, Ferguson, & Katsos (2013) were interested in how rapidly adults computed a PCI as they listened to an utterance unfold in real time. In their paradigm, a speaker watched as an actor placed objects into boxes, and then delivered a statement about her actions. Crucially, there were two conditions, which varied according to whether or not the speaker witnessed the entirety of the actions. In one such scenario, the woman placed a spoon into box A, and then a spoon into box B. In one condition, this was all that the speaker saw, but in another condition, the speaker then saw her then place a fork into box A. In both conditions, the speaker delivered the utterance in (3).

(3) The woman put a spoon into box B and a spoon and a fork into box A.

In the second condition, where the speaker saw the *entire* scene, participants were able to anticipate the correct referent (box B) soon after the onset of the preposition. Since this preposition marked the end of the direct object argument of *put*, participants were able to infer that the speaker indicated that the woman put a spoon [and nothing else] into the box. Since box A had a spoon and a fork in it, and box B only had a spoon in it, box B was the logical choice.

Calculation of the same sort of “and nothing else” implicature has also been observed in very young children. Papafragou and Tantalou (2004) ran a task in which animals were given jobs to perform (e.g., eat a sandwich). Greek-speaking children (age 4-6) were told to reward the animal if he had done his job. For each trial, the animal was asked if he had done what he was asked to do (e.g., “Did you eat the sandwich?”), and he responded with an underinformative statement (e.g., “I ate the cheese.”). Since being a sandwich does not necessarily entail having cheese, children had to infer from the puppet's statement that he

did not eat the entire sandwich, but only the cheese [and nothing else]. In response, children generally refused to give the animal a prize, reporting that he had not done his job. The percentage of correct responses here differs quite remarkably from those reported in the child language literature on entailment-based scalar implicatures.

Perhaps just as remarkable is the fact that the same sort of reasoning process appears to be deployed by even younger children, as young as 3.5 years of age, in an even more minimalist task. Stiller, Goodman, & Frank (2015) presented English-speaking children age 2 to 4 and adults with a forced choice among a set of images differing only in the number of features displayed (e.g., a smiley face, a smiley face with glasses (one feature), and a smiley face with glasses and a hat (two features)). Participants were then given the prompt in (4) and asked to select the intended referent.

(4) My friend has glasses.

By 3.5 years of age, children were more likely to select the image with one feature (the glasses) than either of the other images and more than chance, demonstrating that they interpreted (4) as picking out the smiley face that had glasses [and nothing else]. Since it was true that both the image with one feature and the one with two features had glasses, it was thus pragmatic reasoning – and more specifically a PCI – that led participants to make this selection.

Although these studies provide with evidence that participants of all ages are able to rely upon relevant contextual information to calculate particularized conversational implicatures, two open questions remain. First, how much does this ability generalize beyond object DPs, and second, how rapidly can children (and adults) engage in this kind of reasoning as the target sentences unfold. Our study was intended to address these two questions.

3. Method

3.1. Participants

Native English-speaking adults ($N = 55$) and children ($N = 52$) (2;9 to 5;6) participated. Participants were randomly assigned to one of two between-subject conditions, based on the nature of the subject, as in (1) (Singular or Plural). The singular intransitive frame is our target, while the conjoined subject plural intransitive serves as a control. The mean age of children in the Singular condition was 4;6, and in the Plural condition 3;9. Participants were recruited from the Central NJ, and Boston, MA, areas and were tested either in the lab or in a quiet room at their preschool. Data were excluded from an additional 3 children due to a side bias, 2 for pointing difficulty, 1 due to a developmental disability, 1 for technical error, and 2 for experimenter error (delivering the wrong form of the utterances). Only those child participants whose parents gave consent for the eye gaze to be analyzed were included in the final analysis of eye gaze; however, all children were included in the pointing analysis.

3.2. Materials

Because we wanted to focus on children's pragmatic reasoning, while taking for granted their real world knowledge to ensure proper semantic computation, we avoided proper names, and used definite descriptions involving familiar animals, as in (5a)-(5b).

