

How to Help Best: Infants' Changing Understanding of Multistep Actions Informs their Evaluations of Helping

Brandon M. Woo (bmwoo@g.harvard.edu) and Elizabeth S. Spelke (spelke@wjh.harvard.edu)

Department of Psychology, Harvard University, Cambridge, MA 02138

Center for Brains, Minds, and Machines, McGovern Institute for Brain Research, Cambridge, MA 02139

Abstract

Research beginning with Piaget reveals a change in infants' understanding of multistep, means-end action sequences: Whereas 12-month-old infants reason that (e.g.) one opens a box to access its contents, younger infants are more likely to reason that one's goal is simply to open the box. Here we explore the implications of this developmental change in infants' action understanding for infants' social evaluations. Using a puppet show paradigm, we examined infants' evaluations of two agents who helped another agent to achieve either the end or the means of a means-end sequence, both before and after 12 months of age. In a subsequent preference test, 15-month-old infants reached for an End-Helper over a Means-Helper, whereas 8-month-old infants did the reverse. These findings link infants' evaluation of helpers to their representations of action plans, consistent with recent computational models of naïve psychology.

Keywords: infant development; social cognition; action understanding

Introduction

In our social world, people face the challenge of determining who might help or harm them. Beginning with Hamlin, Wynn, and Bloom (2007), a number of studies have demonstrated that even infants engage in social evaluation of individuals who have helped others fulfill their goals or prevented others from doing so. Specifically, infants preferentially look at and reach for helpers over hinderers by 3 months and 4.5 months, respectively (Hamlin & Wynn, 2011; Hamlin et al., 2010; see Margoni & Surian, 2018, for meta-analysis). In order to determine who has helped or hindered a protagonist, however, one must first understand the protagonist's goal. Here we investigate how infants evaluate agents who help others fulfill their goals in a multistep, means-end action sequence. If one character tries to open a box to get to an object and the object then is moved to a different box, what do infants consider to be the more helpful action: opening the original box that the character previously sought to open, or opening the box that now contains the object?

Beginning with Piaget (1952), past work has demonstrated a developmental change in infants' capacities to engage in multistep, means-end actions that require understanding of second-order goals (e.g., understanding that the goal of opening a box is secondary to the goal of obtaining its contents). Until late in the first year, infants do not reliably execute two-step actions involving second-order goals, such as removing a barrier or pulling on a blanket to gain access to an object that lies out of sight or out of reach; see also, Bates

et al., 1980; Diamond, 1985; Sommerville & Woodward, 2005; Willatts, 1999). Such limits to infants' and children's action planning pervasively modulate their action capacities and understanding of the object-directed actions of others.

Insights into the role of action planning in infants' action understanding comes from a highly productive line of research that places infants in the role of third-party observers of other people's actions. Imagine a person repeatedly opening a box and grasping a truck inside. The truck in the box is switched with a different toy, a duck, and the truck is put into a second box as the person observes. After this switch, will the person be more likely to open the first box, where the truck used to be, or the second box, where the truck now resides?

When infants observe others engage in such means-end actions on objects, their understanding of those actions undergoes systematic developmental changes that track the development of their own action capacities. Just as infants do not reliably produce means-end actions until late in the first year of life, infants only reliably infer that another person performs a means-end action to achieve the end state (e.g., accessing the truck in the above scenario) at about 12 months of age. Infants below 12 months of age are more likely to infer that the goal state desired by the actor is the opening of the first box, where the truck used to be (Sommerville & Woodward, 2005; see also Woodward & Sommerville, 2000; Gergely et al., 2002).

Recent computational models can account for these findings. According to Baker et al. (2009), we understand other people's actions, and the mental states that underlie them, through a Bayesian process of inverse planning. From observing others' actions, we work backwards to infer the mental states that led them to enact those actions. Through computational models using an inverse planning approach, we can recover not only the goals of others' actions and behaviors (Ullman et al., 2010), but other mental states, including the value that people place on different goal objects (Jara-Ettinger et al., 2016), and their beliefs about the existence and locations of different objects (Baker et al., 2017). If Bayesian inverse planning models are correct, then younger infants may fail to understand the first-order goal of other agents in a two-step action for the same reason that they fail to engage in multistep actions to achieve their own goals: They are unable to generate the correct, hierarchically structured two-step action plan.

