
Children's Use of Intentionality Cues in Verb Learning

by Joshua Tabaldo
Psychology

Faculty Mentor: Dare Baldwin, Psychology
Graduate Student Mentor: Eric Olofson,
Ph.D. Candidate, Psychology

Abstract

Attempts to discern how children solve the problem of learning language have tended to focus on nouns, but few studies (Carpenter, Akhtar, & Tomasello, 1998) discuss the acquisition of verbs. While noun learning can be attributed largely to the whole-object and the mutual exclusivity assumptions (Markman, 1985), a possible tool for verb acquisition uses intentionality cues, for which one assumes that novel verbs refer to intentional actions. To test this hypothesis, we presented children aged 22-26 months with a dual display of paired action videos, one representing an accidental action and the other representing an intentional action. The action streams in each pair of videos were identical; they differed only in the object being manipulated and contextual information that distinguished which action was categorized as intentional and which accidental. By asking children to choose the video that portrays the novel verb, we can interpret the child's use of intentional cues. Children pointing to the intentional video would support our hypothesis; children pointing to the accidental videos would dispute it.

Introduction

For decades, psychologists and linguists have been trying to solve one of the premier riddles of childhood development: how do children learn language, or more specifically, how do children acquire novel labels? Quine (1960) developed the “Gavagai” example to illustrate this phenomenon. Imagine a young anthropologist exploring a new land with the help of a native guide. Exposed to this novel native language, the anthropologist has a hard time deciphering what his guide is saying. While the two wander down a dirt path, a rabbit darts in front of them. The guide points at the animal and exclaims, “Gavagai!” To what is the guide referring? Perhaps he is saying, “There it goes!” or even referring to his pet rabbit named “Gavagai.” He could be pointing out that the animal is white, fast, or has big ears. Perhaps these natives hold the rabbit as a sacred religious symbol that is viewed as the center of their universe. The point is that there are an infinite number of possible referents for the word. How can the anthropologist correctly determine the guide’s meaning? Quine defines this as the problem of induction. The task, however, is much more difficult for a language-learning infant than for our anthropologist. While the anthropologist has higher-level cognitive abilities and prior language experience, the infant has neither, which only amplifies the problem. Of course, children find ways to master language at a rapid rate and have access to a range of tools to help them achieve that mastery.

Ellen Markman (1989) proposed a few basic assumptions about mental tools that children have when making judgments about novel linguistic data. First, the *whole object assumption* asserts that when presented with a novel label for an object, a child will

assume that the label refers to the object as a whole. For example, when the guide pointed and said, "Gavagai!" one of the child's first inclinations would be to map that label to the entire rabbit, not just the ears or the manner in which it jumps. Second, the *taxonomic assumption* posits a cognitive constraint suggesting that children assume a word refers to a class of individuals rather than to a single animal or person. Going back to our Gavagai example, this bias shows that we attach the label not to specific rabbits, but to rabbits in general. Markman's third assumption is that of *mutual exclusivity*, which states that children will not assign a new label to an object that already has one. When a novel label is presented to refer to a previously labeled object, the new label could refer either to part of the object or to an adjacent object. For example, if one were to present a 3-year-old with two objects, one familiar (a banana) and one novel (a lemon wedge press) and present the child with a novel label, the child would assign the word to the object that does not yet have a label (Markman & Wachtel, 1988).

The three assumptions also can work in conjunction with one another. Perhaps we already know that "Gavagai" refers to the entire rabbit, because of the whole object assumption, but we are given a new label. Using mutual exclusivity, we assign the word to refer to part of the rabbit, such as the floppy ears. These assumptions are great for understanding how children learn labels for static, tangible objects. Little is known, however, about how children successfully solve the problem of induction when learning novel verbs. The action verbs that children first learn refer to something that usually occurs momentarily. Although Markman's theory of assumptions has been investigated only with regard to children's acquisition of nouns, it is possible that children similarly have assumptions

about the referents of novel verbs, an idea that is the core of this research.