- (5) a. The pig is bending.
- b. The pig and the duck are bending

Each of these utterances was paired with the same two adjacent scenes: one in which there was only a pig bending, and one in which a pig and a duck were both bending. (The events in both scenes corresponded to the target verb.) Note that (5a) is true in each scene described here, as captured by the semantic formalism in (6): (6b) asymmetrically entails (6a). Thus, a speaker who utters (5a) and is taken to be maximally informative can be understood as intending reference to a scene in which only (5a)/(6a) is true, and correspondingly *not* intend reference to a scene in which (6b)/(6b) is true.

- (6) a. $\exists x.(pig)(x) \wedge (bending)(x)$
- b. $\exists x\exists y.(pig)(x) \wedge (bending)(x) \wedge (duck)(y) \wedge (bending)(y)$

The intransitive frame serves as a fitting subject of investigation, given its role in a series of verb learning studies over the years (e.g., Arunachalam, Syrett, & Chen, in press; Gertner & Fisher 2012; Naigles 1990; Naigles & Kako 1993; Noble, Rowland, & Pine 2011; Pozzan, Gleitman, & Trueswell 2015; Yuan, Snedeker, & Fisher 2012). Given the choice between a scene in which two agents are coordinating their actions (e.g., bending) and a causative scene in which an agent acts on a patient (e.g., the agent bending the patient), accompanied by a conjoined subject intransitive (e.g., Mary and Suzie are lorp), children well into their third or fourth year are at chance in choosing between the scenes when asked to find the novel verb referent. Even earlier, at 19-21 months, children hearing a singular intransitive frame are at chance choosing between a causative scene and a non-causative scene in which a singular agent performs an action, with or without a bystander (though they reliably map transitive frames to the causative scene).

However, we might expect just such a response pattern from children who are guided by their semantics and/or have no motivation for engaging in pragmatic reasoning that would adjudicate between the two scenes. On the one hand, the 19-21-month olds might just be too young to choose between the scenes on the basis of the language presented in that task. On the other, the older children are typically given a choice between two scenes, both involving two actors, along with a conjoined-subject intransitive frame. If this frame is compatible with multiple interpretations (see discussion in Arunachalam et al. in

press), then chance performance is not unexpected. Moreover, since no studies have yet given children the choice between a singular and a plural event referent, presented with either the singular or conjoined-subject intransitive, it is not known whether they can compute an “and nothing else” PCI. Collecting evidence that they *can* do so provides us with more insight not only into their developing pragmatic reasoning in general, but more specifically into their expectations about how a speaker intends reference to events.

Our visual stimuli consisted of video clips of live actors who wore animal costumes (always a duck and a pig) and performed simple actions. These events were filmed using a Sony digital camera and edited in iMovie. The subevents performed by the characters were carefully time-locked. We filmed the events with both characters, which served as the 2-participant scene, and then digitally cropped out one of the characters to serve as the 1-participant scene. The visual events were then paired with a pre-recorded voiceover, as indicated in the descriptions of the Familiarization and Test phases below.

Auditory stimuli were recorded in a sound-attenuated booth, and edited in Praat for intensity and length. The grammatical subject lead-ins (e.g., the singular subject and first conjunct *The [pig/duck]...*) were the same sound files re-used throughout the stimuli, spliced together with the same token of the second conjunct (*and the [duck/pig]...*). These sound files of the singular or plural subjects were then spliced together with the verb phrases in Praat, to ensure that the prosodic features and length were identical throughout the stimuli. Auditory and visual stimuli were then assembled in Final Cut Pro.¹

There were seven experimental trials, presented in one of two counterbalanced orders, with seven different corresponding actions. These included frequent/familiar and infrequent verbs (*bend, bonk*, tickle, pat, swat*, poke, wash**).² Each trial consisted of the same tripartite structure: Familiarization Phase, Test Phase, and Pointing Elicitation. Participants viewed the videos on a laptop with a built-in webcam, which recorded their faces.

