

Robotica

<http://journals.cambridge.org/ROB>

Additional services for **Robotica**:

Email alerts: [Click here](#)

Subscriptions: [Click here](#)

Commercial reprints: [Click here](#)

Terms of use : [Click here](#)



Parental scaffolding as a bootstrapping mechanism for learning grasp affordances and imitation skills

Emre Ugur, Yukie Nagai, Hande Celikkanat and Erhan Oztop

Robotica / Volume 33 / Special Issue 05 / June 2015, pp 1163 - 1180
DOI: 10.1017/S0263574714002148, Published online: 19 August 2014

Link to this article: http://journals.cambridge.org/abstract_S0263574714002148

How to cite this article:

Emre Ugur, Yukie Nagai, Hande Celikkanat and Erhan Oztop (2015). Parental scaffolding as a bootstrapping mechanism for learning grasp affordances and imitation skills. Robotica, 33, pp 1163-1180 doi:10.1017/S0263574714002148

Request Permissions : [Click here](#)

Parental scaffolding as a bootstrapping mechanism for learning grasp affordances and imitation skills

Emre Ugur^{†‡*}, Yukie Nagai[§], Hande Celikkanat[¶]
and Erhan Oztop^{||‡}

[†]*Institute of Computer Science, University of Innsbruck, Innsbruck, Austria*

[‡]*Advanced Telecommunications Research Institute, Kyoto, Japan*

[§]*Graduate School of Engineering, Osaka University, Osaka, Japan*

[¶]*Department of Computer Engineering, Middle East Technical University, Turkey*

^{||}*Department of Computer Science, Ozyegin University, Istanbul, Turkey*

(Accepted July 18, 2014. First published online: August 19, 2014)

SUMMARY

Parental scaffolding is an important mechanism that speeds up infant sensorimotor development. Infants pay stronger attention to the features of the objects highlighted by parents, and their manipulation skills develop earlier than they would in isolation due to caregivers' support. Parents are known to make modifications in infant-directed actions, which are often called "motionese".⁷ The features that might be associated with motionese are amplification, repetition and simplification in caregivers' movements, which are often accompanied by increased social signalling. In this paper, we extend our previously developed affordances learning framework to enable our hand-arm robot equipped with a range camera to benefit from parental scaffolding and motionese. We first present our results on how parental scaffolding can be used to guide the robot learning and to modify its crude action execution to speed up the learning of complex skills. For this purpose, an interactive human caregiver-infant scenario was realized with our robotic setup. This setup allowed the caregiver's modification of the ongoing reach and grasp movement of the robot via physical interaction. This enabled the caregiver to make the robot grasp the target object, which in turn could be used by the robot to learn the grasping skill. In addition to this, we also show how parental scaffolding can be used in speeding up imitation learning. We present the details of our work that takes the robot beyond simple goal-level imitation, making it a better imitator with the help of motionese.

KEYWORDS: Developmental robotics; Affordance; Imitation; Parental scaffolding; Motionese.

1. Introduction

Infants, pre- and post-natal, are in constant interaction with their environment for developing sensorimotor skills. The random looking hand arm movements of a newborn serve not only to strengthen his skeleto-muscle system but also to adapt his sensorimotor control. Infants, as young as two weeks old try to create opportunities to enhance their sensorimotor development, for example by keeping their hands in view despite opposing load forces.⁵¹ They realize very early on that they can manipulate their own body as well as their caregivers: a cry brings more care and food, a smile induces a reciprocal smile and so on. In turn, the signals an infant displays in response to the stimulation from the caregiver create a non-trivial infant-caregiver dynamics. This takes the form of so called 'parental scaffolding' when the infant starts to learn more complex movement skills.