In prior research on infants' execution or understanding of multistep actions, all the actions have been directed at inanimate objects: infants either act themselves, or they

observe an agent who acts to change the state of the physical world (e.g., by opening a box to access an object). In the present experiments, we examine what happens when, instead, infants observe agents who help other agents who are attempting, but failing, to complete a multistep, means-end action sequence. In past work in which infants preferentially reached for helpers, the agents sometimes needed help achieving a first-order goal (e.g., climbing a hill to get to the top, as in Hamlin et al., 2007) and sometimes needed help achieving a second-order goal (e.g., opening a box to gain access to a toy inside, as in Hamlin & Wynn, 2011). In the latter case, however, infants might have preferred the helper (who joined an agent in opening the box) without understanding the agent's primary goal of obtaining the toy. To our knowledge, our experiments are the first to investigate whether infants' understanding of multistep actions informs their social evaluations.

From a computational perspective, understanding multistep, means-end actions could be harder, easier, or equally difficult in situations involving helping (as in the study of Hamlin & Wynn, 2011), relative to situations in which an infant observes an agent acting to obtain an object and predicts what the agent will do after the object is moved to a new box (as in the studies of Woodward & Sommerville, 2000). Action understanding could be harder because representations of helping require the coordination of the actions of two distinct agents with different goals that are hierarchically organized: The helper's goal is to aid another person's efforts to achieve his goal. If the latter person's goal is to obtain an object by opening a box, therefore, the helper accordingly pursues a third order goal to help the agent achieve its second-order goal of rendering accessible the object that it seeks to attain. Hamlin et al. (2013) present a computational model of helping, based on inverse planning of actions with hierarchically structured, embedded goals.

In contrast, understanding of means-end action sequences could be easier in social contexts for either of two reasons. First, social contexts allow for an unpacking of multistep actions into their component parts. The agent and the helper can divide and conquer by each engaging in a distinct, single-step, direct action: The helper aids in the opening of the box, and then the protagonist extracts the object. By decomposing the multistep action into two single-step actions, infants need not ascribe a third-order social goal to a helper. Instead, they could associate the helper with the creation of a state that the protagonist desires for some reason, without attributing to the helper any knowledge of what that reason is.

Second, infants, like adults, may attribute third-order goals to a helper in this situation, because goal attribution is a form of mental state inference, and children's mental state inferences may be enhanced in social contexts. Consistent with this possibility, Hamlin et al. (2013) found that 10-month-old infants selectively reached for an agent who opened a door that provided a protagonist with access to its preferred toy over one who opened a different door that provided access to a non-preferred toy only if the agents had seen the protagonist directly reach for and grasp one toy over

the other in familiarization (i.e., the agents were knowledgeable of the protagonist's preference). Critically, obtaining the object at test (but not during familiarization) required a two-step action because the act of helping was not a direct action on the toy, but an action on one of two barriers that appeared at test preventing access to the toys. Because the infants reached for the agent who provided the protagonist with access to its preferred toy only if the agents had knowledge of the protagonist's preference, infants' behavior suggests that they were sensitive to the means-end structure of the action required to obtain the protagonist's goal, and that they represented the protagonist's goal as being embedded in the agent's goal. That is, in contrast to findings that 10-month-olds as a group do not infer the ultimate goal of a single agent's means-end action (Sommerville & Woodward, 2005), 10-month-old infants may be able to reason about hierarchically structured, embedded goals in social evaluative contexts. One notable difference, however, is that Hamlin et al.'s protagonist engaged in a single-step action in familiarization, whereas Sommerville and Woodward's actor engaged in two-step actions. It is unknown if familiarization to the protagonist grasping a toy (i.e., what would later be the ultimate goal in test trials) facilitated infants' means-end reasoning in Hamlin et al.'s experiment.

Finally, action planning could be equally difficult in situations involving helping, relative to those in which infants themselves act on objects or observe the object-directed actions of a single agent. Infants may begin to understand helping in situations requiring multistep actions at the same time that they begin to plan these actions for themselves and understand the actions of others in non-social contexts. This prediction follows from the hypothesis that core social knowledge involves shared experiences: When infants view social characters in an interaction, they may view them as *acting as one*, with a shared body of knowledge, or common ground (Clark, 1996). Consistent with this possibility, evidence for sensitivity to the unitary character of shared social actions comes from studies of adults with no access to a conventional language (Gleitman et al., 2019) as well as recent studies of infants (Papeo et al., 2020).