Familiarization Phase. Each trial began with an introduction to the two characters. Each appeared sequentially in the center of the screen with captivating audio. (See Table 1). The target action was then depicted by each character, one at a time, one on each side of the screen. The accompanying audio directed participants’ attention to the characters and their actions. This phase thus served to familiarize participants with the characters and the dynamic scenes, but did not introduce the target utterances. The order in which the characters were introduced and the choice of which one appeared in the 1-participant scene at test was counterbalanced across trials.

¹ KS recorded and edited the visual and auditory stimuli. SA assembled the visual and auditory files.

² Those marked by * featured instruments (e.g., a toy racquet or scrubber).

			
(dynamic event)	(dynamic event)	(dynamic event)	(dynamic event)
<i>Look at the pig!</i> <i>What a nice pig!</i>	<i>Look at the duck!</i> <i>What a nice duck!</i>	<i>Wow!</i> <i>Watch!</i>	<i>Hey, look!</i> <i>Watch!</i>

Table 1. Sample Familiarization Phase for ‘bending’ event

Test phase. Following the Familiarization Phase, participants proceeded to the Test Phase. (See Table 2.) First, the two dynamic test scenes appeared simultaneously, with audio. This display reinforced the dynamic nature of the events. The scenes then froze on a still frame, with the target utterance: “Look! The pig is bending!” in the Singular condition, and “Look! The pig and the duck are bending!” in the Plural condition. These still scenes allowed us to collect participants’ eye gaze patterns as they heard the target sentence unfold, while ensuring that their gaze was not affected by the dynamic properties of the scenes. The dynamic scenes then replayed to remind participants of the actions, and the audio repeated. The scenes froze again and the target utterance repeated.

1-participant scene	2-participant scene
	
1. (dynamic event)	} <i>repeated</i>
<i>Look! Did you see that?</i>	
2. (still scene)	
<i>[The pig is / The pig and the duck are] bending</i>	
>elicited pointing<	

Table 2. Sample Test Phase for ‘bending’ event

Pointing elicitation. Immediately upon completion of the test phase, the experimenter paused the video and elicited pointing by asking, “Can you point? Can you show me [target utterance]?” We made sure to deliver the target utterance in the declarative intransitive frame, and not ask, e.g., “Where do you see the duck bending?” or “Where is the duck bending?” so that the target utterance was always in the same syntactic frame. The experimenter recorded the direction of pointing/scene selection on a coding sheet. The pointing

responses served as an offline measure, which we anticipated would serve as a reflection of participants' interpretation of the speaker's intended referent.

3.3. Coding and Analysis

Gaze direction during each presentation of the target sentence was coded offline, frame-by-frame, by trained coders naïve to study hypotheses. Coders noted whether participants were looking to the *left* or *right* of the screen or elsewhere, as well as whether gaze was uncodable due to blinks or obstruction. Saccades were coded as beginning one frame before the shift was visible; blinks and occlusion were coded as beginning on the first frame on which the pupils were not visible. Each participant was coded by at least two coders. Intercoder agreement was high. For *pointing*, a 1 was assigned to each trial on which a participant pointed to the 2-participant scene (the correct response in the Plural condition) and a 0 for each trial on which they pointed to the 1-participant scene (the correct response in the Singular condition *if* participants calculate the PCI).

3.4. Predictions

We made the following predictions. First, participants who did not engage in pragmatic reasoning about PCIs should perform at chance in the Singular condition, since the proposition expressed by the speaker's utterance is true in both contexts. We did not expect this pattern for adults, but based on previous research we thought it might be evident in children younger than 3.5 years of age. We further predicted that both adults' and children's performance would be near ceiling in the Plural condition, since only the scene in which both agents have the property expressed by the VP allows that sentence to be both true and felicitous. (The definite description *the pig* in (5b) renders that utterance infelicitous in a scene in which there is only a duck bending, since there is a failure of the presupposition of existence (Strawson 1950).)

