

Do Infants Think That Agents Choose What's Best?

Laura Schlingloff¹ (schlingloff_laura@phd.ceu.edu)

Denis Tatone¹ (denis.tatone@gmail.com)

Barbara Pomiechowska¹ (barbara.pomiechowska@gmail.com)

Gergely Csibra^{1,2} (CsibraG@ceu.edu)

¹Cognitive Development Center, Department of Cognitive Science, Central European University,
Október 6. u. 7, 1051 Budapest, Hungary

²Department of Psychological Sciences, Birkbeck, University of London,
Malet Street, London WC1E 7HX, UK

Abstract

The naïve utility calculus theory of early social cognition argues that by relating an agent's incurred effort to the expected value of a goal state, young children and infants can reason about observed behaviors. Here we report a series of experiments that tested the scope of such utility-based reasoning adopted to choice situations in the first year of life. We found that 10-month-olds (1) did not expect an agent to prefer a higher quantity of goal objects, given equal action cost (Experiment 1) and (2) did not expect an agent to prefer a goal item that can be reached at lower cost, given equal rewards (Experiment 2a and 2b). Our results thus suggest that young infants' utility calculus for action understanding may be more limited than previously thought in situations where an agent faces a choice between outcome options.

Keywords: infant social cognition; action understanding; teleological reasoning; naïve utility calculus

Introduction

For members of a social species like humans, making sense of others' behaviors and adjusting responses to them appropriately is a crucial skill. However, such a task is not easy: interpreting observed actions and inferring others' goals is a backward inference that is always underdetermined, as similar-looking actions can be motivated by different underlying intentions. The problem is particularly pressing for children and infants: they have to learn from and about others with limited prior experience to go on.

Following previous work from computational cognitive modelling (e.g., Baker, Saxe & Tenenbaum, 2009), Jara-Ettinger et al. (2016) have proposed that young children's action understanding can be modeled as a case of Bayesian inverse planning through which an observer can infer an agent's underlying world model and/or utility functions. According to the theory of the naïve utility calculus, we see other agents as utility maximizers who act in a way that maximizes the trade-off between expected costs of an action and expected benefits of a goal (i.e., the net utility). An efficient agent can thus be understood as one who minimizes cost and maximizes benefits.

This theory is supported by a growing body of research in social cognitive development. Studies with preverbal infants suggest that a form of utility-based social reasoning may be already present from a young age: infants expect agents to maximize utility by behaving efficiently. In particular, they assume that agents will minimize the action costs required to bring about a goal, for example by taking the shortest

available path towards it (Csibra, 2008; Gergely et al., 1995; Liu & Spelke, 2017).

To date, most of the developmental literature on action interpretation has focused on infants' understanding of the efficiency of goal-directed actions. However, there is another way to maximize utility: by choosing among different alternative goal options the one that yields the highest net utility. For instance, someone might reach for a cracker over a cookie when both items are equally far away, while going for a cookie over a cracker when the cookie is closer. From an observer's perspective, this behavior pattern may seem inconsistent at first glance, but it makes sense if one understands that the person flexibly chooses what is most beneficial depending on the context. Indeed, from watching such a scenario, five-year-olds infer that the agent assigns higher value to crackers over cookies, but is not willing to incur much extra effort to reach them if cookies are less costly to obtain (Jara-Ettinger et al., 2015).

Can infants also perform such computations? That is, can they compare the relative utility of different goal options available to an agent, and attribute a goal based on this? A recent study by Liu et al. (2017) suggests that they may. Here, 10-month-olds saw an agent approach two different goals equally often but at varying costs, and subsequently expected the agent to prefer the goal it had previously reached through a costlier action.

Building on this finding, the present study aimed to test whether infants at this age would use the naïve utility calculus productively to infer an agent's behavior in a novel context where the agent could choose between two goals of different utility. We implemented this choice context in two ways: (1) an agent could approach a higher quantity of goal objects (Experiment 1), and (2) an agent could reach one of two identical goal objects at relatively lower costs (Experiments 2a and 2b). For instance: (1) if infants understand that someone likes bananas, would they think that this person should prefer *more* rather than *fewer* bananas, given the choice? And (2) would infants assume that the person should go for a banana that is *easier* rather than one that is *harder* to reach?

