

Mother–Child Conversations About Pictures and Objects: Referring to Categories and Individuals

Susan A. Gelman and Robert J. Chesnick
University of Michigan

Sandra R. Waxman
Northwestern University

The distinction between individuals (e.g., Rin-Tin-Tin) and categories (e.g., dogs) is fundamental in human thought. Two studies examined factors that influence when 2- to 3-year-old children and adults focus on individuals versus categories. Mother–child dyads were presented with pictures and toys (e.g., a picture of a boat or a toy boat). Conversations were coded for references to generic categories (“Dogs are furry”), ostensive labels (“This is a dog”), or specific individuals (“Lassie”). Overall, pictures generated more talk about categories; objects generated more talk about individuals. However, when objects could not be manipulated, speakers expressed relatively more category references. These results suggest that representations (in the form of pictures or objects-on-display) encourage young children and parents alike to think about categories.

The distinction between individuals and categories is fundamental to logic, philosophy, and a range of psychological and linguistic issues (e.g., Biernat, 1991; Gelman, Coley, Rosengren, Hartman, & Pappas, 1998; Jackendoff, 1983; Leslie & Kaldy, 2001; Macnamara, 1986; Needham, Dueker, & Lockhead, 2005). These issues include social identity (do you think about your new boss as an individual or as a woman?), reasoning (when learning a new fact about Fido, when do you extend it to other animals?), and conceptual development (when do children focus on categories versus individuals?). At core, the issue is whether we construe an individual (e.g., Fido, or your new boss) as an individual or as a category member.

Within developmental psychology, the psychological and linguistic issues have both received attention. One key question has been when and how children come to consider an item as both an individual and a member of its kind (Needham & Bailargeon, 2000; Xu, Carey, & Quint, 2004). Another key question has been when and how children become sensitive to the linguistic conventions that distin-

guish individuals from kinds. There is healthy debate surrounding each of these questions. Yet there are also several points of consensus. For example, by the time children are 2 years of age, they successfully represent both individuals and kinds (Hall, Lee, & Bélanger, 2001), and use language effectively (e.g., proper nouns, count nouns, generics) to convey these distinct representations.

The goal of the present research is to discover the extent to which differences in how an item is presented affect how readily people think about individuals versus kinds. Specifically, we ask whether objects and pictures differently contribute to children’s and adults’ tendency to focus on kinds. This question is important for at least three reasons. First, kind-based (“generic”) information is central to children’s developing knowledge base (Prasada, 2000); therefore, it would be useful to understand the factors that support it (a previously unexamined issue). Second, the study can yield a deeper understanding of what objects and pictures mean to young children, offering further insights into their understanding of symbols (Sigel, 1999). Prior studies of children’s interpretation of nonverbal symbols (symbolic understanding) focused primarily on the link between symbols and individual objects. This is the first study we know of that directly examines the link between a symbol and a more abstract category. Third, by examining the relation between individual/kind representations on the one hand and picture/object representations on the other, we may identify connections between two aspects of development that were previously considered separate.

This research was supported by NICHD grant HD-36043 to Gelman and NICHD grant HD-28730 to Waxman. Study 1 was conducted as an honor’s thesis by Chesnick at the University of Michigan. We are very grateful to the parents and children who participated in the studies. We also thank Felicia Kleinberg, Brook McCloud, Bethany Gorka, Julie Carpenter, and Rob Palazzo for transcribing and coding the videotapes, and Sarah Glauser for creating the drawings that were used in the studies. We also appreciate Tracy Lavin’s comments on an earlier draft.

Correspondence concerning this article should be addressed to Susan Gelman, Department of Psychology, University of Michigan, 530 Church St., Ann Arbor, MI 48109-1043. Electronic mail may be sent to gelman@umich.edu.

To begin, we outline the view that there is a conceptual continuum running from individuals to kinds, and that points along this continuum are systematically reflected in language. In other words, the linguistic form employed by speakers can be seen as a consequence—or marker—of their construal.

Conceptual Continuum and Linguistic Conventions

There are several steps along a continuum of reference from wholly individual-centered to wholly kind-centered, and different types of noun phrases mark points along this continuum. At one end of the continuum, proper nouns (“This is *Fido*”) highlight the uniqueness of an individual (Hall et al., 2001). At the other end are generic noun phrases, which refer to kinds without reference to any distinct individual (“*Dogs bark*”; “*A dog is a mammal*”; Carlson & Pelletier, 1995). Generic noun phrases are expressed in English with multiple formal devices, including bare plurals (e.g., “*Bats live in caves*”), indefinite singulars (e.g., “*A wok is how people in China cook*”), and definite singulars (e.g., “*The elephant is found in Africa and Asia*”).

Occupying the middle ground between clearly individuating expressions like proper nouns on the one hand and clearly categorical expressions like generic nouns on the other, are a variety of noun-phrase types. For example, singular pronouns (e.g., *this . . . she . . .*) refer to distinct individuals and make no explicit reference to a category (e.g., “*This one*”), but may make implicit reference to a category (e.g., “*he*” refers to a male). Moving further along the continuum, nongeneric uses of count nouns may refer to an individual (e.g., “*I love my dog*”) or a group of individuals (e.g., “*The dogs are hungry*”). Other nongeneric uses of count nouns serve to place an individual in a category (e.g., “*That is a dog*”). We refer to the latter as ostensive labeling (see Goldin-Meadow, Gelman, & Mylander’s [2005] “nongeneric categoricals”). What is common to these nongeneric noun phrases is that they signal, either implicitly or explicitly, membership in a kind, and do so by means of the count noun.

Noun phrases provide a critical source of information to children as they develop concepts along this continuum. From a logical analysis, it is impossible to refer unambiguously to a kind in the absence of language. No process of enumerating or displaying examples can convey that birds (as a kind) have hollow bones. Yet all human languages render this an uncomplicated affair. Generic noun phrases (“*Birds have hollow bones*”) refer unambiguously to

kinds. Moreover, linguistic analyses and empirical evidence reveal several further patterns regarding generic use (Gelman, 2004). First, generics are typically invoked to refer to qualities of a kind that are relatively inherent, enduring, and timeless—not accidental, transient, or tied to context (Lyons, 1977). Second, these generic properties need not be essential or biological (e.g., “*Chairs have backs*”), but they are distinctive in highlighting properties that are broad in scope and central to the kind in question (Gelman, Star, & Flukes, 2002; Hollander, Gelman, & Star, 2002). Third, generics are also more commonly used for animals than artifacts, suggesting that they are recruited for those kinds that are more (vs. less) richly structured (Gelman et al., 1998). Therefore, the tendency to produce generics is guided by domain.

In short, there is a conceptual continuum from individuals to kinds, and this continuum is reflected in language by different types of noun phrases. Generics and ostensive labeling phrases reveal a focus on kinds, and proper names reveal a focus on individuals.