Second, we predicted that participants who did engage in this reasoning would do so incrementally, as the target sentence unfolded. In the Singular condition, the listener has sufficient information at the auxiliary *is* to know that the subject is indeed singular.³ (Recall that because we spliced together the two conjuncts, no prosodic cues were present that could signal this before the auxiliary.) Thus, if listeners use this cue as it unfolds, they might shift attention to the 1-participant scene as early as 200 ms after the auxiliary (200 ms being the minimum time required to program and launch a saccade, Hallett 1986), though of course calculation of the implicature could take considerably longer.

³ Naturally, in principle one could continue the sentences differently (e.g., "The pig is bending and so is the duck."). We suspected such a possibility would not be readily entertained, especially as the trials proceeded, and the participants were only exposed to one target syntactic frame at test.

4. Results

We first present the pointing results, followed by the eye gaze results. Within each, we begin with the results from the adults, which serve as the backdrop against which we compare the results from the children.

4.1. Pointing Results

As anticipated, adults showed a clear split in their selection of the ‘2-participant scene’ between the experimental conditions, as evident in Fig. 1. In the Plural condition, with the exception of one adult on one trial, the 2-participant scene was always chosen. In the Singular condition, we see a consistent preference for the 1-participant scene that increases over the course of the experimental session, with preference for the 2-participant scene at 26% on the first trial and just 8% on the final trial. Not surprisingly, a logistic regression model (binomial family) on the Singular condition data with subject and trial as random effects yielded a significant intercept parameter estimate (-7.96 , $z = -4.22$, $p < 0.001$), indicating that performance differed significantly from chance.

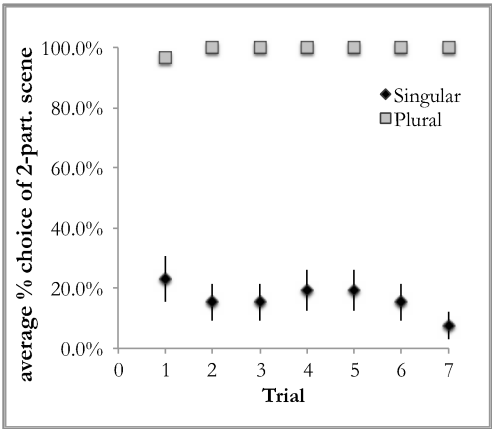


Fig 1. Adults' pointing responses to the 2-participant scene by condition

Children showed a similar pattern, selecting the 2-participant scene on average 97% of the time in the Plural condition. In the Singular condition, mean preference for the 2-participant scene was, like adults, below chance overall, 35%, and a logistic regression model also yielded a significant intercept parameter (-0.97 , $z = -2.04$, $p < 0.05$). However, as is evident from Fig. 2, their performance improved over the course of the experimental session. On the first trial, they selected the 2-participant scene 57% of the time, but by the final trial, they did so just 35% of the time—a trend we take to indicate a gradual increase in pragmatic reasoning over the course of the task. Children's preference for the 2-participant scene did not correlate with age ($R^2 = 0.0057$).

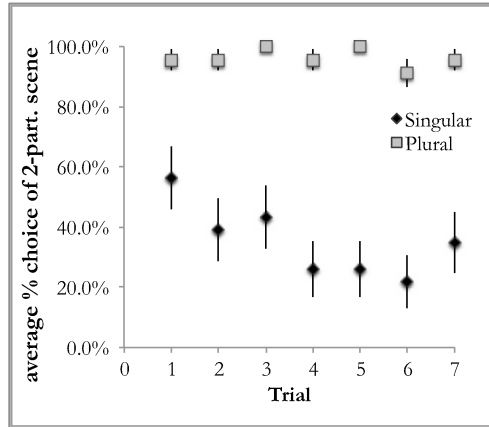


Fig 2. Children’s pointing responses to the 2-participant scene by condition

4.2. Eye Gaze Results

Here, we report on eye gaze from the Singular condition only, as this condition allows us to track pragmatic reasoning online, as the sentence unfolds. Because there were two presentations of the test sentence per trial, and seven trials, we present multiple analyses for each age group to demonstrate patterns within and across trials. A word of caution: we present our results as preliminary and suggestive only, since the fact that not all participants contributed gaze data means that we lack statistical power.