Experiment 1

Since the naïve utility calculus for young infants is (in part) about energetic costs and benefits, as indicated by their

expectation that shorter and less effortful paths are preferable (Gergely et al., 1995), infants might expect an agent to maximize her utility by choosing a larger amount of a valued resource. This hypothesis rests on the assumption that obtaining more of something confers higher benefits than less of it. This assumption has been shown to guide infants' own decision-making: when 10-month-olds have a choice between two different amounts of crackers hidden in opaque containers, they reliably crawl toward the relatively higher amount (Feigenson, Carey, & Hauser, 2002).

We conducted a looking-time study that aimed to test whether infants would expect others to similarly prefer a relatively higher quantity of goal objects of the same kind. Having watched an agent selectively approach a specific kind of goal object (e.g., 1 banana over 1 strawberry), would infants expect the agent to opt for a larger amount of the goal object (3 bananas over 1 banana) when it becomes available? If infants assign a higher utility to three goal objects grouped together relative to one, they should look longer when the agent approaches the single item. On the other hand, if infants do not have such an expectation, they should respond as in a classical outcome-choice task (Woodward, 1998), looking longer to the three-object outcome as this represents a novel goal for the agent.

Methods

Participants Twenty-four 10-month-old infants (age range: 9 m 16 d – 10 m 12 d, $M = 10.03$ m) participated in Experiment 1. An additional 9 infants were tested, but had to be excluded due to experimenter error ($n = 4$), failure to meet the predefined attention criteria ($n = 2$), fussiness ($n = 2$), or parental interference ($n = 1$). Participants were healthy, full-term infants recruited through a local database. Written informed consent was obtained from the parents before the experiment. The study received full ethical approval from the local ethics board.

Apparatus Infants were seated in their caregiver's lap in a darkened, soundproof room, 80 cm away from a wide-screen 102 cm LCD monitor. The stimuli were 3D animated videos created with Blender animation software (Stichting Blender Foundation, 2018) and presented from a Mac mini computer with MATLAB (Mathworks) using the Psychophysics toolbox extension (Brainard, 1997). Infants were video recorded during the session, and an experimenter watched the video live for online coding to determine the onset and termination of trials.

Procedure and Stimuli Caregivers were instructed to hold infants by their hips without impeding their ability to attend or disengage from the screen. Caregivers' eyes were covered with opaque sunglasses. Before each trial, a short attention-getting clip was shown until the infant attended to the screen. Trials ended either when the infant looked away for 2 consecutive seconds after a video had stopped, or if 8 seconds (in familiarization trials) or 60 seconds (in test trials) had passed since a video ended.

Familiarization. Infants watched a total of 8 familiarization trials. Each trial consisted of a video lasting

7.5 seconds, and the display of the last frame as a still image. In all videos, an agent (a pear-shaped blue character with eyes) approached a goal object. Initially, the agent was located at the top of the screen in a narrow hallway opening up at the bottom, and two goal objects (a banana and a strawberry) were located at the bottom, on the left and right side of the screen, respectively. The agent first moved downward in a straight line, then turned left or right, approaching the goal object on that side, and came to a standstill after having made physical contact with the object (upon which a ringing sound was played). The agent always approached the same type of goal object, which was sometimes located on the left, and other times on the right side of the screen.

Test. Two test videos were shown to the infants. The videos were identical to the familiarization videos in terms of duration, behavior of the agent, and layout, with the exception of the goal objects. There were only tokens of the previously approached goal object kind present; on one side, there was a single item, on the other, three. The single item and the topmost item of the three-item set were placed equidistantly from the agent.

We counterbalanced across participants the type of goal object approached (banana vs. strawberry); the order of the locations approached during familiarization (LRRLLLRR vs. RLLRRLL); the location of the three items at test (left vs. right); and the order of test events (approach-3 first vs. approach-1 first).

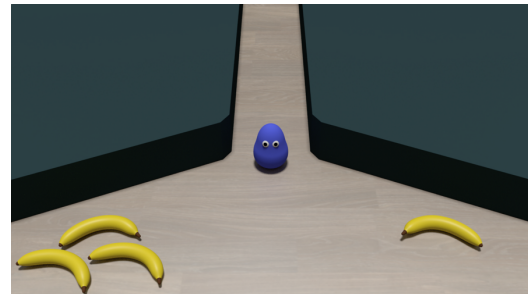


Figure 1: Layout of the stimuli videos (Experiment 1) at test.