Factors That Contribute to Speakers’ Focus on Kinds

The broader purpose of the present set of studies was to identify more closely factors that influence our tendency to construe an item along this continuum from individual to kind. We focus on one factor in particular: the *representational status* of the item itself or the facility with which that item can stand in for (or represent) something else. For example, a drawing of a dog can represent (stand in for) a real dog. This work makes contact with an extensive body of work on representational status (DeLoache, 2004). DeLoache and her colleagues have demonstrated that between 2 and 3 years of age, there is rapid developmental change in children’s representational capacities (DeLoache, 1987, 2004). For example, when young 2-year-olds are shown a small model room, and are explicitly shown that it has roughly the same dimensions and features as a larger actual room, they fail rather dramatically to understand that the small model room can serve as a representation of the larger room. Thus, if an experimenter hides a small toy dog in the small model room, and states explicitly that the location of the small dog is a clue to the location of another larger toy dog that has been hidden in the larger room, young 2-year-olds fail to understand the clue. Moreover, DeLoache and her colleagues have demonstrated that pictures are more readily appreciated as representations than are objects (like the model

room). For example, if young 2-year-olds are shown a photograph of the large room (rather than a model) that depicts the location of the hidden toy dog, they often are able to use this as a clue to the location of the toy dog in the actual room (DeLoache, 1991). Indeed, even 18- and 24-month-old children recognize that a picture symbolically represents an object in the world (Preissler & Carey, 2004): when they learn a new word for a *picture* of a whisk, they immediately extend it to the real-world referent.

This line of work has focused thus far on individuals, asking whether and when children can construe one individual (a small model vs. a picture) as a representation of another (the full-sized room). Yet the same set of distinctions could be brought to bear in thinking about representations of kinds. Thus, we ask whether and how representational status (presenting an item as either an *object* or a *picture*) influences the way in which it is construed (as an individual or as a category member).

Preliminary evidence suggests that pictures and objects differ in the extent to which they are seen to represent kinds. For example, mothers in both the U.S. and China produced nearly 6 times more generics when reading picture books than when playing with toy objects (Gelman & Tardif, 1998). This difference is intriguing because it is consistent with the hypothesis that pictures are more likely than objects to invoke kinds. However, there are several other potential explanations. For example, not only were different items presented in the two contexts (picture-book reading and toy play), but also the books included a larger and more varied set of items than did the toys. Also, whereas the pictures in the book were presented sequentially (one page at a time), the toy objects were all available simultaneously and were therefore more readily incorporated into an ongoing event. A further limitation of the earlier study is that, because the children were so young (under 2 years of age), we had data only from the mothers.

The Present Studies: Rationale and Overview

The present studies were designed to examine more systematically whether representational status (pictures vs. objects) plays an instrumental role in the construals of both mothers and their young children. We asked whether pictures are more likely to be represented as kinds and objects are more likely to be represented as individuals. Our goal was to elicit mother–child conversations as they interacted with

either pictures or objects, and to use their language production as a means of gauging their construals. We chose this path for two reasons: (a) because language production is a sensitive index of young children's concepts, particularly for the distinction of kinds versus individuals, given the tight links between language and concepts noted earlier, and (b) because mothers are more successful than experimenters in eliciting conversation from young children. Two-year-olds can be quite forthcoming in conversations with their mothers (Bartsch & Wellman, 1995; Gelman, Taylor, & Nguyen, 2004), but are notoriously taciturn when talking with experimenters and other strangers.

We hypothesize that pictures will foster a focus on kinds and will therefore elicit a relatively high proportion of generic noun phrases and ostensive labeling phrases, whereas objects will foster a focus on individuals, and will therefore elicit a relatively high proportion of proper nouns and other individuating phrases, including those in which a speaker treats an object as a conversational partner (e.g., talking to or for an item).

To get at this issue, the very same items were presented either as objects or as pictures. Each mother–child pair saw 12 items presented as objects (e.g., a toy dog) and another 12 presented as pictures (e.g., a picture of a toy lion). We counterbalanced across participants the medium in which each item was presented (toy vs. picture), so that for any given item (e.g., the dog), half of the mother–child pairs were presented with a toy version (e.g., a toy dog) and half with a picture version of the same item (e.g., a drawing of the toy dog). The items were drawn equally from three domains (animals, food, and artifacts) to provide a more general test of the representational status hypothesis. Although prior research has found that generics are more frequently used for animals than artifacts (Gelman et al., 1998), we predicted that the contextual effects would hold across domains.

We predicted that mother–child dyads would focus more on kinds when they were interacting with pictures than with objects. We reasoned that the toy objects would be viewed as individuals in their own right, but that pictures of these same objects would be viewed as representations and that this would invoke discussion of the kind. We therefore expected to find more utterances concerning kinds (e.g., generic noun phrases and ostensive labeling) in talk about pictures, and more utterances about individuals (e.g., proper nouns and utterances directed to or from an item) in talk about objects.

Table 1
Items Used in Studies 1 and 2

Animals	Food	Artifacts
<i>Gorilla</i>	<i>Corn</i>	<i>Dresser</i>
<i>Horse</i>	<i>Hot dog</i>	<i>Hammer</i>
<i>Rat</i>	<i>Ice cream</i>	<i>Tambourine</i>
<i>Snake</i>	<i>Pear</i>	<i>Watch</i>
Crocodile	French fries	Boat
Dog	Lemon	Football
Frog	Pizza	Hat
Lion	Watermelon	Pot

Note. Items in italics appeared in Set A.

Study 1
Method

Participants

Fifteen mother–child pairs participated, with children ranging in age from 2 years 7 months to 2 years 11 months (mean age 2 years 10 months). Eight of the children were girls and 7 were boys. The sample was recruited from a midwestern university town, and was primarily White.

Items

Materials included 24 toy objects (8 animals, 8 artifacts, and 8 foods) and 24 realistic color drawings of those objects, created by an artist to look as similar as possible to the actual objects (same shape, parts, color, and details); see Table 1 and Figure 1. We used toys rather than actual objects, so that we could present a range of objects from a range of domains. Half the items were assigned to Set A and half were assigned to Set B. Each set included 4 animals,

4 artifacts, and 4 foods. For a given mother–child dyad, one set was presented as pictures and the other as objects. Across dyads, each set appeared as objects for 7 or 8 dyads, and as pictures for 7 or 8 dyads. As a result, then, each item appeared roughly equally often as a picture or as an object.

Procedure

Mother–child dyads were tested in our on-campus laboratory. Mothers were informed ahead of time that they would be videotaped and that we were interested in mother–child interactions; however, they were not specifically told that we were interested in language or differences between objects and pictures. At the beginning of the testing session, mother and child were seated on a couch, and mothers were encouraged to look at and talk about the pictures and objects, as they would normally do at home. Each dyad saw 12 items presented as objects and 12 items presented as drawings. Pictures and objects were presented in counterbalanced blocks. The amount of time spent with each item was self-paced. Sessions were videotaped. At the end of the session, the child received a small toy or book.

Because there was one picture on each page of a book, pictures were seen sequentially. We therefore sought to impose the same sequential attention to the objects. To this end, we stored the objects in a chest of drawers, with each object covered with a cloth that had on it a number (from 1 to 12). Mothers were instructed to take out each object one at a time in numerical order, and to return it before taking out the next. The cloths ensured that children could not see more than one object at a time. Occasionally (on fewer than 15% of all utterances) more than one object was taken out of a drawer at once; such sequences were noted. These trials did not appear



Figure 1. Sample object (photographed), drawing, and object encased in plastic box, used in the studies.

to differ from those in which only one object was present.

Transcribing and Coding

Each videotaped session was transcribed verbatim by one coder and then checked by two additional coders. Transcriptions were then coded to identify all utterances. Utterances were identified on the basis of intonational contour and timing; any continuous unit of conversation that was free of full stops or interruptions from the other speaker was identified as an utterance. As such, utterances could consist of sentences, phrases, or even single words if they were pronounced with final pitch (rising or falling intonation). Once the utterances were identified, a two-phase coding system was implemented.