As a first step toward testing these contrasting predictions, we assessed infants' evaluations of helpers in multistep, means-end actions at 15 months of age (Experiment 1) and 8 months of age (Experiment 2). We focused on these ages, because past work has demonstrated that infants of these two ages should have different capacities to enact and reason about others' two-step, means-end actions (Sommerville & Woodward, 2005), yet infants at both ages show preferences for helpers in a broad array of situations (Margoni & Surian, 2018). In both of the present experiments, we showed infants videotaped events in which a protagonist—a bear puppet—tried to open one of two boxes, each with a toy inside. The bear was only able to open the box that it had chosen with help, in alternation, from each of two helpers—rabbit puppets who wore distinctive clothing and appeared on the two sides of the stage. Once the box was opened, the bear pounced on the toy, in a manner that indicated to adults that the toy was

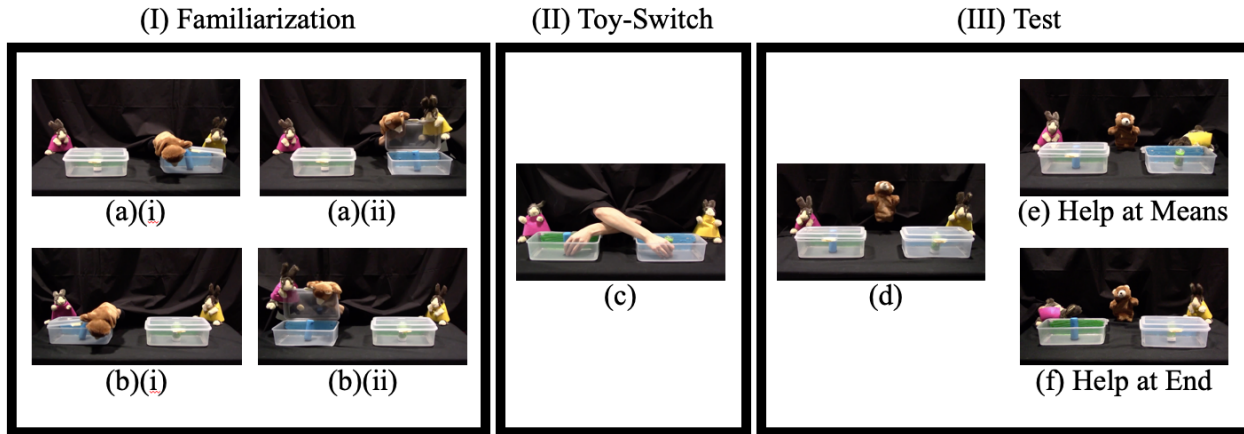


Figure 1. Stimuli in Experiments 1 and 2, including familiarization (I), the toy-switch (II), and test events (III). In familiarization events (I), the Protagonist repeatedly tried and failed to open the same box, which switched sides on the stage between events ((a)(i), (b)(i)). The Helper on that side of the stage helped the Protagonist lift the box's lid ((a)(ii), (b)(ii)), and the Protagonist grasped the toy inside. In the toy-switch event (II), a pair of hands switched the toys' positions as the Helpers were present (c). During test events (III), the Protagonist first jumped between the two boxes. Next, in alternating events, the Means-Helper opened the original box that the Protagonist had tried to open in familiarization (e), and the End-Helper opened the box containing the original toy that the Protagonist had grasped in familiarization (f).

its second-order goal. After this familiarization, we switched the locations of the toys in the presence of the rabbits, who observed this event. Thus, the first- and second-order goals of the bear's original means-end action in familiarization were now separated. In alternation, we presented test events in which each rabbit opened one of the boxes for the bear; that is, one opened the original box, while the other opened the second box that now contained the toy that the bear had pounced on. Thus, one rabbit (the Means-Helper) opened the same box that the bear had sought to open before, whereas the other rabbit (the End-Helper) opened the box containing the object that the bear had sought before. All action was paused after the boxes were opened in the test trials (i.e., infants never saw the bear react to the box being opened). To prefer the End-Helper, therefore, infants would need to reason that the End-Helper had created a state in which the bear could achieve its first-order goal, but without seeing this goal being achieved. In all previous work on infants' evaluations of helping (including Hamlin et al., 2013), by contrast, the protagonist always acted on a goal after a helper created a state in which the protagonist could achieve its goal. By preventing infants from seeing the bear's final action, our method prevented infants from evaluating either rabbit's action based on the valence of the bear's action that followed. Thus, mature performance on our task required that infants take account of the hierarchical structure of the helpers' goals: They could not succeed by representing the helpful rabbit's action and the bear's action as sequence of two acts with distinct first-order goals.