Coding and Analysis Infants' looking behavior was manually coded off-line using the same criteria as online coding. The looking times of 50% of the participants was reanalyzed by an independent second coder who was blind to the hypothesis and to the stimuli condition. The average absolute difference between coders was 0.28 s; data from the first coder was used for analyses.

The raw looking times were base-10 log-transformed for analysis (Csibra et al., 2016), but raw data is used for descriptive statistics and plots. As specified in the preregistration (see below), we conducted both Bayesian and frequentist statistical analyses. For the Bayesian analysis, we used the method recommended by Csibra et al. (2016). This method calculates a Bayes Factor which compares a null model to an alternative model that assumes a moderate

increase or decrease in looking times between conditions. For the frequentist statistical analyses, we conducted a paired-sample two-tailed t-test. Moreover, we conducted a 2x2 mixed ANOVA to check for order effects, which is a pattern commonly found in looking-time studies with infants (e.g., Liu et al., 2017). Statistical analyses were performed in R (version 3.4.1; R Development Core Team, 2017), plots were created using the ggplot2 package (Wickham, 2009).

This study was preregistered at the OSF [<https://osf.io/9h7yg>]; stimuli, data, and analyses can be accessed at [<https://osf.io/6pf3b/>].

Results

The Bayesian analysis suggested that the infants looked longer when the agent approached the three goal objects ($M = 19.31$ s, $SD = 15.14$ s) than when the agent approached the single object ($M = 14.98$ s, $SD = 14.89$ s). This data yielded a BF of 73.45, which indicates a strong effect. In a t-test, this looking-time difference was not significant $t(23) = 1.993$, $p = .058$.

A 2x2 mixed ANOVA did not show a significant Order x Condition interaction, $F(1,22) = 0.919$, $p = .348$.

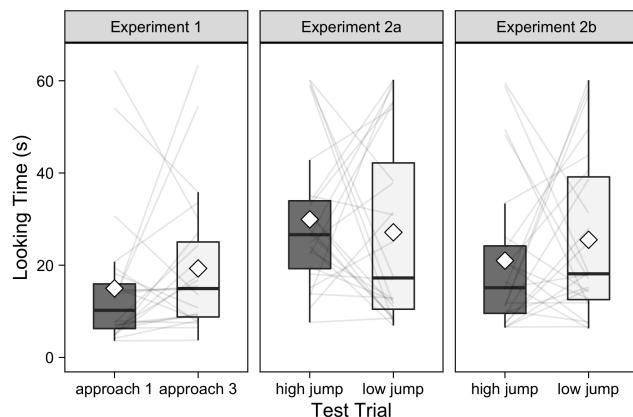


Figure 2: Boxplots of average looking times toward the test events in Experiments 1, 2a, and 2b. Light grey lines connect the looking times of individual participants, white diamonds indicate means, horizontal lines indicate medians, boxes indicate middle quartiles, and whiskers indicate points within 1.5 times the interquartile range from the upper and lower edges of the middle quartiles.

Discussion

The results of Experiment 1 failed to provide evidence that infants expect an agent to maximize her utility by approaching a higher quantity of goal objects. On the contrary, infants looked longer when the agent selected more goal objects compared to when she continued to pursue the same goal as during familiarization, which consisted in approaching a single goal object. This pattern of results is consistent with the interpretation that infants did not interpret the higher quantity of goal objects as an indicator of higher expected value to the agent.

Several explanations may account for these results. One possibility is that infants do not see quantity as a default indicator of how much an agent will value a goal: they may not assume that more of a good thing is necessarily better. Another possibility is that the novelty of the higher quantity of goal objects (which was never shown during familiarization) was disruptive and prevented infants from comparing the relative reward magnitude of the two options. If this is the case, infants may have problems using the naïve utility calculus productively to reason about novel situations (at least in the domain of benefit maximization). Third, it is conceivable that infants may not even represent the scenario as a choice among alternatives. During familiarization, they might have merely attributed a goal to the agent (e.g., approach the banana) without taking into account the non-approached object (cf. Feiman, Carey, & Cushman, 2015), and at test looked longer at the outcome less consistent with the previously attributed goal. Experiment 2 addressed this account.

Experiment 2a

Here we asked whether infants attribute a goal on the basis of the relative expected utility of two goal options when the *benefits* are kept constant, but the *costs* vary, as prior research suggests that infants at this age have a firm grasp on efficient (that is, cost-minimizing) action (e.g., Liu & Spelke, 2017). If infants' reasoning about utility maximization is limited to cost comparisons, they might still expect an agent to choose among equally valuable goals the one that can be reached at the lowest cost.