Phase I of coding involved selecting all utterances in which the target item was visible to at least one member of the dyad (mother or child), and in which the utterance included a noun or pronoun referring to the item, a part of the item, or other item(s) or parts of the same type. These utterances, hereafter referred to as “on-task” utterances, were submitted to further analysis. Examples are provided in Appendix A. Utterances produced when the item was nonvisible occurred primarily when an object was temporarily covered by a cloth. Nonvisible items were excluded from further consideration so as not to inflate the number of object-related utterances that were included on the object trials while taking the objects out of the toy cabinet. This was deemed to be a conservative coding decision, as otherwise we might be in danger of inflating the number of object-trial utterances that could not be generic (e.g., utterances referring to an object whose identity is not yet known almost certainly will not be generic).

Phase II of coding involved categorizing all on-task utterances into one or more of the following types: generic phrases, ostensive labeling phrases, individuating phrases, or other. Examples are provided in Appendix A.

Generic phrases. Generics were utterances in which the noun phrase or pronoun referred to a category rather than an individual or set of individuals, and includes examples such as “Horses don’t bite” (maternal utterance), “What sound does a snake make?” (maternal utterance), “I don’t like [watermelon] seeds” (child utterance). Coding of generics is discussed in detail in prior publications (e.g., Gelman et al., 1998; Gelman et al., 2004; Pappas & Gelman, 1998) and Appendix B.

Ostensive labeling phrases. Ostensive labeling phrases were those in which a speaker provided or

requested a label but offered no further information (e.g., “What is that?” or “French fries”). This coding category was *not* used if the utterance included some other proposition (e.g., “Little hot dog” would not be considered ostensive labeling).

Individuating phrases. Individuating phrases were those in which a speaker singled out the individual item and did not invoke its membership in a kind. Examples included providing a proper name for the item, asking what proper name to give the item, or pretending that the item was itself a conversational partner (i.e., talking directly to the item or pretending to talk as that item).

Other. All other utterances were coded as “other” (e.g., “Little hot dog”; “Let me see what it looks like”).

Cohen’s kappas were calculated for each coding category, on at least 12 of the dyads (80%) and are as follows: .78 generics (98% agreement), .79 ostensive labeling (91% agreement), and .89 individuating (99% agreement).

Results

The procedure was successful in eliciting conversation within the dyads ($M_s = 67$ on-task utterances per child and 197 on-task utterances per mother). Overall, mothers produced more utterances than did their children (paired- $t(14) = 7.19$, $p < .001$), and objects elicited more utterances than did pictures ($M_s = 166$ and 98, respectively, paired- $t(14) = 3.89$, $p < .01$).

Our main question is whether representational status (objects vs. pictures) differentially elicited a focus on individuals versus kinds. We used the utterances produced by mothers and children as an index of their focus. More specifically, we tabulated the number of utterances that singled out individuals (i.e., individuating phrases) and those that referred to kinds (i.e., generic and ostensive labeling phrases). We then calculated the proportion of each type of utterance within each domain (animal, food, artifact). Therefore, for example, to derive the proportion of generic animal phrases, we tabulated the number of generic phrases produced on trials involving animals as the target item, and then divided this by the total number of on-task utterances produced for animal trials (including generic, ostensive labeling, individuating, and other). The data for each type of phrase were submitted to an analysis of variance (ANOVA), using speaker (child, mother) as a between-participants factor and representational status (picture, object) and domain (animal, food, artifact) as within-participants factors.

Table 2
Pictures Versus Objects, Mean % of On-Task Utterances That Were Generic Phrases

	Study 1—Pictures		Study 1—Objects		Study 2—Pictures		Study 2—Objects-in-Boxes	
	M	SD	M	SD	M	SD	M	SD
<i>Children</i>								
Animals	2.8	6.4	1.3	2.9	6.4	9.9	3.7	8.6
Foods	9.9	10.3	1.7	4.5	4.9	9.6	2.6	4.8
Artifacts	5.1	9.4	1.1	4.3	1.4	3.6	3.5	5.3
<i>Mothers</i>								
Animals	8.2	9.9	4.7	4.8	16.8	12.6	10.8	9.3
Foods	15.6	9.8	7.1	6.7	19.5	14.3	13.2	9.7
Artifacts	7.2	11.5	1.8	2.8	3.8	6.5	7.1	8.0
<i>Overall</i>								
Animals	5.5	8.6	3.0	4.3	11.6	12.4	7.3	9.5
Foods	12.7	10.3	4.4	6.2	12.2	14.1	7.9	9.3
Artifacts	6.2	10.4	1.4	3.6	2.6	5.3	5.3	7.0

Generic Phrases

The dependent measure in this analysis was the mean proportion of generic utterances produced (see Table 2). We conducted a 2 (speaker) $\times$ 2 (representational status) $\times$ 3 (domain) repeated measures ANOVA, with speaker as a between-participants factor and representational status and domain as within-participants factors. As predicted, a main effect of representational status, $F(1,28) = 15.36$, $p < .001$, $\eta^2 = .35$, revealed that participants were more likely to produce generics when they were interacting with pictures than with objects (8.13% and 2.97% of on-task utterances, respectively). (All eta-squared (η^2) results that we report use the partial η^2 formula [SSeffect/(SSeffect+SSerror)]. Tabachnick and Fidell [1989] suggest that partial η^2 is an appropriate alternate computation of η^2 .)

There was also a main effect of domain, $F(2,56) = 10.52$, $p < .001$, $\eta^2 = .27$, indicating that generics were higher for food than for either animals or artifacts (M s = 8.6%, 4.3%, and 3.8% of on-task utterances, respectively), $ps < .01$ using Bonferroni's adjustment; animals and artifacts did not differ significantly from one another. These two main effects were qualified by a Representational Status $\times$ Domain interaction, $F(2,56) = 4.47$, $p < .05$, $\eta^2 = .14$, indicating that the effect of representational status held up significantly in each domain (ps ranging from $< .05$ to $< .001$, one-tailed), but that it was more pronounced when target items represented food and artifacts than animals. In addition, there was a main effect of speaker, $F(1,28) = 4.99$, $p < .05$, $\eta^2 = .15$, indicating that mothers produced a higher proportion of generics than their children (7.43% and 3.66% of on-task utterances, respectively).

Ostensive Labeling Phrases

The dependent measure in this analysis was the mean proportion of ostensive labeling phrases produced (see Table 3). We conducted a 2 (speaker) $\times$ 2 (representational status) $\times$ 3 (domain) repeated measures ANOVA, as we did for the generic phrases, above. As predicted, there was a main effect of representational status, $F(1,28) = 32.27$, $p < .001$, $\eta^2 = .53$, with participants producing a higher proportion of ostensive labeling phrases when they were interacting with pictures than with objects (57% and 39%, respectively). There was also a main effect of domain, $F(2,56) = 17.62$, $p < .001$, $\eta^2 = .39$, indicating that ostensive labeling was significantly more frequent for artifacts and food than for animals, $ps < .02$, Bonferroni's; animals and food did not differ significantly from one another. These two effects were mediated by a Representational Status $\times$ Domain interaction, $F(2,56) = 4.67$, $p < .05$, $\eta^2 = .14$, indicating that although the effect of representational status held up in all three domains, ps ranging from $< .05$ to $< .001$ by Bonferroni's, it was greatest for artifacts and smallest for foods. In addition, there was a main effect of speaker, $F(1,28) = 27.31$, $p < .001$, $\eta^2 = .49$, indicating that children produced a higher proportion of ostensive labeling phrases than did their mothers (62% and 33%, respectively).