When trying to understand how children acquire verbs, it is important to know what else they are capable of doing because one skill might aid in the formation and mastery of another. Children have an extraordinary ability to understand a person's goals and intentions (Meltzoff, 1995; Sommerville & Woodward, 2005; Woodward, 1998). That ability, as a study by Carpenter, Akhtar, and Tomasello (1998) has shown, is apparent in children as young as 14 months of age. That study also showed just how action and intention are connected. In the study, 14- to 18-month-olds watched an experimenter interact with an apparatus with two features that could be manipulated. The experimenter would engage both items on the apparatus, but the action was designed so that either manipulation could be deemed intentional or accidental. The only evidence the children had to infer intention was the experimenter saying either "Woops!" or "There!"—making only one of the words refer to one of the manipulations and the other word refer to the remaining action. Of course, the different scenarios were counterbalanced. Saying "Woops!" was supposed to connote an accident, and saying "There!" meant an intentional action. After the ambiguous action engaged both items on the apparatus and the two terms were verbalized, children were asked to make the apparatus work. Children preferred to imitate only the action that was identified with the label "There!" suggesting they prefer to imitate intentional actions. They might have known that the accidental manipulation of one of the items (marked by the "Woops!") would not activate the apparatus, and thus, did not attempt to recreate that action. This suggests that infants have some

knowledge of dealing with action in terms of an actor's intentions and goals. Infants appear to show the ability to use language to help classify action as either intentional or accidental. Therefore, the data suggest that children might be able to use verbs, a piece of language that strictly refers to action, to make similar judgments.

A different study by Tomasello and Barton (1994) connected intentions and action with the addition of verbs to display children's use of intention in verb acquisition. In their study, Tomasello and Barton presented 2-year-old children with two novel verbs (either "twang" or "hoist"), with each verb being associated with a toy. Two structures were used for this experiment, a crane and a merry-go-round, with each structure being capable of both actions (twang or hoist). In the intentional-first condition, the experimenter said, "Let's (either twang or hoist) Big Bird. Let's (either twang or hoist) him." They performed the target action and then the accidental action in the same general action sequence. The only information the children had in order to base judgments of intentionality was the experimenter saying, "There!" right after producing the target action and either "Woops!" or "Uh-oh!" directly after the accident. As one kind of control, the researchers simply had the accident precede the target action. The assignment of the target words (either plunk or whirl), the order of conditions, and the order of the toys (either Big Bird or Ernie) were counterbalanced across all conditions. Two comprehension tests had the child reproduce the two target actions in the different tests, one in each test. Only the child's first response was recorded. The results show that no child in any condition reproduced the accidental action but would equally reproduce the intentional action whether it came first or last. This study shows that the two-year-olds could rely on the

social-pragmatic information about the experimenter's intentional state to learn new verbs.

These two studies suggest that children assume that novel verbs refer to intentional actions. The Carpenter et al. study (1998) suggests that infants readily construe action in intentional terms, and this information is important in their own engagement with the world. The Tomasello and Barton (1994) study suggests that infants prefer to map verbs onto intentional actions. It is unclear, however, from the Tomasello and Barton (1994) article whether infants assumed that novel verbs refer to intentional acts, or whether intentional acts appear more salient, and children simply attached the verb to the salient action. By replicating this research with a new methodology, we can hope to gain a better understanding of whether and how children use intentionality cues in verb learning. If replication of the research asserts this claim, further questions can probe the basis of the intentional bias. Would it be harder to learn accidental verbs? If so, then we could test this by perhaps teaching children novel verbs for new actions (some intentional and some accidental) and then do a vocabulary test on those words later. Another question could explore whether children have a means to resist mapping novel verbs to accidents? If so, we could devise ways of finding these means of resistance. These are some of the questions we intend to investigate once preliminary research is conducted. And all of these questions further the fundamental query of how children acquire language, but more specifically, verbs.

Thus, the current research is part one of a two-part study. This first part is intended to replicate the previous Tomasello and Barton (1994) study with a new methodology. The next part of the study

will explore the nature of this difference (if one is found) and whether there is a resistance towards mapping novel verbs onto accidental actions.

Method

Participants

Eight children were the participants for the study. The age of the participants ranged from 23 months and 29 days to 25 months and 26 days ($M = 24$ months and 23 days, $SD = 22$ days). Five children were female and three were male. Three participants were not included in data analysis because they did not point to any of the trials (2) or equipment failure (1). The participants were recruited by the telephone, using information from birth announcements in a local newspaper. The children (and adults) were compensated for their cooperation and time with either a small T-shirt or a gift certificate to a local toy store.