Preliminary support for this hypothesis comes from a study by Scott and Baillargeon (2013), where 16-month-olds infants expected a person to choose between two identical objects the one that would require fewer steps to be accessed. This result is consistent with infants applying utility reasoning to attribute a goal preference to an agent. In Experiment 2a, we aimed to build on this result. We hypothesized that if infants attribute a utility-maximizing strategy to an agent who is seen performing an efficient action directed at the same goal object at both high and low cost, they should expect the agent to minimize costs by choosing the cheaper option when both options are available. Thus, if there are two identical goal objects present that can be reached by investing relatively more or less effort, infants should look longer when the agent chooses to perform a higher-cost action (jumping over a high rather than a low wall) for the same reward.

Methods

Participants Twenty-four 10-month-old infants (age range: 9 m 18 d - 10 m 15 d, $M = 10.0$ m) participated in Experiment 2a. An additional 14 infants were tested, but were excluded due to failure to meet the attention criteria ($n = 6$), experimenter error ($n = 2$), or ceiling looking times at both test events ($n = 6$). Recruitment, consent, and ethical approval were the same as in Experiment 1.

Apparatus The apparatus was the same as in Experiment 1, with the exception that stimuli were presented with PyHab 0.7.2 habituation software (Kominsky, 2019) in PsychoPy 3.0.6 (Peirce et al., 2019).

Procedure and Stimuli As in Experiment 1, caregivers were instructed to hold infants by their hips, and caregivers' eyes were covered with opaque sunglasses. Before each trial, an attention-getting clip was shown. During familiarization, this was a short clip (2 s); before each of the two test trials, the attention-getter was a longer clip (15 s) to recapture infants' attention. Each trial contained multiple instances of an event, such that the events were shown in a (quasi-) looped manner (see *Familiarization* for details). Trials ended either when the infant looked away for a minimum of two consecutive seconds, or if 46 seconds (familiarization) or 60 seconds (test) had passed since the trial onset. In contrast to Experiment 1, we did not measure looking time to a still image of the video's last frame, but kept playing the looped stimuli until either of the aforementioned termination criteria was met (see for example Csibra et al., 2003, and Liu et al., 2017). The rationale for this procedure change was that infants here were supposed to contrast two actions, and not two goals, about which the action outcomes provide insufficient information.

Familiarization. Infants watched a total of six familiarization trials. Each trial consisted of a maximum of five events (less if the infant ended a trial by looking away for two seconds before a trial ended), which each had a duration of 8.5 s. In each trial, there were two high jump events, two low jumps, and one straight approach. After each event, a black screen was briefly displayed (0.5 s).

The scene shown in the stimuli always contained an agent initially located in the middle of the screen. There was always either a low or a high dark grey wall to the left or right side of the agent. The walls were always in the same location and did not change sides within a subject (such that, for example, the low wall always appeared on the left side). Goal objects were yellow bananas.

Each video within a trial began with a bell sound, indicating the onset of an event. Then a banana fell to the ground (onto a pink landmark), after which the agent turned and moved towards it. If there was a wall in the way of the agent's path, the agent jumped over it. Upon making contact with the banana, the agent came to a standstill, and a ringing sound was played. The timing of approach was kept constant for all familiarization events (low jump, high jump, straight approach). The height of the jump was adjusted to the height of the wall.

Test. Infants watched two test trials. Each trial consisted of the same event, which was looped for a maximum of 60 s. As in familiarization, the videos within a trial were interspersed with a brief 0.5 s display of a black screen.

The layout was similar to the familiarization trials, except that both walls (high and low) were present in the same locations as before. The event played out the same as in familiarization, with the exception that two bananas fell down, and that there was an additional 0.5 s pause before the

agent started moving. In the inconsistent test event, the agent approached the banana behind the higher wall; in the consistent test event, she approached the banana behind the lower wall.

We counterbalanced across participants the location of the high and low walls (high left vs. high right); the side of the first approach during familiarization (LHLHLH vs. HLHLHL); and the order of test events (inconsistent first vs. consistent first).

Coding and Analysis Infants' looking behaviors were coded and analyzed the same way as in Experiment 1. Again, a second coder, blind to the experimental condition, recoded 50% of participants' looking times. The average absolute difference between coders was 0.4 s.