Individuating Phrases

Individuating phrases (including both proper names and episodes of talking to or for the item) were found exclusively in conversations involving items from the animal domain. We therefore collapsed the data over domain, summing the number

Table 3

Pictures Versus Objects, Mean % of On-Task Utterances That Were Ostensive Labeling Phrases

	Study 1—Pictures		Study 1—Objects		Study 2—Pictures		Study 2—Objects-in-Boxes	
	<i>M</i>	<i>SD</i>	<i>M</i>	<i>SD</i>	<i>M</i>	<i>SD</i>	<i>M</i>	<i>SD</i>
<i>Children</i>								
Animals	64.3	28.9	42.7	24.8	62.5	24.8	36.7	26.1
Foods	75.7	18.8	67.1	29.1	73.2	25.3	56.2	26.2
Artifacts	79.1	22.1	45.5	26.4	70.3	19.7	55.7	25.3
<i>Mothers</i>								
Animals	29.2	22.6	17.6	11.5	38.8	20.8	18.6	16.0
Foods	45.5	16.9	34.9	15.5	48.2	15.0	33.7	16.7
Artifacts	45.5	25.5	24.7	12.1	46.3	18.7	33.9	15.1
<i>Overall</i>								
Animals	46.8	31.2	30.1	22.9	50.6	25.6	27.7	23.3
Foods	60.6	23.4	51.0	28.1	60.7	24.1	45.0	24.5
Artifacts	62.3	29.0	35.1	22.8	58.3	22.5	44.8	23.3

of individuating phrases divided by the total number of on-task utterances, and conducted a 2 (speaker) $\times$ 2 (representational status) repeated measures ANOVA, with speaker as the between-participants factor and representational status as the within-participants factor. A marginal main effect for representational status, $F(1,28) = 3.90$, $p = .058$, $\eta^2 = .12$, was consistent with the hypothesis that speakers are more likely to construe objects than pictures as individuals: speakers produced a higher proportion of individuating phrases when interacting with objects than with pictures (2.94% vs. 1.53%, respectively).

Discussion

The results of this study are consistent with the proposal that the representational status of an item has consequences for our tendency to view it as an individual or as an index of its kind. Mothers and children alike produced a higher proportion of kind-relevant utterances (generic phrases and ostensive labeling phrases) when they interacted with pictures than with objects. Conversely, they offered a higher proportion of individualizing utterances (proper names, talking to/for) when they interacted with animal objects than with animal pictures. This suggests that mothers and their children share an intuition that objects fall closer to the individuating end of the continuum and pictures fall closer to the category end. These shared intuitions should serve them well in the natural course of communication, supporting their ability to focus on both kinds and individuals, and to direct the scope of inductive inferences accordingly.

The difference between speakers' intuitions regarding objects versus pictures is intriguing, but it is not entirely clear that it stems from the representational status of the items. The objects and pictures differed not only in their representational status, but also in that only the objects could be handled and manipulated. Perhaps this difference is responsible for the greater focus on individuals in the object than the picture context of Study 1. To examine this possibility, in Study 2 we uncouple these factors. We encase the objects from Study 1 in plexiglass boxes, thus preserving their representational status as objects but eliminating any possibility of direct manipulation. If the encased objects elicit the same patterns as unencased objects, this will constitute evidence that it is the representational status of the items, rather than the ability to manipulate them, that underlies their construal as individuals.

Study 2

The purpose of Study 2 was to eliminate direct manipulation of the objects and therefore examine directly the effect of differences in the representational status of pictures versus objects. If manipulation is central to the tendency to construe objects as individuals, then the object/picture difference should disappear when participants are prevented from manipulating them. In contrast, if representational status is central, then the object/picture difference should persist even when the objects are encased. As in Study 1, we predicted greater attention to kinds, and therefore greater use of generics and ostensive labeling in the Picture condition. Also as in Study 1, we predicted greater attention to individuals, and

therefore greater use of proper names and talking to/for the item in the Object condition.

Method

Participants

Twenty mother–child pairs participated, with children ranging in age from 2 years 7 months to 3 years 2 months (mean age 2 years 11 months). Ten of the children were girls and 10 were boys. The sample was recruited from a midwestern university town, and was primarily White.

Items

Study 2 used the same items as in Study 1, except that each object was now encased in a custom-made transparent plexiglass box. Boxes varied in size to accommodate the objects they held. An example appears in Figure 1. The boxes were sealed shut, so that participants were unable to touch or manipulate the objects themselves.

Procedure, Transcribing, and Coding

The procedure was identical to that of Study 1, except that mothers were further instructed that neither they nor their children should try to open the plexiglass boxes. Transcribing and coding also proceeded as in Study 1. Cohen's kappas were calculated for Phase I and Phase II separately, yielding .94 for Phase I (96% agreement, based on all 20 dyads) and .83 for Phase II (90% agreement, based on 17 dyads).

Results

Once again, the procedure elicited considerable conversation within the dyads ($M_s = 74$ and 154 utterances per child and mother, respectively). As in Study 1, mothers produced more utterances than did their children, paired- $t(19) = 8.60$, $p < .001$, and objects generated more utterances than did pictures, $M_s = 143$ and 85, respectively, paired- $t(19) = 4.36$, $p < .001$.

Our main question is whether objects that are sealed in plexiglass cases continue to elicit a focus on individuals (vs. kinds). As in Study 1, we used the utterances produced by mothers and children as an index of their focus, tabulating the number of utterances that singled out the individuals (i.e., individuating phrases) and those that referred to kinds (i.e., generic and ostensive labeling phrases). As in

Study 1, we then calculated the proportion of each type of phrase within each domain (animal, food, artifact), using the number of on-task utterances within each domain as the denominator, and submitted these scores to an ANOVA, using speaker (child, mother) as a between-participants factor and representational status (picture, object) and domain (animal, food, artifact) as within-participants factors. To examine our hypotheses about participants' construals, we conducted three separate ANOVAs: one each for the generic, ostensive labeling, and talking-to scores, respectively.

Generic Phrases

The dependent measure in this analysis was the mean proportion of generic utterances produced. We conducted a 2 (speaker) $\times$ 2 (representational status) $\times$ 3 (domain) repeated measures ANOVA, with speaker as the between-subjects factor, and representational status and domain as the within-subjects factors (see Table 2). Results revealed a main effect of speaker, $F(1, 38) = 26.37$, $p < .001$, $\eta^2 = .41$, indicating that mothers produced a higher proportion of generics than children ($M_s = 11.9\%$ and 3.78%, respectively). There was also a main effect of domain, $F(2, 76) = 11.25$, $p < .001$, $\eta^2 = .23$, indicating that generics were more frequent for animals and food than for artifacts ($M_s = 9.44\%$, 10.1%, and 3.98%, respectively). There was also a significant Domain $\times$ Speaker interaction, $F(2, 76) = 5.82$, $p < .01$, $\eta^2 = .13$, indicating that the domain effect reached significance only among the mothers—although the children's scores were also in the predicted direction.

Finally, most relevant for the purposes of the study, there was a Representational Status $\times$ Domain interaction, $F(2, 76) = 6.27$, $p < .01$, $\eta^2 = .14$, as well as a nonsignificant trend toward a main effect of representational status, $F(1, 38) = 2.98$, $p = .092$, $\eta^2 = .07$. Overall, speakers tended to produce a higher proportion of generics in response to pictures than objects ($M_s = 8.8\%$ vs. 6.8%, respectively). However, this effect differed as a function of domain. For both animals and foods, pictures elicited a higher proportion of generics than did objects, $p_s = .025$ and .053, respectively. In contrast, for artifacts, the effect was reversed: objects elicited a higher proportion of generics than did pictures, $p = .013$. This last result was unexpected, and we return to it in the Discussion section.