Adults appeared to show a relatively steady preference for the 2-participant scene, both times they heard each target sentence. In Fig. 3, the mean proportion of looks directed to the 2-participant scene at each frame is plotted, collapsing across trials. The vertical line indicates the onset of the auxiliary “is”.

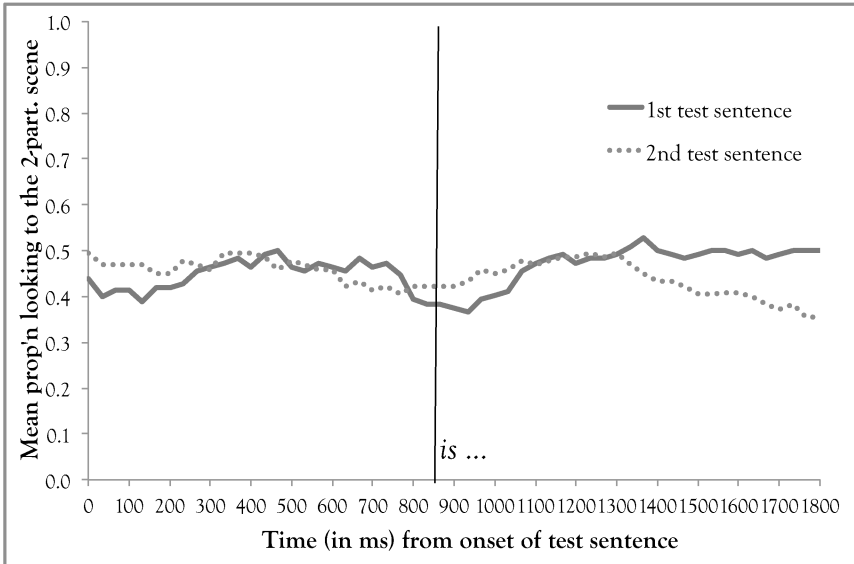
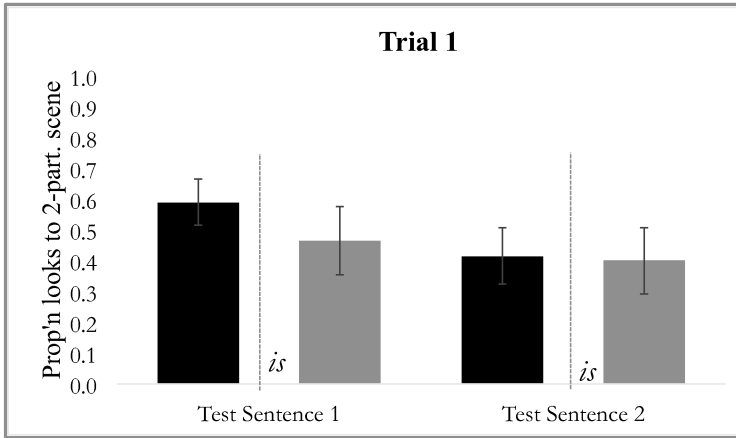


Fig 3. Adults' gaze to the 2-participant scene in the Singular condition over time, by target sentence presentation (first vs. second)

While these results collapse over trials, we can restrict our attention to the first and last trials to observe what happens over the course of the experiment. On the first trial, adults do not yet know how the sentence will unfold, but by the seventh, they should be able to anticipate the structure they will hear. Figure 4 depicts gaze to the 2-participant scene on just those two trials, now aggregated into two time windows: 600 ms prior to, and 600 ms following the onset of “is”.