Apparatus

The child was placed in front of two separate television sets with a DVD player embedded in each television. The TVs were separated by about a foot on two identical wooden stools. The experimenter used a remote control for the television sets to control the starting of the next trials during the experiment. Each child sat in a small chair so that the viewing level was the same as the level of the television sets, but right in front of a parent, who sat in another chair. The participant sat about 2 feet in front of and directly between the two televisions. The experimenter sat on the floor between the two televisions in order to (1) control the starting and stopping of the videos and (2) engage the child for responses. Two video cameras recorded each session. One camera, placed

behind the televisions, recorded the responses of the children; the second camera, placed in the back corner of the room, recorded the stimuli on the televisions. A curtain hanging behind the televisions occluded everything, but had a circular opening for the camera lens (See Figure 1, 2 and, 3).

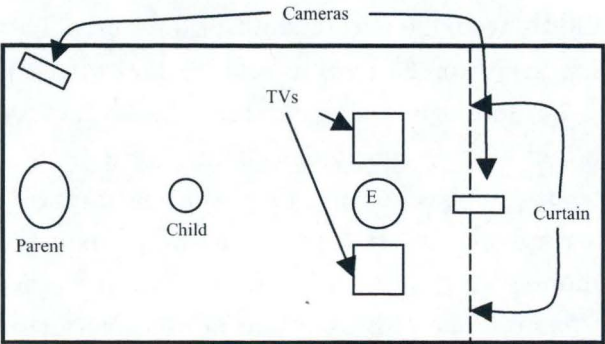


Figure 1: Diagram of experimental setup.



Figure 2: Experimental setup showing the back of the room.

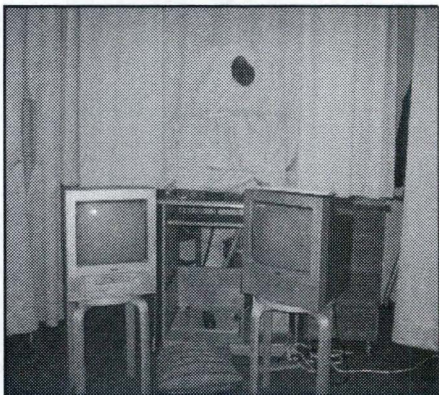


Figure 3: Experimental set up showing the front of the room.

Procedure

Participants came to our laboratory for a single half-hour visit. The current study included four different trials (each a different action). When the child participants entered the room with their guardians, we let them become acquainted with the room and the experimenter while the guardians provided informed consent. Once each child was ready to begin, the experimenter gave the guardian the MacArthur-Bates Communicative Development Inventory [CDI], a parental checklist of vocabulary used and/or known by the child, along with a supplemental vocabulary sheet to fill out during the remainder of the experiment. Meanwhile, the child sat in front of the monitors after the experimenter asked: "Would you like to watch some movies?"

Warm-up trial

In order to get the children used to pointing to the televisions, once the child sat down, the experimenter simultaneously placed two small figures, one a cat and the other a dog of similar size and color, on top of each of the televisions. The experimenter then asked the child to point to the television with the dog on top of it. Once the child produces the correct response, the experimenter asked the child to point to the television with the cat on top of it. If no responses are made, the question can be repeated and perhaps rephrased to elicit a response from the child.

Study

Once the child correctly pointed to the televisions in the warm-up trial, the experimenter took the small animal figures down and then started one of the DVDs. Each video started with a black screen. The experimenter would say, "Watch, she's going to blink it! She's going to blink it!" and then push play on one television

(counterbalanced either left or right). A short action video would then begin that depicted a woman performing an action such as clinking a teacup or ripping a sheet of paper. Starting in a neutral state, the video would zoom in on the action being completed so that the actress' face could not be seen. After the action was completed, she would say an utterance to assign intentionality. The actress would say either "Oops," "Okay," "Uh-oh," or "There" for the *clink*, *break*, *dump*, and *rip* videos with "Oops" and "Uh-oh" being accidental utterances and "Okay" and "There" representing more intentional sayings. The video looped three times with a second of a black screen between each loop before ending on a still frame of the end state of the action. Then, while the still frame was showing, the experimenter pressed play on the other DVD player to start its video of the same action, but the video was different in two ways. First, the object being manipulated to complete the action was different from the object depicted in the first video. Second, while the utterance of the actress was either the intentional or accidental saying in each video, the exact utterance was different from the one used in the first video loop, (See Table 1.) After playing the second

Table 1-Different versions of actions

Action	With object 1	With object 2
Dump	Legos in bucket	Fish crackers in fish container
Rip	Magazine	Fabric
Clink	A bottle and glass	Cereal bowl and spoon
Break	Salamander toy	Lego tower

video three times before freezing on its end state, the experimenter asked the child, "In one of these, she's blicking. Can you point to