The present experiments were therefore uniquely poised to investigate infants' reasoning about the two-step action plans of other agents in a social evaluative context. If social evaluations complicate the task of action understanding,

because they involve the coordination of the actions of two agents with hierarchically structured goals, then even 15-month-old infants might fail to represent the opening of the second box containing the original goal object as the more helpful action, and therefore fail to prefer the End-Helper. If social evaluations enhance infants' reasoning of means-end actions, then even 8-month-old infants might successfully represent the opening of the second box containing the original goal object as the more helpful action, and therefore prefer the End-Helper. Finally, if the only factor bearing on the development of infants' action prediction and social evaluation is their capacity for formulating and inferring two-step action plans, then we should observe the same developmental change in the present experiments as in the studies of Piaget, and Sommerville and Woodward: 8-month-old infants should judge the puppet who opens the original box that previously held the original goal object to be more helpful (the Means-Helper), because that was the box that the bear sought to open during familiarization. In contrast, 15-month-old infants should judge the puppet who opens the second box that now holds the original goal object to be more helpful (the End-Helper), because that was the object that the bear consistently sought during familiarization.

Experiment 1

Methods

Participants Twenty-four full-term 15-month-old infants contributed data to this experiment (11 girls; mean age = 15.05 months; range = 14;1–15;15). An additional 9 participants began the experiment but were not included in the final sample due to fussiness ($n = 5$), inattentiveness ($n = 3$), and parental interference ($n = 1$). Blind experimenters

determined exclusions using pre-set criteria. For all studies, participants were tested with parental informed consent.

Displays Infants sat on their caregiver's lap before a 40" by 52" LCD projector screen. Two speakers located on the sides of the screen played all stimuli-related sounds. Parents were instructed to close their eyes and not influence their infants. Each infant viewed 6 familiarization events, 1 event in which toys switched positions, and 4 test events, for a total of 11 events. All events depicted 2 transparent boxes (one blue, one green), and 2 toys (one blue, one green) that were inside the boxes. Events are outlined below (see Fig. 1).

All *familiarization* events began with two bunnies (the Means-Helper and the End-Helper; one wearing pink, one wearing yellow) sitting at a stage's rear corners. At the start of each familiarization event, a bear puppet (the Protagonist) jumped onto the stage behind one of the boxes and approached the side of the box that was closer to the stage's center. The Protagonist looked towards the box twice. The Protagonist then jumped on a corner of the lid of the box, grasped it, and made 4 attempts to open the box, failing each time. On the Protagonist's fifth attempt to open the box, as the Protagonist grasped one corner of the lid of the box, and the Helper on the other side of the box moved forward and jumped to grasp the remaining corner of the lid. The acting Helper and the Protagonist then opened the box together. Once the box was open, the Protagonist jumped on top of the open lid and laid its head down to grasp the toy inside. The acting Helper then jumped forward to the stage next to the box, and returned to its corner of the stage. Once the acting Helper returned to its corner, all action paused. These helping actions were based on a show that has been used in prior literature to depict helping (e.g., Hamlin & Wynn, 2011).

In all 6 familiarization events, the Protagonist always approached and tried to open the same box, demonstrating that it had a preference. Between familiarization events, while all puppets were off stage, a pair of hands came out and switched the boxes' positions. Thus, the box that the Protagonist approached appeared alternately on the left and right, and the helper puppets consistently helped when the box was on its side. Because only the boxes' positions switched between events, and the Helpers' corners did not change between events, each Helper took turns being the acting Helper that helped the Protagonist to open the box.

After familiarization, infants saw a single *toy-switch* event in which a pair of hands opened the two boxes while the Helpers were present on stage. The hands took the toys in the boxes, and switched them. Thus, the original box that the Protagonist had tried to open now contained a different toy, and the box that the Protagonist had not tried to open but now contained the original toy that the Protagonist had sought.

Lastly, infants saw *test* events, which began with the Means-Helper and the End-Helper sitting at the stage's rear corners. The boxes were on stage as in familiarization events, except that their contents had been switched. At the start of each test event, the Protagonist jumped onto the stage at the center, and jumped up and down two times as though requesting help. In each event, one of the Helpers moved

forward and jumped to grasp the corner of the lid of one box. If the acting Helper was the Means-Helper, it opened the original box that the Protagonist had tried to open, although the toy inside was now different. If the acting Helper was the End-Helper, it opened the box that the Protagonist had not tried to open but now contained the original toy that the Protagonist had grasped. All action was paused: i.e., infants never saw the bear grasp either toy in test events.