Ostensive Labeling Phrases

The dependent measure in this analysis was the mean proportion of ostensive labeling phrases pro-

duced. We conducted a 2 (speaker) $\times$ 2 (representational status) $\times$ 3 (domain) repeated measures ANOVA, with speaker as the between-subjects factor and representational status and domain as the within-subjects factors. The data are shown in Table 3. Results indicated a main effect of representational status, $F(1,38) = 40.37, p < .001, \eta^2 = .51$, indicating a higher proportion of ostensive labeling when talking about pictures than about objects ($M_s = 56\%$ and 39% , respectively), as predicted. In other words, pictures are more likely than objects to elicit references to kinds (in the form of ostensive labeling). This result supports the representational status hypothesis. There was also a main effect of speaker, $F(1,38) = 23.40, p < .001, \eta^2 = .38$, indicating a higher proportion of ostensive labeling among children than mothers ($M_s = 59\%$ and 37% , respectively). Finally, there was a main effect of domain, $F(2,76) = 19.21, p < .001, \eta^2 = .34$, indicating more ostensive labeling concerning artifacts and foods than animals ($M_s = 52\%, 53\%$, and 39% of on-task utterances, respectively).

Individuating Phrases

In contrast to Study 1, individuating the item did occasionally take place in all three domains. Therefore, we were able to conduct a 2 (speaker) $\times$ 2 (representational status) $\times$ 3 (domain) repeated measures ANOVA, with speaker as the between-subjects factor, and representational status and domain as the within-subjects factors. Results indicated a main effect of representational status, $F(1,38) = 5.61, p < .05, \eta^2 = .13$, a main effect of domain, $F(2,76) = 11.64, p < .001, \eta^2 = .23$, and a significant Representational Status $\times$ Domain interaction, $F(2,76) = 10.82, p < .001, \eta^2 = .22$. There were no significant effects involving speaker. As can be seen in Table 4, individuating phrases constituted a higher proportion of talk about animal objects than talk about animal pictures. This difference was significant only for the animal domain, $p < .001$, as we would expect, given that artifacts and food rarely are spoken to or given proper names. Thus, the Representational Status Hypothesis is supported for this measure, within the animal domain.

Comparison of Studies 1 and 2

In a further set of analyses, we compared the results of Studies 1 and 2 directly to determine whether there were significant differences in response. These analyses therefore provide a direct test of some comparisons that were only implicit in considering

Table 4

Study 2, Pictures Versus Objects-Encased-in-Boxes, Mean % of On-Task Utterances That Were Individuating Phrases

	Pictures		Objects-in-Boxes	
	M	SD	M	SD
<i>Children</i>				
Animals	1.73	5.7	5.57	7.8
Foods	1.43	6.4	1.15	5.2
Artifacts	1.07	4.8	1.29	4.6
<i>Mothers</i>				
Animals	0.16	0.7	5.19	5.6
Foods	0.50	2.2	0.91	4.1
Artifacts	0.00	0.0	0.24	1.1
<i>Overall</i>				
Animals	0.95	4.1	5.38	6.7
Foods	0.96	4.7	1.03	4.6
Artifacts	0.54	3.4	0.77	3.3

Studies 1 and 2 separately. Specifically, we predicted that talk about objects would differ by study, with Study 2 eliciting relatively more kind-focused talk (generic and ostensive labeling phrases) and relatively less individual-focused talk (proper names or talking to/for the object), but that talk about pictures would not differ by study.

To address this prediction, we conducted a study (2: Study 1, Study 2) $\times$ speaker (2: mother, child) $\times$ representational status (2: pictures, objects) $\times$ domain (3: animals, artifacts, food) ANOVA for each dependent variable (generic, ostensive, individuating phrases). There were significant effects involving study for both generic phrases (Domain $\times$ Study interaction, $F(2,132) = 3.68, p < .05, \eta^2 = .05$; Representational Status $\times$ Domain $\times$ Study interaction, $F(2,132) = 4.49, p < .05, \eta^2 = .06$) and ostensive labeling phrases (Representational Status $\times$ Domain $\times$ study, $F(2,132) = 4.59, p < .05, \eta^2 = .06$). However, there were no significant effects involving individuating phrases. Because our predictions vary by representational status (i.e., study effects are predicted for objects only, not for pictures), we present the results for objects and pictures separately, on the basis of follow-up Bonferroni's tests.

For *objects*, there were four significant study differences, all supporting the idea that encasing objects in boxes led speakers to construe them more as kinds. Encasing the objects in boxes elicited a significantly higher proportion of generics than when objects were not encased in boxes for all three domains ($p_s < .05, .05$, and $.01$ for animals, foods, and artifacts, respectively). Moreover, artifacts received a higher proportion of ostensive labeling when

encased in boxes (Study 2) than when they were not (Study 1), $p < .01$. In contrast, for *pictures*, there was only one significant study difference: animals received a higher proportion of generics in Study 2 than in Study 1 ($p < .05$). This difference was not predicted, but neither was it part of any consistent pattern of differences.

Examination of Children and Mothers Separately

In all of the analyses to this point, speaker (mother vs. child) was included as a factor in the analyses, and invariably the results held across this factor, indicating that the effects were consistent across both mothers and children. However, we also wish to test the effects within mothers and children separately. Combining Studies 1 and 2 gives us sufficient power to do so. Specifically, for both mothers and children, and separately for each coding category, we conducted a 2 (study) $\times$ 2 (representational status) $\times$ 3 (domain) ANOVA.

Mothers. Mothers showed representational status effects in the predicted direction for all three dependent variables. For generic phrases, there was a main effect of representational status, $F(1, 33) = 9.88$, $p < .01$, $\eta^2 = .23$, a main effect of domain, $F(2, 66) = 18.66$, $p < .001$, $\eta^2 = .36$, and a main effect of study, $F(1, 33) = 5.63$, $p < .05$, $\eta^2 = .15$. For ostensive labeling phrases, there was a main effect of representational status, $F(1, 33) = 40.41$, $p < .001$, $\eta^2 = .55$, and a main effect of domain, $F(2, 66) = 22.84$, $p < .001$, $\eta^2 = .41$. Finally, for individuating phrases, there was a main effect of representational status, $F(1, 33) = 5.27$, $p < .05$, $\eta^2 = .14$.

Children. Children also showed representational status effects in the predicted direction for all three dependent variables. For generic phrases, there was a main effect of representational status, $F(1, 33) = 6.75$, $p < .05$, $\eta^2 = .17$, indicating that overall children produced a greater proportion of generic phrases for pictures than for objects. There was also a Representational Status $\times$ Domain interaction, $F(2, 66) = 3.32$, $p < .05$, $\eta^2 = .09$, indicating that although the means were higher for pictures than objects in all three domains, the difference reached significance only within the food domain, $p < .005$, Bonferroni's. For ostensive labeling phrases, there was a main effect of representational status, $F(1, 33) = 34.65$, $p < .001$, $\eta^2 = .51$, indicating that children produced a greater proportion of ostensive labeling phrases for pictures than objects. There was also a main effect of domain, $F(2, 66) = 15.37$, $p < .001$, $\eta^2 = .32$, and a Representational Status $\times$ Domain $\times$ Study interaction, $F(2, 66) = 3.38$, $p < .05$, $\eta^2 = .09$,

indicating that the picture–object difference was significant for all three domains in both studies, ps ranging from $< .05$ to $< .001$, with the exception of food in Study 1. Finally, for individuating phrases, there was a main effect of representational status, $F(1, 33) = 4.57$, $p < .05$, $\eta^2 = .12$, indicating a greater proportion of individuating phrases for objects than pictures.