Because of the unexpectedly high number of infants who did not disengage and look away from the screen during the experimental procedure, we used an additional exclusion criteria to avoid the problem of ceiling effects: we excluded participants who did not end at least one of the two test trials with a 2 second look-away.

The preregistration for this study can be found at [https://osf.io/pvy37], stimuli, data, and analyses at [https://osf.io/7j58z/].

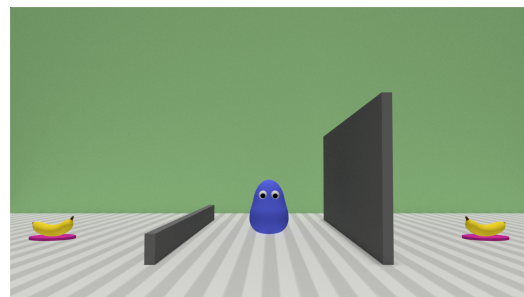


Figure 3: Layout of the stimuli videos (Experiment 2a and 2b) at test.

Results

With the sample of 24 infants we tested, the Bayesian analysis provided some evidence supporting our hypothesis: we obtained a BF of 5.38. However, a t-test did not yield a significant result ($t(23) = 1.272, p = 0.216$), suggesting that infants did not look longer at the high jump action ($M = 29.93$ s, $SD = 15.9$ s) compared to the low jump action ($M = 27.13$ s, $SD = 20.13$ s).

As in Experiment 1, we analyzed whether there was an effect of the order of test video presentation. A 2x2 mixed ANOVA did not show a significant Order x Condition interaction ($F(1,22) = 0.00044, p = .983$).

Discussion

The results of Experiment 2a were not conclusive with respect to our hypothesis. While the Bayesian analysis suggested that the infants indeed looked longer at the agent's

high-jump action at test, and thus had expected her to choose the goal object that can be obtained at lower cost, this effect was weak and the looking-time pattern was only shown by 14 of 24 infants.

Because we used a procedure where the test events were looped, there is a possibility that repeatedly playing the auditory cues that accompanied the events may have driven the infants' attention back to the screen, preventing them from disengaging from the event even after they had lost interest. Analyzing only the data from the first looks, i.e., their looking times until they disengaged from the screen for the first time, yielded a BF of 245.47 (high jump: $M = 23.55$ s, $SD = 15$ s; low jump: $M = 18.64$ s, $SD = 16.76$ s), indicating a strong effect. However, this analysis was post-hoc and cannot be considered as confirmatory. Therefore, we decided to conduct a replication of Experiment 2a, removing the sound effects from the stimuli.

Experiment 2b

Methods

Participants Twenty-four 10-month-old infants (age range: 9 m 18 d - 10 m 14 d, $M = 10.0$ m) participated in Experiment 2b. An additional 13 infants were tested, but were excluded due to parental interference ($n = 1$), fussiness ($n = 1$), failure to meet the attention criteria ($n = 5$), technical failure ($n = 5$), or ceiling looking times at both test events ($n = 1$). Recruitment, consent, and ethical approval were the same as in the previous experiments.

Apparatus The apparatus was the same as in Experiment 2a.

Procedure and Stimuli The procedure and stimuli were identical to the ones in Experiment 2a, except that we removed all sound cues from the familiarization and test trial stimuli.

Coding and Analysis Infants' looking behaviors were coded and analyzed the same way as in Experiment 2a. The average absolute difference between coders was 0.46 s. The preregistration for this experiment is accessible at [<https://osf.io/2h78y>], data and analyses at [<https://osf.io/7j58z/>].

Results

Infants did not look longer at either of the two test events: the Bayesian analysis resulted in a BF of 2.59, providing neither support for our hypothesis nor for the null hypothesis of no effect. The looking times were not significantly different between conditions (high jump: $M = 21.01$ s, $SD = 16.7$ s; low jump: $M = 25.51$, $SD = 17.59$ s; $t(23) = -1.083$, $p = .29$). A 2x2 mixed ANOVA did not show a significant Order x Condition interaction ($F(1,22) = 2.033$, $p = .33$).

An analysis of the data from the first looks yielded the same conclusion: infants did not look significantly longer at either test event (BF: 0.49; high jump: $M = 16.52$ s, $SD = 11.53$ s; low jump: $M = 20.15$ s, $SD = 16.48$ s). Unlike in Experiment 2a, here the pattern of first looks did not differ substantially from that of the overall looking time.