Counterbalancing The following were counterbalanced across infants: color of the box that the Protagonist tried to open in familiarization events (blue/green); End-Helper show side (left/right); End-Helper order (first/second); End-Helper color (pink/yellow); and End-Helper choice side (left/right).

Measures and analyses Our principal measure was selective reaching. Following test events, we presented infants with a *choice* between the Means-Helper and the End-Helper. First, parents turned 90 degrees to the left so that they were no longer facing the screen, and closed their eyes. An experimenter blind to the puppets' identities kneeled in front of infants and held the Means-Helper and the End-Helper approximately 30 cm apart and initially out of reach for infants. Infants were required to look at both puppets before looking back to the experimenter; the puppets were then moved within reach. The experimenter determined a choice as the first puppet infants touched via a visually guided reach (i.e., a touch preceded by a look).

Additionally, looking time data were coded online for familiarization and test events using Xhab64 (Pinto, 1995) software from the moment that all action paused during an event until infants looked away for 2 consecutive seconds or until 30 seconds elapsed. The observer watched infants through a live video feed in a separate room, could not hear or see events, and was blind to counterbalancing. We examined whether looking time differed after test events involving the Means-Helper and the End-Helper.

Results

All reported *p*-values are two-tailed. Fifteen-month-old infants preferred the End-Helper to the Means-Helper (21/24 infants chose the End-Helper, binomial $p < .001$, relative risk = 1.75; see Fig. 2). None of the low-level cues of our display (box color, puppet sides, puppet orders) predicted infants' choice behavior ($ps < .526$).

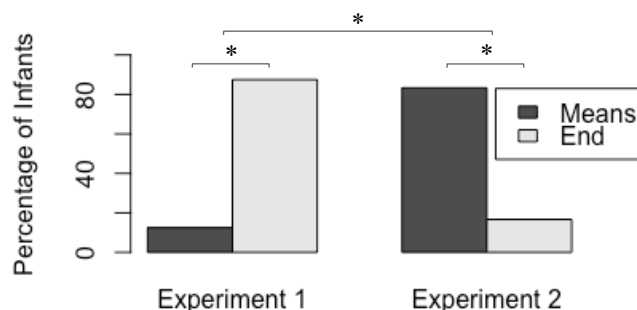


Figure 2. Percentage infants choosing Means-Helpers and End-Helpers in Experiment 1 ($n = 24$) and Experiment 2 ($n = 24$). Asterisks indicate significant differences ($*p < .001$).

To determine whether looking time differed in test events when the Means-Helper and the End-Helper acted, we examined infants' looking times in a gamma mixed-effects model (see Table 1 for summary of descriptive statistics). The dependent variable was looking time. The fixed effects were Helper type (Means/End), event pair (first/second), and the interaction between them. The random effect was participant ID. We found that Helper type did not predict looking time ($b = -0.03$, $\beta = -0.00$, $t = -0.26$, $p = .790$, 95% CI of β [-0.02, 0.02]), that looking time was lower during the second event pair vs. the first event pair ($b = -0.27$, $\beta = -0.02$, $t = -2.20$, $p = .027$, 95% CI of β [-0.05, -0.003]), and that there was no significant interaction between Helper type and event pair ($b = 0.25$, $\beta = 0.01$, $t = 1.02$, $p = .303$, 95% CI of β [-0.01, 0.03]).

Table 1: Infants' mean looking time (s) ¹ after test events

	Attention to test events involving helping at means	Attention to test events involving helping at end
Experiment 1: 15-month-olds		
Event Pair 1	8.95 (1.12)	7.20 (0.72)
Event Pair 2	6.11 (1.10)	6.40 (0.76)
All Events	7.56 (0.80)	6.81 (0.76)
Experiment 2: 8-month-olds		
Event Pair 1	5.26 (0.74)	5.23 (0.50)
Event Pair 2	4.72 (1.27)	4.14 (0.77)
All Events	4.99 (0.78)	4.70 (0.64)

Discussion

In Experiment 1, 15-month-old infants selectively reached for the End-Helper over the Means-Helper. By the measure that is standardly used to assess infants' social evaluations, these findings provide evidence that infants determined that the puppet who opened the new box with the goal object was more helpful, and therefore, a more valuable social partner. Our present findings suggest that 15-month-old infants' evaluations of helpers in means-end actions depend on their successful analysis of the two-step action that satisfied the protagonist's first-order goal. Fifteen-month-old infants' preference for the End-Helper was based on an understanding that the Protagonist's first-order goal in familiarization was to reach the specific toy that had been in the original box.