(a)



(b)

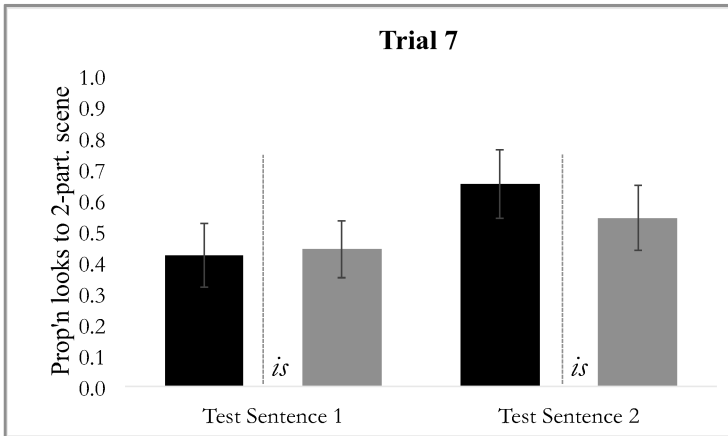


Fig 4. Adults' gaze to the 2-participant scene in the Singular condition over time, on Trial 1 (a) and Trial 7 (b) only

On the first trial, during the first presentation of the target sentence, adults begin with a preference for the 2-participant scene. This is not surprising; the 2-participant scene is more interesting, both because it has more characters and because of the synchronized action they had performed before the still frame. But as the sentence unfolds, preference for the 2-participant scene drops after the onset of *is* and remains low for the rest of the trial, as they hear the sentence again, suggesting that they shifted their gaze to the 1-participant scene on hearing the auxiliary; if this is the correct interpretation, it would indicate a rapid calculation of the implicature and a corresponding gaze shift, even on their very first experience with the trial structure. On the seventh and final trial, adults show nearly the opposite pattern. From the beginning of the target sentence, they

show a preference for the 1-participant scene, though this reverses by the second presentation. However, regardless of where adults are *looking* at the end of the second presentation of the test sentence, they *select* the 1-participant scene when queried with the singular intransitive frame immediately afterward.

We turn now to the children, whom we divided according to pointing performance: those who pointed to the 1-participant scene on four or more of the seven trials were labeled as “good” pointers. These children, who appeared to have calculated the PCI, can provide insight into how rapidly they did so. Their gaze is depicted in Fig. 5. Here, the first and second presentations of the target sentence show distinctly different profiles, with the first indicating a preference for the 2-participant scene, and the second indicating a preference for the 1-participant scene (which they ultimately pointed to). Importantly, gaze patterns for the two sentences diverge at the auxiliary *is*, which suggests that they are in response to the unfolding sentence. The “bad” pointers, not depicted here, showed relatively flat lines, hovering at 0.5 throughout both presentations.