Discussion

The results from Study 2b indicate that, contrary to our prediction, infants did not look longer when an agent chose to perform a costlier action over a less costly action to obtain the same benefit. Analyzing the data from Experiments 2a and 2b together, a mixed ANOVA with Trial (consistent vs. inconsistent test event) as within-subject and Experiment (2a vs. 2b) as between-subject factors showed no main effects (Trial: $F(1,46) = 0.001$, $p = .916$; Experiment: $F(1,46) = 2.802$, $p = .101$) and no interaction effect ($F(1,46) = 2.763$, $p = .103$), which supports this conclusion.

One possible reason for this null result is that infants did not assign equal benefits to the two identical-looking goal objects at test, which would be required for them to evaluate the relative utility of the outcomes. In fact, infants have a propensity to rationalize seemingly irrational actions: for instance, in Liu et al. (2017) infants resolved the apparent inconsistency of an agent sometimes performing a costly action (for goal A) and other times refusing to (for goal B) by inferring that goal A was more valuable to the agent than goal B. In our study, infants may have similarly reasoned that the object behind the higher wall provided a larger benefit, which made the agent approaching that object plausible.

Under this account, both the "consistent" and "inconsistent" test events may have satisfied infants' rationality criteria. The actions of the agent in the two events were not equally efficient with respect to the goal description we had posited ("*reach a banana with as little cost as possible*"); however, since the agent only ever jumped as high over each barrier as was needed, each action was efficient with respect to the goal realized under another description ("*reach this banana with as little cost as possible*").

General Discussion

Reasoning about what other people find valuable and how they might likely behave in order to bring about desirable goals is a crucial component of social cognition and one that children need to master to become proficient social agents. In the present study, we tested whether 10-month-olds would draw goal inferences from the behavior of an animated agent, and expect utility-maximizing actions in a novel context which afforded a choice between goals. Building our predictions on the literature on infants' understanding of goal-directed actions, as well as the theory of the naïve utility calculus, we investigated whether infants would expect an agent to (1) maximize her benefits by choosing more of a preferred goal object, and (2) minimize action costs by choosing one of two identical goal objects that can be reached with relatively less effort.

Our results supported neither of these two hypotheses. While the findings presented here are preliminary and need to be bolstered with conceptual replications and follow-up work, they raise the possibility that infants do not productively generate expectations regarding which goal an agent will approach in a novel context, even if they themselves have been shown to flexibly behave in a utility-maximizing way (e.g., Feigenson et al., 2002; Lucca, Horton

& Sommerville, 2020). As the benefits that agents obtain from a given goal may be opaque to naïve observers and often cannot be perceived from the properties of an object, her actual behavior may represent a more reliable cue for inferring how much an agent values a given outcome. In other words, the fact that an agent incurs higher costs to obtain A over B (as in Liu et al., 2017) is a reliable indicator that the agent assigned a high(er) net utility to A. Under this account, infants may remain agnostic about the magnitude of rewards that given goal objects confer to an agent, when the agent is facing a choice between goal objects differing in quantity or costs required for obtaining them.

If this interpretation is on the right track, it would follow that infants' social reasoning in the first year of life may rely to a greater extent on the agents' actual behavior and the effort they incur, and only gradually begin to integrate information about costs and benefits, enabling more flexible reasoning about the various ways in which utility maximization can be brought about.

Another potential explanation of the present results is that the concept of *choice* available to infants at this age may not be sophisticated enough to allow them to solve the task posed by our studies. These in fact required infants to compute the utilities of two potential outcomes, compare them, and undergird their representation of the agent's goal with whichever is higher. While the usual interpretation of Woodward's (1998) seminal study is that this is exactly what infants do when they observe an agent confronted with a choice situation, our finding raises the alternative possibility that infants may solve that task without contrasting observed and counterfactual outcome options (see also Feiman et al., 2015). Rather, they might disregard the non-chosen item and merely represent the agent's goal.

This account is compatible with infants' success in tasks like those of Gergely et al. (1995), which require that infants, upon observing a goal-directed action, reconstruct whether the same goal could have been brought about in a more efficient way. When the means to accomplish a goal and the corresponding costs are transparent (for instance, a linear relationship between the path length an agent traverses and the energetic cost she incurs), a more efficient alternative action can be straightforwardly identified from visual analysis of the environment the agent acted in. Juxtaposing multiple potential goals, on the other hand, is arguably a more computationally demanding task.