Nevertheless, 15-month-old infants' looking times in test events did not vary depending on which box was opened. This finding suggests that infants did not hold expectations for how the Helpers would act during test events, in conflict with past findings that infants, aged 12 months, will expect that an agent who has opened a box in order to grasp a toy will open a new box if the toy inside has changed location (Woodward & Sommerville, 2000). This inconsistency may be explained by two competing effects on infants' looking in test trials. By 3 months of age, infants prefer looking at

helpers over hinderers (Hamlin & Wynn, 2011). In addition to being surprised to see a helper open the original box because its actions are focused on the first-order goal, then, infants may prefer to look at the helper who acts on the new box because it has done something that is more helpful. These two competing reasons to look longer at one helper vs. the other helper could lead to the lack of a difference in looking times in our test trials.

In sum, we found that 15-month-old infants selectively reached for the End-Helper to the Means-Helper. In Experiment 2, we investigated 8-month-old infants' evaluations. Past research suggests that infants of this age are less capable of formulating two-step action plans, and that they are unlikely to infer the first-order goal of a means-end action (Piaget, 1952; Sommerville & Woodward, 2005). If 8-month-old infants' evaluations reflect their capacity for generating and recovering two-step action plans, then they should prefer the Means-Helper. If, however, a social context enhances their means-end understanding or decomposes the multistep action, then they may also prefer the End-Helper.

Experiment 2

Methods

Participants Twenty-four full-term 8-month-old infants contributed data to the experiment (8 girls; mean age = 7.86 months; range = 7;9–8;11). An additional 9 participants began Study 2 but were not included in the final sample due to inattentiveness ($n = 3$), failure to choose between puppets ($n = 3$), parental interference ($n = 2$), and fussiness ($n = 1$).

Procedure This was identical to that of Experiment 1.

Results

Eight-month-old infants preferred the Means-Helper to the End-Helper (20/24 infants chose the Means-Helper, binomial $p < .001$, relative risk = 1.66). Patterns of choice differed significantly across Studies 1 and 2 ($\chi^2(1) = 21.37$, $p < .001$, Cohen's $h = 1.49$, Wald's odds ratio = 35, 95% CI [6.94, 176.39]; see Fig. 2). As in Experiment 1, low-level cues of our display did not predict infants' choices ($ps < .282$).

To determine whether looking time differed in test events when the Means-Helper and the End-Helper acted (see Table 2 for summary of descriptive statistics), we ran a gamma mixed-effects model with the same dependent variable, fixed effects, and random effect as in Study 1. As in Experiment 1, we found that Helper type did not predict looking time ($b = -0.01$, $\beta = -0.00$, $t = -0.14$, $p = .844$, 95% CI of β [-0.03, 0.02]). Additionally, we found that looking time did not vary by event pair ($b = -0.08$, $\beta = -0.01$, $t = -0.68$, $p = .495$, 95% CI of β [-0.03, 0.01]) and that there was no significant interaction between Helper type and event pair ($b = -0.26$, $\beta = -0.01$, $t = -1.00$, $p = .315$, 95% CI of β [-0.04, 0.01]).

¹ All numbers in parentheses are standard errors.

Discussion

In Experiment 2, 8-month-old infants showed the opposite reaching preference to that of the 15-month-old infants in Experiment 1. Whereas the older infants reached for the End-Helper over the Means-Helper, the younger infants reached for the Means-Helper over the End-Helper.

These findings provide evidence that the 8-month-old infants valued the helper who fostered the attainment of the Protagonist's first action in a two-step action sequence over the helper who fostered the Protagonist's ultimate goal. This finding suggests that 8-month-old infants did not infer the ultimate goal of the Protagonist, and instead only inferred that the Protagonist had the goal of opening the original box that it had tried to open in familiarization. Our finding encourages a reanalysis of past work (Hamlin & Wynn, 2011) that examined 5- and 9-month-old infants' evaluations of helpers and hinderers using a show depicting a protagonist trying to open a box containing a toy: Younger infants may have positively evaluated an agent who opened a box not because they viewed the protagonist as wanting the object inside, but as wanting the box to be open.

General Discussion

In two experiments, we investigated whether infants' understanding of multistep actions informs their social evaluations. In Experiment 1, 15-month-old infants selectively reached for a helper who opened a new box containing the toy that the Protagonist had sought to attain in familiarization, over a helper who opened the original box that the Protagonist had opened in familiarization. By contrast, in Experiment 2, 8-month-old infants demonstrated the opposite preference when presented with exactly the same events. These findings suggest that infants of both ages evaluated helpers in accord with their changing understanding of the hierarchical structure of means-end action sequences.