Table 2-Number of Points by Order and Intentional Placement

		Order of Point		Total
		First	Second	
Intentional	First	4	10	14
Placement	Second	2	12	14
Total		6	22	

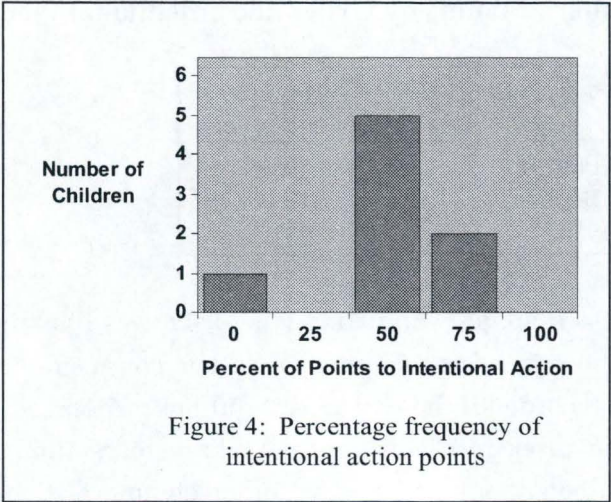
where she was blinking? Show me where she was blinking.” The participant’s points were recorded as either correct or incorrect; nothing was recorded if the participant did not respond.

The same procedure was repeated three more times for the remaining three actions. Of course, other obvious variables such as the order of the intentional sound, which video was played on the right or left, and what was said by the actress are all counterbalanced.

Results

Analysis of the data explored various aspects of the study, but did not include analysis by action or by verb pair (oops/okay or there/uh oh) because the study was not fully counterbalanced. To do a full analysis by action, the entire counterbalancing scheme would be needed.

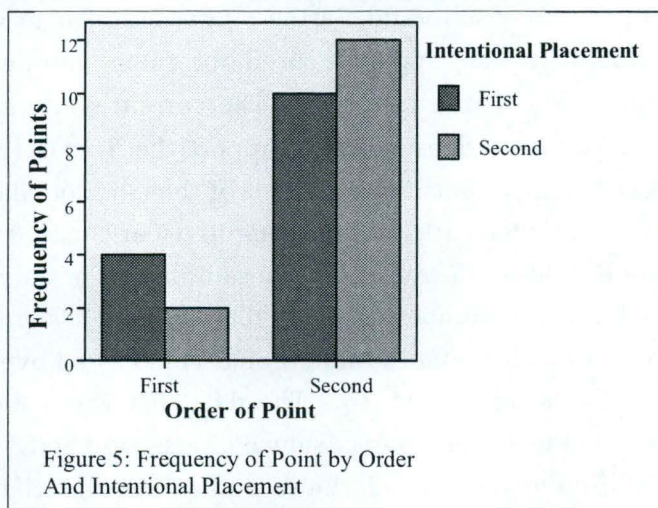
A between-subjects, one-sample t-test analysis of the percentage of correct points ($M = 50\%$, $SD = 23.15$) shows that there was not a significant number of points to the intentional video. The data thus far show that children are pointing to the intentional video at a level expected by chance (50%), $t(7) = 0.00$, *ns*. This result, however, does not accurately describe the nature of the data because one subject with a score of 0% correct (two non-responses and two incorrect responses) drove down the sample mean. In Figure 4, we



see that the majority of children pointed to the intentional action at about chance levels or slightly above. At this point in the study, the participant who had 0% correct points is being included in the dataset because the counterbalancing scheme is not yet complete. Once the entire data has been collected, the participant's inclusion in the data can be further decided.

A prediction interaction looked into order of point and intentional placement. It was found that whether participants pointed to the first or second action did not differ on the intentional placement of the action (whether the intentional video came first or second), $\chi^2(1) = 0.848$, $p > 0.3$ (See Figure 5). Taking a closer look at Figure 5, we see the beginning of an interaction occurring, though nothing close to significance. We would expect to see a lot of points to the first action when the intentional video is first and a lot of points to the second action when the intentional video is second. Unfortunately, children tended to point to the second

action regardless of whether it was intentional. The participants, however, had a slight preference to point to the first action when the



intentional video was shown first. That slight preference provided a hint that our research may be on the right track, especially when we examine the planned binomials of marginal frequency, which showed that children had a bias to point to the second action regardless of other factors ($p < 0.005$). The results revealed no side bias; that is, no preference to point at one side over another.