Contingency analyses. The analyses above indicate that children themselves show the same patterns as adults, with overall effects of representational status for generic phrases, ostensive labeling phrases, and individuating phrases. However, because the children were in conversation with their mothers, it is possible that they were simply imitating or mimicking that which the mothers produced. We therefore conducted another set of analyses to determine the extent to which children's utterances were spontaneously produced versus prompted by the mothers. Every child utterance that was coded as generic, ostensive, or individuating was further coded into one of three categories: *repeat* of a prior utterance from the mother, *response* to a maternal question that prompted the utterance, or *spontaneous* utterance. A second person coded the data from 25% of the participants, yielding agreement of 95%. Examples are provided in Appendix C.

Results indicate that children are robust contributors to these conversations, with roughly half of their coded utterances being spontaneous, neither prompted by maternal questions nor repeating a prior maternal utterance. This is perhaps not surprising with ostensive labeling (47% spontaneous), because prior research has shown that young children spontaneously use this form in their own speech (Brown, 1973). More surprisingly, roughly the same rates of spontaneous production are found with generic phrases (58% spontaneous) and individuating phrases (57% spontaneous). Responding to a maternal prompt question was also fairly common, accounting for 43% of children's ostensive labeling phrases, 41% of their generic phrases, and 33% of their individuating phrases. In contrast, simply repeating all or part of a maternal utterance accounted for very few child utterances: 10% of their ostensive phrases, 1% of generic phrases, and 10% of individuating phrases.

Finally, we compared the rate of spontaneous, response, or repeating utterances as a function of condition (object vs. picture) for each type of utterance (ostensive, generic, and individuating). To provide a thorough test, we conducted both parametric (t tests) and nonparametric (chi-square) analyses. In none of these comparisons were there

any significant condition effects. This result indicates that children are no more likely to repeat a maternal utterance in one condition versus another. Therefore, we can infer that the condition effects obtained earlier (e.g., more generics in the picture condition, more individuating phrases in the object condition) were not because of children mimicking their mothers.

Discussion

For the most part, the results of Study 2 replicate those of Study 1. Although the objects were encased in clear containers and therefore could not be touched or manipulated directly, several significant representational status differences were still obtained. We found a higher proportion of generic phrases and ostensive labeling phrases for pictures than objects and, conversely, a higher proportion of individuating phrases for objects than pictures. These results support the view that representational status is instrumental in our construals. Objects were more readily construed as individuals (as indicated by a higher proportion of individuating phrases [proper nouns and talking to/for the item]), whereas pictures were more readily construed as kinds (as indicated by a higher proportion of ostensive labeling phrases and, for two of the three domains, generic phrases).

At the same time, there were also significant study differences. Speakers generated a higher proportion of generics when the objects were encased in boxes, for all three of the domains. There was also a higher proportion of ostensive labeling of artifacts when they were encased in boxes. This suggests that making items less manipulable increases their status as representations. Interestingly, other researchers seem to share this intuition. For example, in a study of 2½- to 3-year-old children's interpretation of graphic symbols, Callaghan (2000) mounted materials (both pictures and objects) on foam board in order "to highlight their symbolic status" (p. 190).

The current findings are consistent with those of DeLoache (2000), who found that placing a model room behind a window also heightened children's ability to treat the model room as a representation of the larger one. With her task, children were more successful at using the model room as a basis for finding a hidden toy when it was behind glass. DeLoache interpreted this as evidence that the window decreased the salience of a model *as an object* itself, and accordingly increased its status as a representation. Furthermore, DeLoache found that increasing the salience of a model as an object (e.g., by

allowing children to play with it) decreased its representational status. Taken together with the present studies, these results suggest that factors that increase (or decrease) the tendency to construe an item as an object will decrease (or increase) its representational status.

Moreover, encasing the objects in boxes yielded one result that differed qualitatively from Study 1, namely, that artifacts received proportionately more generics in the object condition than in the picture condition. Why might this be? One possibility is that artifacts are especially likely (compared with other domains, including animals or foods) to be understood in terms of how people interact with them (e.g., function, use, or purpose may be especially central to artifacts; Kemler Nelson, Egan, & Holt, 2004; but see Sloman & Malt, 2003), so that when such interactions are disrupted, there is more dramatic change in how they are construed. Thus, artifacts that are put on display may be especially likely to be construed as symbols rather than manipulables (much as museum exhibits of old coins or ancient pottery indicate their significance as symbols, not as objects to be played with).

General Discussion

In two studies, mothers and their 2½-year-old children viewed and talked about items presented either as two-dimensional color drawings or as three-dimensional toy objects. We predicted that talk about objects would tend to focus on individuals, whereas talk about pictures would more often extend to a broader focus on categories (kinds). These predictions were indeed borne out: talk about pictures elicited a higher proportion of ostensive labeling phrases (e.g., "That's a watch") and generic phrases (e.g., "What do you make with lemons?"), and a lower proportion of individuating phrases (either proper names (e.g., "Mr. Frog") or talking to or for the item (e.g., "Oh, hi, Annabelle" [to horse]; "Hi, [child's name]" [as if frog is talking])). These patterns held up for both the mothers' talk and the children's talk.

What accounts for these differences between objects and pictures? It cannot be the content of the items, because the very same items were presented as both pictures and objects (e.g., the boat item was a red fire-boat with railings and writing on the side, both when presented as a toy object and when presented as a colored drawing). Consequently, the pictures and objects were as comparable as possible in both amount of detail and degree of prototypicality.

We have argued that it is the representational status that evokes these differences. However, one additional difference is that objects can be manipulated in ways that pictures cannot. We therefore conducted Study 2, in which the same pictures and objects were used as in Study 1, except that the objects were encased in translucent plexiglass boxes. With this modification, the objects could no longer be manipulated, making them functionally more comparable with the pictures. The results of Study 2 indicate that the object–picture differences were still significantly maintained: more talk about *kinds* (ostension and generics) for pictures than objects, more talk about *individuals* (proper names or talking-to) for objects than pictures. Therefore, being able to manipulate the objects cannot fully account for the effect. At the same time, placing objects in boxes did increase the tendency of children and their mothers to treat the objects like pictures (significantly more ostension and generics), showing that this factor did have some effect. This last result is consistent with that of DeLoache (2000), in which placing a model behind a window led to improved use of the model to find a hidden object (i.e., increased ability to treat the model as a representation).

Interestingly, this effect seemed to be strongest in the artifact domain, such that encasing artifacts especially altered their construal, increasing the tendency for them to be treated as kinds (with corresponding increases in both generic nouns and ostensive labeling). At this point we can only speculate as to why artifacts would be more susceptible to this effect than either animals or foods. One possibility is that artifacts may be understood in terms of functional interactions with people, so that disrupting such interactions may lead to more dramatic change in how they are construed.

The results of the current experiments support DeLoache's (1991) claim that pictures are more readily construed as representations than are objects, and that making an object less objectlike also increases how readily it is construed as a representation. When talking about objects, participants focused on the items as individuals: providing individuating names and taking the item as a conversational partner (talking to or for it). In contrast, when talking about pictures, participants were relatively less likely to individuate the item and more likely to treat it as representing a broader category (e.g., the focus was no longer on *this* creature, but on the fact that it is a member of the *dog* category [ostension], or even on *dogs* in general [generic]).