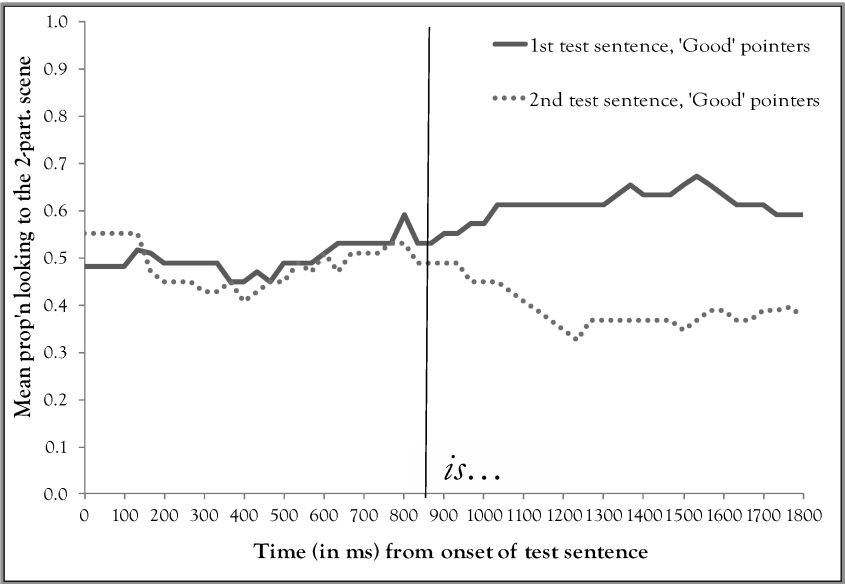
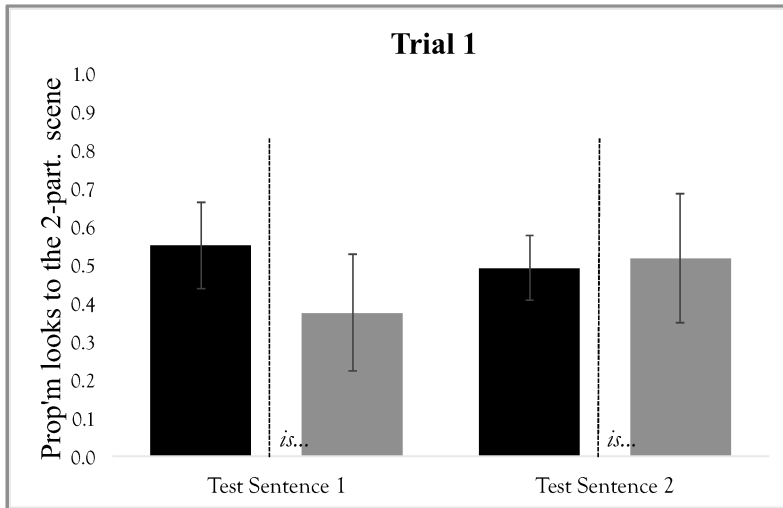


Fig 5. Children’s gaze (“good” pointers) to the 2-participant scene in the Singular condition over time, by target sentence presentation

The first and last trials (Fig. 6) show a strikingly similar pattern to the adults. On the first trial, children prefer the 2-participant scene and then shift attention to the 1-participant scene. The last trial shows the opposite pattern.

(a)



(b)

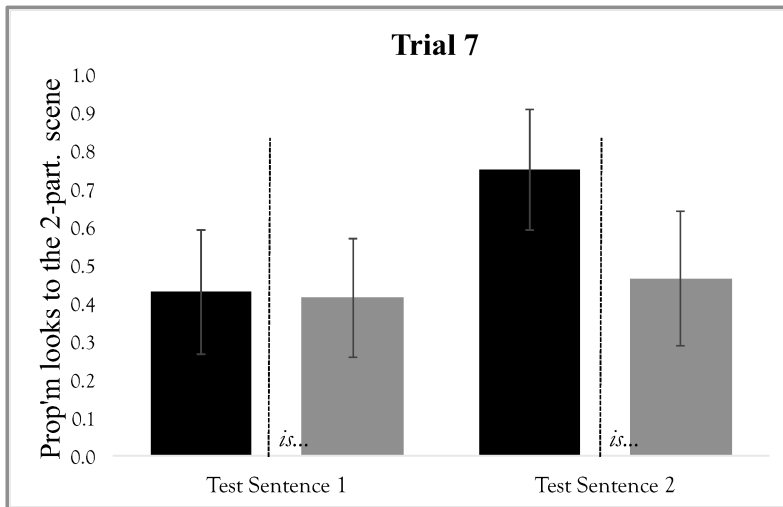


Fig 6. Children's gaze ("good" pointers) to the 2-participant scene in the Singular condition over time, on Trial 1 (a) and Trial 7 (b) only