The present findings are consistent with the hypothesis that infants' evaluations of helpful actions depend on their capacity for generating and recovering two-step action plans. Specifically, past work has demonstrated that infants do not reliably plan two-step, means-end actions until later in the first year of life (Bates et al., 1980; Diamond, 1985, 1991; Piaget, 1952; Sommerville & Woodward, 2005; Willatts, 1999), and that they do not reliably infer the first-order goal of a two-step, means-end action until 12 months of age (Sommerville & Woodward, 2005; Woodward & Sommerville, 2000). The present experiments provide evidence for the same developmental limits to action planning in a social evaluative context.

Our findings are striking given that infants never saw the Protagonist respond to the box being opened in test events. To prefer the End-Helper, the older infants evidently reasoned that the End-Helper had created a state in which the Protagonist could grasp the original toy, despite not having seen the Protagonist actually act on the original toy in test events. To prefer the Means-Helper, moreover, the younger infants evidently reasoned that the Means-Helper fulfilled the

Protagonist's desired goal state of a particular open box, even though the Protagonist did not act on either box in test events.

Although 15- and 8-month-old infants showed opposite evaluations of End- and Means-Helpers, it is an open question as to when in development this reversal of preference occurs, relative to the age at which infants reliably demonstrate the ability to infer the first-order goal of a two-step action: i.e., 12 months (Sommerville & Woodward, 2005; Woodward & Sommerville, 2000). It may be that the ability to evaluate helpers who foster the attainment of the ultimate goal-state in a means-end action develops: (i) some time after 12 months and before 15 months, if a social evaluative context makes means-end understanding more difficult; (ii) some time after 8 months and before 12 months, if a social evaluative context facilitates means-end understanding; or (iii) at 12 months, if means-end understanding is equally manifest in social evaluative and non-evaluative contexts. The application of our methods to infants at intermediate ages could tease apart these possibilities.

Future work should explore whether infants' evaluations were only based on what they saw as a more positive outcome, based on their action understanding (an associative account), or whether infants' evaluations also depended on an analysis of the helpers' mental states. That is, 15-month-olds could have preferred the End-Helper only because it had opened the box with the toy that the Protagonist had sought in familiarization. Although infants must have represented the protagonist's goal, they may not have reasoned that the helpers represented also represented the Protagonist's goal. Similarly, 8-month-olds could have preferred the Means-Helper only because it had opened the box that the Protagonist had tried to open in familiarization; it is unknown whether 8-month-olds also reasoned that the helpers represented the Protagonist's goal of opening that box.

Evidence against an associative account comes from two sources. First, by 8 months of age, infants' evaluations privilege helpers' intentions over action outcomes: They do not distinguish between an agent who unsuccessfully attempts to help and an agent who successfully helps, even though only the successful helper is associated with a positive outcome (Hamlin, 2013). Moreover, by 10 months of age, infants selectively reach for an agent who provides a protagonist with access to a preferred toy over an agent who provides the protagonist with access to a non-preferred toy only if the helpers have knowledge of the protagonist's preference (Hamlin et al., 2013). Future research could test the associative account by studying infants' evaluations, in means-end action contexts, of helpers who manifest different intentions or states of knowledge.

If the associative account of the present findings is wrong, then our findings suggest capacities for action planning beyond what has been demonstrated in past work. In Experiment 1, in addition to being able to represent the Protagonist's first- and second-order goals in familiarization, 15-month-old infants would need to infer that the Protagonist's first- and second-order goals are embedded in

the End-Helper's social goal. Likewise, in Experiment 2, in addition to representing the Protagonist's goal of opening the original box in familiarization, 8-month-old infants would need to infer that this goal is embedded in the Means-Helper's social goal. That is, if the associative account is wrong, then 15- and 8-month-old infants can be credited with representing the third-order goal of an End-Helper, and the second-order goal in a Means-Helper, respectively. Such abilities are notable, because children as old as 3 years struggle to engage in three-step action sequences (i.e., sequences that would require formulating a first-order, second-order, and a third-order goal; Metevier, 2006), and infants do not reliably formulate and recover two-step action plans until later in the first year of life (Bates et al., 1980; Diamond, 1985; Piaget, 1952; Sommerville & Woodward, 2005; Willatts, 1999; Woodward & Sommerville, 2000). Future work is critical to determine whether the associative account is responsible for findings here, or whether we may indeed be observing third- and second-order goals in Experiments 1 and 2, respectively.