Although we interpret our data as supporting DeLoache's argument, there is also an important

difference. DeLoache found that pictures more readily than objects serve as representations of specific individuals (e.g., Snoopy). In contrast, we have found that pictures more readily than objects serve as representations of generic categories (e.g., dogs). Thus, both DeLoache's work and the present set of studies are alike in demonstrating that pictures are more easily interpreted as representations, but the type of representation differs in the two cases.

One complication of this work is that the objects we used were in fact representations themselves, because they were all toys. That is, the toys had a dual status as both concrete objects and abstract representations (DeLoache, 2000, p. 330). As with pictures, the toy pizza *represented* a slice of pizza, the toy boat represented a real boat, and so forth. So it is not the case that our stimuli presented a sharp contrast between representations (pictures) and non-representations (objects), but rather that we included items that varied along a representational continuum. Not only were the objects partly representational, but conversely pictures also had a non-representational aspect to them (e.g., a rectangle of laminated paper is a type of object unto itself). In other words, pictures are more on the representational end of a continuum, with toys on the more "real" end of that continuum. In fact, we suspect that there are additional points along this continuum (see Callaghan, 2000, for evidence with children 2½–3 years of age). A black-and-white simple line drawing, for example, may be a more abstract representation than a full-color detailed drawing, which may be a more abstract representation than a photograph. At times, even a fully functional object could be considered a representation if it is small and placed in an appropriate context (e.g., a doll's stroller, i.e., three fourths the size of an actual stroller). For current purposes, what is important is that these differences have measurable effects on how children and mothers view an item. However, in future research it would be interesting to explore various points along the continuum to determine children's sensitivity to these differences with the present task.

The object–picture differences have implications for how children reason in other tasks and domains. Liben (1999) discusses preschool children's "iconic realism" as when a 3-year-old reports that the picture of an ice cream cone will be cold (Beilin & Pearlman, 1991). The present data shed new light on this error. Our data suggest that children's difficulty may not reflect confusion about the status of pictures per se, but rather their difficulty inhibiting the generic knowledge that pictures call to mind. Thus, although a child appropriately understands that

pictures of ice cream cones are representations and therefore not to be eaten, their knowledge about the generic properties of ice cream cones is prominent in a picture context, and not easily suppressed. If this is the case, then (a) iconic realism errors should primarily concern generic properties (not idiosyncratic properties) and (b) such errors should be less frequent when children are presented with toy objects rather than pictures.

Another example of how these results may have broader implications comes from a study by Guthrie, Rapoport, and Wardle (2000) concerning food preferences in preschool children, comparing real foods, three-dimensional models of foods, and photographs of foods. Although judgments were most reliable when children were asked to judge real foods, photographs came in a close second—and food models produced unreliable ratings. The problem was not with identifying the foods in the models, as children did well on that. Rather, they seemed to have difficulty linking up the food in the models to their knowledge of the categories that these models represented. This finding seems consistent with our argument that photographs more readily link to generic or category knowledge.

A further possible implication of this work is that objects may be less ideal than pictures for teaching children generic knowledge or facts, because objects may interfere with reasoning about the broader category. This speculation receives some support in work by DeLoache, Uttal, and Pierroutsakos (1998) showing that children at times do less well with manipulable objects, when the goal is to use the manipulables as representing a more abstract concept (e.g., letters or numbers). Conversely, books and museumlike displays may be particularly effective in conveying generic information.

Finally, these data may contribute to debates concerning the level at which infants form categories. Whereas some researchers find categorization primarily limited to the global level (e.g., Mandler & McDonough, 2000), others find more extensive categorization at global, basic, and subordinate levels (e.g., Quinn, 2002). Among numerous methodological differences between the studies at issue, one key difference is that Mandler and McDonough used toy objects as stimuli, and Quinn and colleagues used photographs. Extrapolating from the current data, it may be that infants are more likely to categorize pictures than objects. Our findings may also have implications for the report that it is not until 14 months of age that infants generalize from a small toy model to a video of that model (Younger & Johnson, 2004). The present findings suggest that

infants may demonstrate an earlier understanding of this relationship if pictures were used rather than toy models.

One important point is that the representational status effect held up with both children and mothers examined separately. Of greatest interest in the current context is that children showed sensitivity at such an early age. By $2\frac{1}{2}$ –3 years of age, children do not yet fully appreciate that representations “stand in” for their referents (Liben, 1999), and are just starting to appreciate the dual nature of objects as both entities and representations (DeLoache, 2004)—yet they are also capable of appreciating the dual nature of pictures (as simultaneously representing individuals and kinds). One factor that may have helped children in this set of studies is that all the items were familiar and have known verbal labels (e.g., lion, hammer, ice cream). Callaghan (2000) found that, in children of this age group, the ability to appreciate the symbolic status of pictures was significantly higher when verbal labels were available.

It is also notable that the effect persisted into adulthood (for the mothers). Adults are certainly flexible and capable of construing an item (whether picture or object) as either an individual or representing a kind; nonetheless, adults are swayed by the same factors that influence preschool-aged children. One major question this finding raises (but does not address) is whether parental input influences children’s behavior in this realm (either children’s attention to kinds or children’s understanding of representational symbols). Prior research suggests intriguing links between parental talk about pictures and children’s grasp of symbols (Callaghan, Rochat, MacGillivray, & MacLellan, 2004; Szechter & Liben, 2004). It will be important in future work to try to tease apart the influences of parents on children’s development.

References

- Bartsch, K., & Wellman, H. M. (1995). *Children talk about the mind*. New York: Oxford.
- Beilin, J., & Pearlman, E. G. (1991). Children’s iconic realism: Object versus property realism. In H. W. Reese (Ed.), *Advances in child development and behavior* (Vol. 23, pp. 73–111). New York: Academic Press.
- Biernat, M. (1991). Gender stereotypes and the relationship between masculinity and femininity: A developmental analysis. *Journal of Personality and Social Psychology*, 61, 351–365.
- Brown, R. (1973). *A first language: The early stages*. Cambridge, MA: Harvard.