5. General Discussion

The results of our study demonstrate that preschoolers well before age five, like adults, calculate the "and nothing else" PCI associated with a speaker's use of a singular subject in an intransitive frame, just as they do with singular direct objects (Papafragou & Tantalou, 2004; Stiller et al., 2015). The pointing data not

only demonstrate this point, but reveal an increasing level of pragmatic reasoning over time on the part of the children. Moreover, we also obtained preliminary evidence that this implicature is calculated in real-time, as the sentence unfolds. We analyzed participants' eye gaze at the 1- and 2-participant scenes as they heard the target sentences. Though our results can only be taken as suggestive at this point, it is notable that both children and adults appeared to shift their attention over the course of the sentences, and specifically just after the singular auxiliary *is*, which was their cue that the referential expression labeling the subject was complete. Those children who pointed correctly on more than half of the trials showed a similar pattern to adults on the first and last trials, and when all of their trials are taken together, their gaze appears to shift at the auxiliary. Further research will be needed to (a) see if this pattern is statistically robust, and (b) determine if the breakpoint we have visually identified in the figure is indeed where this shift occurs.

References

- Arunachalam, Sudha, Syrett, Kristen, & Chen, Yongxiang. (in press). Lexical disambiguation in verb learning: Evidence from the conjoined-subject intransitive frame in English and Mandarin Chinese. *Frontiers in Psychology: Language Sciences*.
- Breheny, Richard, Ferguson, Heather, & Katsos, Napoleon. (2013). Taking the epistemic step: Toward a model of on-line access to conversational implicatures. *Cognition*, 126, 423-440.
- Gertner, Yael, & Fisher, Cynthia. (2012). Predicted errors in children's early sentence comprehension. *Cognition*, 124, 85-94.
- Geurts, Bart. (2010). *Quantity implicature*. Cambridge: Cambridge University Press.
- Grice, H. Paul. (1975). Logic and conversation. In H. Paul Grice (Ed.), *Studies in the Ways of Words*. Cambridge: Harvard University Press. Reprinted in P. Cole & J. L. Morgan (Eds.) *Syntax and Semantics, Vol. 3: Speech Acts* (pp. 41-58). New York: Academic Press.
- Hallett, Peter E. (1986). Eye movements. In K.R. Boff, L. Kaufman, & J.P. Thomas (Eds.), *Handbook of perception and human performance* (pp. 10.1-10.112). New York: Wiley.
- Horn, Laurence R. (1984). Toward a new taxonomy for pragmatic inference. In D. Schiffrin (Ed.), *Form and use in context: Linguistic applications (Georgetown University Round Table)* (pp. 11-42). Washington, DC: Georgetown University Press.
- Horn, Laurence R. (2006) Implicature. In L. R. Horn and G. Ward (Eds.), *The Handbook of Pragmatics* (pp. 2-28). Oxford: Blackwell Publishing Ltd.
- Levinson, Stephen. (1983). *Pragmatics*. Cambridge, MA: MIT Press.
- Levinson, Stephen. (2000). *Presumptive meanings*. Cambridge, MA: MIT Press.
- Naigles, Letitia. (1990). Children use syntax to learn verb meanings. *Journal of Child Language*, 17, 357-374.
- Naigles, Letitia, & Kako, Edward. (1993). First contact in verb acquisition: Defining a role for syntax. *Child Development*, 64, 1665-1687.

- Noble, Claire H., Rowland, Caroline F., & Pine, Julian M. (2011). Comprehension of argument structure and semantic roles: Evidence from English-learning children and the forced-choice pointing paradigm. *Cognitive Science*, 35, 963-998.
- Papafragou, Anna, & Tantalou, Nicki. (2004). Children's computation of implicatures. *Language Acquisition*, 12, 71-82.
- Pozzan, Lucia, Gleitman, Lila R., Trueswell, John C. (2015). Semantic ambiguity and syntactic bootstrapping: The case of the conjoined-subject intransitive sentences. *Language Learning and Development*. Online first.
- Stiller, Alex, Goodman, Noah, & Frank, Michael C. (2015). Ad-hoc implicature in preschool children. *Language Learning and Development*, 11, 176-190.
- Strawson, Peter F. (1950). On referring. *Mind*, 59, 320-344.
- Yuan, Sylvia, Fisher, Cynthia, & Snedeker, Jesse. (2012). Counting the nouns: Simple structural cues to verb meaning. *Child development*, 83, 1382-1399.