In sum, the present experiments examined infants' action planning in a social evaluative context. Our findings suggest that infants' evaluations are consistent with their capacity for formulating and recovering two-step action plans. The present paper calls for future work to probe the development of infants' evaluations of helpers in multistep action contexts more finely, and to test whether infants' evaluations of helpers depend on associative learning about the outcomes of helpers' actions or their attributions, to helpers, of social intentions and social knowledge.

Acknowledgements

This material is based upon work supported by the Center for Brains, Minds and Machines, funded by National Science Foundation STC award CCF-1231216. We thank the families who volunteered to participate, and we thank Tomer Ullman for providing feedback on an earlier draft.

References

- Baker, C. L., Jara-Ettinger, J., Saxe, R., & Tenenbaum, J. B. (2017). Rational quantitative attribution of beliefs, desires and percepts in human mentalizing. *Nature Human Behaviour*, 1(4), 1-10.
- Baker, C. L., Saxe, R., & Tenenbaum, J. B. (2009). Action understanding as inverse planning. *Cognition*, 113(3), 329-349.
- Bates, E., Carlson-Luden, V., & Bretherton, I. (1980). Perceptual aspects of tool using in infancy. *Infant Behavior and Development*, 3, 127-140.
- Clark, H. H. (1996). Using Language. Cambridge, UK: Cambridge University Press.
- Diamond, A. (1985). Development of the ability to use recall to guide action, as indicated by infants' performance on AB. *Child Development*, 56, 868-883.
- Gergely, G., Bekkering, H., & Király, I. (2002). Developmental psychology: Rational imitation in preverbal infants. *Nature*, 415(6873), 755.

- Gleitman, L., Senghas, A., Flaherty, M., Coppola, M., & Goldin-Meadow, S. (2019). The emergence of the formal category "symmetry" in a new sign language. *Proceedings of the National Academy of Sciences*, 116(24), 11705-11711.
- Hamlin, J. K., Ullman, T., Tenenbaum, J., Goodman, N., & Baker, C. (2013). The mentalistic basis of core social cognition: Experiments in preverbal infants and a computational model. *Developmental Science*, 16(2), 209-226.
- Hamlin, J. K., & Wynn, K. (2011). Young infants prefer prosocial to antisocial others. *Cognitive Development*, 26(1), 30-39.
- Hamlin, J. K., Wynn, K., & Bloom, P. (2007). Social evaluation by preverbal infants. *Nature*, 450(7169), 557-559.
- Hamlin, J. K., Wynn, K., & Bloom, P. (2010). Three-month-olds show a negativity bias in their social evaluations. *Developmental Science*, 13(6), 923-929.
- Jaeger, J., Gweon, H., Schulz, L. E., & Tenenbaum, J. B. (2016). The naïve utility calculus: Computational principles underlying commonsense psychology. *Trends in Cognitive Sciences*, 20(8), 589-604.
- Margoni, F., & Surian, L. (2018). Infants' evaluation of prosocial and antisocial agents: A meta-analysis. *Developmental Psychology*, 54(8), 1445-1455.
- Metevier, C. M. (2006). *Tool-using in rhesus monkeys and 36-month-old children: acquisition, comprehension, and individual differences*. University of Massachusetts Amherst.
- Papeo, L., Nicolas, G., & Hochmann, J. (2020). Visual perception grounding of social cognition in preverbal infants. <https://doi.org/10.31219/osf.io/v7y9x>
- Piaget, J. (1952). *The origins of intelligence in the child*. London: Routledge & Kegan Paul.
- Pinto, J. (1995). Xhab64 [Computer software]. Palo Alto, CA: Stanford University.
- 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), 1-30.
- Sommerville, J. A., Woodward, A. L., & Needham, A. (2005). Action experience alters 3-month-old infants' perception of others' actions. *Cognition*, 96(1), B1-B11.
- Ullman, T.D., Baker, C.L., Macindoe, O., Evans, O., Goodman, N.D., & Tenenbaum, J.B. (2010). Help or hinder: Bayesian models of social goal inference. *Advances in Neural Information Processing Systems*, 22, 1874-1882.
- Willatts, P. (1999). Development of means-end behavior in young infants: Pulling a support to retrieve a distant object. *Developmental Psychology*, 35(3), 651-657.
- Woodward, A. L., & Sommerville, J. A. (2000). Twelve-month-old infants interpret action in context. *Psychological Science*, 11(1), 73-77.