- Callaghan, T. C. (2000). Factors affecting children's graphic symbol use in the third year: Language, similarity, and iconicity. *Cognitive Development*, 15, 185–214.
- Callaghan, T. C., Rochat, P., MacGillivray, T., & MacLellan, C. (2004). Modeling referential actions in 6- to 18-month-old infants: A precursor to symbolic understanding. *Child Development*, 75, 1733–1744.
- Carlson, G. N., & Pelletier, F. J. (Eds.) (1995). *The generic book*. University of Chicago Press.
- DeLoache, J. S. (1987). Rapid change in the symbolic functioning of very young children. *Science*, 238, 1556–1557.
- DeLoache, J. S. (1991). Symbolic functioning in very young children: Understanding of pictures and models. *Child Development*, 62, 736–752.
- DeLoache, J. S. (2000). Dual representation and young children's use of scale models. *Child Development*, 71, 329–338.
- DeLoache, J. S. (2004). Becoming symbol-minded. *Trends in Cognitive Sciences*, 8, 66–70.
- DeLoache, J. S., Uttal, D. H., & Pierroutsakos, S. L. (1998). The development of early symbolization: Educational implications. *Learning and Instruction*, 8, 325–339.
- Gelman, S. A. (2004). Learning words for kinds: Generic noun phrases in acquisition. In D. G. Hall & S. R. Waxman (Eds.), *Weaving a lexicon*. Cambridge, MA: MIT Press.
- Gelman, S. A., Coley, J. D., Rosengren, K., Hartman, E., & Pappas, A. (1998). Beyond labeling: The role of maternal input in the acquisition of richly-structured categories. *Monographs of the Society for Research in Child Development*, 63(1).
- Gelman, S. A., Star, J., & Flukes, J. (2002). Children's use of generics in inductive inferences. *Journal of Cognition and Development*, 3, 179–199.
- Gelman, S. A., & Tardif, T. Z. (1998). Generic noun phrases in English and Mandarin: An examination of child-directed speech. *Cognition*, 66, 215–248.
- Gelman, S. A., Taylor, M. G., & Nguyen, S. (2004). Mother–child conversations about gender: Understanding the acquisition of essentialist beliefs. *Monographs of the Society for Research in Child Development*, 69(1).
- Goldin-Meadow, S., Gelman, S. A., & Mylander, C. (2005). Expressing generic concepts with and without a language model. *Cognition*, 96, 109–126.
- Guthrie, C., Rapoport, L., & Wardle, J. (2000). Young children's food preferences: A comparison of three modalities of food stimuli. *Appetite*, 35, 73–77.
- Hall, D. G., Lee, S. C., & Bélanger, J. (2001). Young children's use of syntactic cues to learn proper names and count nouns. *Developmental Psychology*, 37, 298–307.
- Hollander, M. A., Gelman, S. A., & Star, J. (2002). Children's interpretation of generic noun phrases. *Developmental Psychology*, 38, 883–894.
- Jackendoff, R. (1983). *Semantics and cognition*. Cambridge, MA: MIT Press.
- Kemler Nelson, D. C., Egan, L. C., & Holt, M. B. (2004). When children ask, "what is it?" what do they want to know about artifacts? *Psychological Science*, 15, 384–389.
- Leslie, A. M., & Kaldy, Z. (2001). Indexing individual objects in infant working memory. *Journal of Experimental Child Psychology*, 78, 61–74.
- Liben, L. S. (1999). Developing an understanding of external spatial representations. In I. E. Sigel (Ed.), *Development of mental representation: Theories and applications* (pp. 297–321). Mahwah, NJ: Erlbaum.
- Lyons, J. (1977). *Semantics* (Vol. 1). New York: Cambridge University Press.
- Macnamara, J. (1986). *A border dispute*. Cambridge, MA: MIT Press.
- Mandler, J. M., & McDonough, L. (2000). Advancing downward to the basic level. *Journal of Cognition and Development*, 1, 379–403.
- Needham, A., & Baillargeon, R. (2000). Infants' use of featural and experiential information in segregating and individuating objects: A reply to Xu, Carey and Welch (2000). *Cognition*, 74, 255–284.
- Needham, A., Dueker, G., & Lockhead, G. (2005). Infants' formation and use of categories to segregate objects. *Cognition*, 94, 215–240.
- Pappas, A., & Gelman, S. A. (1998). Generic noun phrases in mother–child conversations. *Journal of Child Language*, 25, 19–33.
- Prasada, S. (2000). Acquiring generic knowledge. *Trends in Cognitive Sciences*, 4, 66–72.
- Preissler, M. A., & Carey, S. (2004). Do both pictures and words function as symbols for 18- and 24-month-old children? *Journal of Cognition and Development*, 5, 185–212.
- Quinn, P. C. (2002). Category representation in young infants. *Current Directions in Psychological Science*, 11, 66–70.
- Sigel, I. E. (Ed.) (1999). *Development of mental representation: Theories and applications*. Mahwah, NJ: Erlbaum.
- Sloman, S. A., & Malt, B. C. (2003). Artifacts are not ascribed essences, nor are they treated as belonging to kinds. *Language and Cognitive Processes*, 18, 563–582.
- Szechter, L. E., & Liben, L. S. (2004). Parental guidance in preschoolers' understanding of spatial-graphic representations. *Child Development*, 75, 869–885.
- Tabachnick, B. G., & Fidell, L. S. (1989). *Using multivariate statistics*. Philadelphia: Harper & Row.
- Xu, F., Carey, S., & Quint, N. (2004). The emergence of kind-based object individuation in infancy. *Cognitive Psychology*, 49, 155–190.
- Younger, B. A., & Johnson, K. E. (2004). Infants' comprehension of toy replicas as symbols for real objects. *Cognitive Psychology*, 48, 207–242.

Appendix A. Examples of Coding

Phase I

Not visible	"Do you think it's going to be heavy or light?"; "Roll this one back up"; "Let's take out this one"; "Let's unwrap it to see what it is."
On task	"An alligator!" "Do you like lemons?" "Can you make him jump?" "I see his belly."
Other	"Rabbit, rabbit." "Why?" "Sure, OK." "One, two, three."

Phase II

Generic	"What do froggies say?" "Ice cream's for eating."
Ostensive labeling	"It's a hat." "What—do you know what this is?"
Individuating	"Oh, hi, Annabelle." [to horse]; "Bye, Mr. Frog!" [to frog]; "Does this doggie have a name, too?"
Other	"Now, can you tell me about that one?"; "He's a funny frog."

Appendix B. Guidelines for Identifying Generics

(Note: None of the examples in Appendix B were drawn from the present studies.)

A. Generics have two major properties:

(1) There is a general category the speaker refers to. The speaker is not referring to any particular individual or instance. Thus, generics typically do *not* have any of the following before the noun: (a) a number (e.g., "two birdies"), (b) a pronoun (e.g., "my marbles"), (c) the word "some" (e.g., "some more balloons"), and (d) the word "the" (e.g., "the doggies").

(2) The statement or question is not tied to a particular situation or point in time. This means that the statement or question is in present tense. It usually cannot be in the past, in the future, or in the progressive (-ing) form. (An exception to this point is that the historical past can be generic; e.g., "Dinosaurs were cold-blooded").

B. Examples of generics:

"Turtles are green"; "Boys don't like carrots"; "If you play with cords, dat's very dangerous"; "That looks like

for boys and girls"; "Shoes don't go on the table"; "Do airplanes have wheels?"; "I'm just afraid of animals"; "I like jelly beans."

C. Examples that are *not* generics:

"My turtle is green"; "Some boys didn't like carrots"; "If you play with this cord, dat's very dangerous"; "Your shoes don't go on the table"; "Does this airplane have wheels?"; "I'm looking at animals"; "I'm eating jelly beans"; "Yesterday we saw tigers at the zoo"; "I see a dog"; "This is a drum"; "There are mice in our house."

Appendix C. Examples of Child Utterances That Were Spontaneous, Responses to, and Repeats of Prior Maternal Utterances

Spontaneous:

- Mother: It's a lemon.
Child: Lemon.
Mother: Do you know when . . .
Child: I like lemons! [generic]
- Child: Oooh! A football!!!! [ostensive label]
- Mother: Oh, they are french fries, aren't they!
Child: Bye, Mr. French Fries. [individuating]

Response:

- Mother: Where do lions live?
Child: They live in the jungle! [generic]
- Mother: What is it?
Child: Watch. [ostensive label]
- Mother: Tell me about the puppy.
Mother: What do you think his name is?
Child: Pat. [individuating]

Repeat:

- Mother: No, you know what it is, it's corn.
Child: Oh. It's a corn. [ostensive label]
- Mother: His name is Douglas!
Child: Douglas! [individuating]