Prior research on infants' understanding of efficient, goal-directed actions shows that they can easily compute utility with respect to a particular goal. While they seem to be able to contrast the relative utility of different means aimed to bring about a particular goal state, and generate expectations accordingly for how a rational agent ought to behave, they may not track either the relative utility of all potential outcomes or the requisite actions to accomplish them, and may not attribute goals on the basis of such a comparison.

Acknowledgments

We thank the families who participated in these experiments, and Petra Kármán for help with data collection. This research has received partial funding from the European Research Council (ERC) under the European Union's Horizon 2020 research and innovation programme under grant agreement no. 742231 (PARTNERS).

References

- Baker, C.L.; Saxe, R. & Tenenbaum, J. B. (2009). Action understanding as inverse planning. *Cognition*, 113(3), 329-349.
- Brainard, D. H. (1997). The Psychophysics Toolbox. [Computer Software] *Spatial Vision*, 10, 433-436.
- Csibra, G. (2008). Goal attribution to inanimate agents by 6.5-month-old infants. *Cognition*, 107, 705-717.
- Csibra, G.; Bíró, S.; Koós, O. & Gergely, G. (2003). One-year-old infants use teleological representations of actions productively. *Cognitive Science*, 27(1), 111-133.
- Csibra, G.; Hernik, M.; Mascaro, O.; Tatone, D. & Lengyel, M. (2016). Statistical treatment of looking-time data. *Developmental Psychology*, 52(4), 521-536.
- Feigenson, L.; Carey, S. & Hauser, M. (2002). The representations underlying infants' choice of more: Object files versus analog magnitudes. *Psychological Science*, 13(2), 150-156.
- Feiman, R.; Carey, S. & Cushman, F. (2015). Infants' representations of others' goals: Representing approach over avoidance. *Cognition*, 136, 2014-214.
- Gergely, G., Nádasdy, Z.; Csibra, G. & Bíró, S. (1995). Taking the intentional stance at 12 months of age. *Cognition*, 56(2), 165-193.
- Jara-Ettinger, J.; Gweon, H.; Schulz, L. & Tenenbaum, J. B. (2016). The Naïve Utility Calculus: Computational Principles Underlying Commonsense Psychology. *TRENDS in Cognitive Sciences*, 20(8), 589-604.
- Jara-Ettinger, J.; Gweon, H.; Tenenbaum, J. B. & Schulz, L. E. (2015a). Children's understanding of the costs and rewards underlying rational action. *Cognition*, 140, 14-23.
- Kominsky, J. (2019). PyHab: Open-source real time infant gaze coding and stimulus presentation software. [Computer software] *Infant Behavior & Development*, 54, 114-119.
- Liu, S. & Spelke, E. S. (2017). Six-month-old infants expect agents to minimize the cost of their actions. *Cognition*, 160, 35-42.
- Liu, S.; Ullman, T. D.; Tenenbaum, J. B. & Spelke, E. S. (2017). Ten-month-old infants infer the value of goals from the costs of actions. *Science*, 358(6366), 1038-1041.
- Lucca, K.; Horton, R. & Sommerville, J. A. (2020). Infants rationally decide when and how to deploy effort. *Nature Human Behaviour*, 4, 372-379.
- Pearce, J. W.; Gray, J. R.; Simpson, S.; MacAskill, M. R.; Höchenberger, R.; Sogo, H.; Kastman, E. & Lindeløv, J. (2019). PsychoPy2: experiments in behavior made easy. [Computer software] *Behavior Research Methods*. 10.3758/s13428-018-01193-y

- R Core Team (2017). R: A language and environment for statistical computing (version 3.4.1). R Foundation for Statistical Computing, [Computer software] Vienna, Austria. Retrieved from <https://www.R-project.org>.
- Scott, R. M. & Baillargeon, R. (2013). Do infants really expect agents to act efficiently? A critical test of the rationality principle. *Psychological Science*, 24(4), 466-474.
- Stichting Blender Foundation (2018). Blender (version 2.79b) [Computer software]. Retrieved from <https://www.blender.org/download>.
- Wickham, H. (2009). ggplot2: Elegant Graphics for Data Analysis. [Computer software] Springer-Verlag.
- Woodward, A. L. (1998). Infants selectively encode the goal object of an actor's reach. *Cognition*, 69(1), 1-34.