Discussion

This experiment examined one of the key problems in language development: how do children acquire language when there are an infinite number of possible referents for each object (Quine, 1960)? More specifically, the study focused on the under-researched enigma of the acquisition of verbs and how children's knowledge about intentions varies their learning judgments. Based on research

by Carpenter, Akhtar, and Tomasello (1998) and also in Tomasello and Barton (1994), the study meant not only to replicate the principles in the Tomasello and Barton study, but also to utilize a new methodology that might be used for future studies and to look at the problem in a new way. The current study and its methodology, however, does not yet support the Tomasello and Barton (1994) findings. Because only half of the counterbalancing order was filled, it is too early to draw definitive conclusions about the verb-intention bias. Currently, with a sample size of only eight participants (half the number needed) there is not a significant number of children choosing the intentional video as a novel verb referent over the accidental video. The data that was collected, however, seemed to suggest the existence of a second action bias, by which children tended to pick the second action of a trial more frequently than the first action. With more participants, results may begin to improve.

Even once the counterbalancing scheme is filled, there are a few limitations to this study that merit further exploration. One possible way to improve the methodology of the current study would be to use live actors instead of videotaped actions. That change might make the actions more naturalistic and more salient to the children.

Perhaps a better way to structure the study would be to include a control trial of some sort to identify and possibly eliminate any confounding variables. One way to do this would be to ask the children to point to the video without giving them any novel verbs. Basically, we would ask the children to point to the video that they like most. Another possible control would omit any verbal contextual information ("oops," "there," etc.) during the videos, but

still ask the children to point while giving them the novel verbs.

Another possible improvement would be to include more participants of an older age group. Although the Tomasello and Barton (1994) study had an age range similar to the current study, the subjects had a different task. Therefore, the experiment might be worth trying with an older child population to see if there is an intentional bias in verb learning. It is possible that for some children the task was too simple and boring. In that case, children may have chosen not to respond or not to pay enough attention to point correctly. Either way, a developmental trend might occur so that some ages would excel in the task while other ages would not.

The current study is basically one of replication with the added goal of exploring whether the intentional bias, if present is related to a resistance toward an accidental interpretation over the intentional action. Further, what passes for an intentional bias may be due just as well to a higher salience or more acceptance for an intentional interpretation. This intentional bias may, with more investigation, be shown to be a type of assumption used in verb acquisition, similar to the Markman (1989) assumptions in noun learning. This research may open up further venues for other possible assumptions in word learning, dealing with intentions, emotional states, and motion streams. Perhaps, like the interaction that occurs between different noun assumptions, verbs also have assumptions that work with one another to assist children in their journey toward language acquisition. Taking that idea another step, it would also be economical to have the different assumptions for nouns and verbs work in conjunction with one another to aid in the acquisition of other words.

With the completion of this study, researchers come a step

closer to a full understanding of how children find the correct label when each thing has an infinite number of possible referents; or how children solve the problem of induction. Although we offer no solution at present, we hope to have more evidence of children utilizing intentionality cues for their language and cognitive development once the entire counterbalancing scheme is complete. For the time being, there is a slight second-action bias with no significant preference to point to an intentional video over an accidental one. Knowing that there is a preference towards a more intentional mapping of verb referents will lead to more research possibilities while exploring the foundations and extent of this apparent bias in children's preferences. Looking at how this possible verb learning assumption emerges and interacts with other biases will be a worthwhile endeavor.

References

- Carpenter, M., Akhtar, N., & Tomasello, M. (1998). Fourteen- through 18-month-old infants differentially imitate intentional and accidental actions. *Infant Behavior and Development, 21*, 315-330.
- Markman, E. M. (1989). *Categorization and naming in children*. Cambridge, MA: MIT Press.
- Markman, E. M., & Wachtel, G. E. (1988). Children's use of mutual exclusivity to constrain the meanings of words. *Cognitive Psychology, 20*, 121-157.
- Meltzoff, A. (1995). Understanding the intentions of others: Re-enactment of intended act by 18-month-old children. *Developmental Psychology, 31*, 838-850.
- Quine, W. V. O. (1960). *Word and object*. Cambridge, MA: MIT Press.
- Sommerville, J. A., & Woodward, A. L. (2005). Pulling out the intentional structure of action: The relation between action processing and action production in infancy. *Cognition, 95*, 1-30.
- Tomasello, M. & Barton, M. (1994). Learning words in nonostensive contexts. *Developmental Psychology, 30*(5), 639-650.
- Woodward, A. L. (1998). Infants selectively encode the goal object of an actor's reach. *Cognition, 69*, 1-34.