In developmental psychology, the support given by a caregiver to speed up the skill and knowledge acquisition of an infant or a child is generally referred to as scaffolding.⁴ For example, a caregiver attracts and maintains a child's attention as well as shapes the environment in order to ease a task to teach to the child. For some skills, scaffolding may include physical interaction such as guidance

* Corresponding author. E-mail: emre.ugur@uibk.ac.at

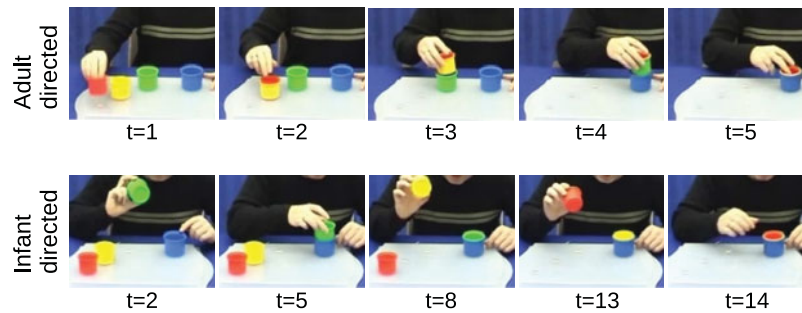


Fig. 1. Snapshots from adult-directed (upper panel) and infant-directed (lower panel) demonstrations of a stacking task. Compared to adult-directed demonstration, infant-directed one took longer time (14 s) and included a number of motionese strategies such as inserting pauses and directing attention to important features.

<http://cnr.ams.eng.osaka-u.ac.jp/yukie/videos-en.html>

from the caregiver to position and orient the child, or to guide the hands and arms to help the child complete a manual task. This guidance is usually performed in conjunction with various signals that highlight the important features or subgoals of a task. In this context, scaffolding can be seen as the process of providing feedback or reinforcement to the child.⁵²

Assisted imitation experiments by Zukow-Goldring and Arbib⁵³ show that caregivers interfere with an ongoing task execution of an infant only at parts where the infant fails. For example in the tasks of concatenating objects or orange peeling, caregivers initially demonstrate the goal, while trying to direct the attention of the child towards the task-relevant features, and then ‘embody’ the child to co-achieve the goal. Throughout this process, caregivers allow the infant to utilize the full sensorimotor experience for learning. The infant, in some sense, learns how it ‘feels’ to execute and complete the task through his/her own proprioceptive, vibrotactile, visual and auditory system. The communication between a caregiver and a child becomes more multimodal and explanatory rather than being simply regulatory especially when the child is on a ‘mission’ to achieve a task.⁴³

In the course of development, infants learn that caregivers would respond to their request for help, and become active seekers of scaffolding opportunities for learning. By the maturation of the ‘joint attention’ ability, which emerges towards the end of the first postnatal year, the infants use more and more of communicative signals,¹⁸ which facilitate, amongst other things, support from caregivers. In general, scaffolding appears to be a critical mechanism for human skill development, because it is known that infants can exhibit certain skills in game-contexts with their mothers long before they can display the skills alone in controlled settings.²¹

Caregivers modify their movement patterns when teaching a task to an infant, much the same way as when they make changes in their vocal production to accommodate the infant auditory system. As the latter is commonly called ‘motherese’, Brand *et al.*⁷ analogously coined the term ‘motionese’ to describe the higher interactiveness, enthusiasm, proximity, range of motion, repetitiveness and simplicity that may be observed when teaching an infant a new task. Reinforcing this notion, Nagai and Rohlfing²⁸ found that when interacting with infants, as opposed to other adults, parents instinctively try to employ a number of strategies in order to increase the saliency of the objects and their initial and final states (see Fig. 1). These strategies include suppressing their own body movements before starting to execute the task, and generating additional movements on the object, such as tapping it on the table. The resulting behavior is qualitatively different from adult-adult interaction in observable parameters, such as the pace or the smoothness of the movement.³⁷ It is shown that infants also benefit from this behavior. Koterba and Overson²³ showed that infants exposed to a higher repetition of demonstration exhibit longer bangs and shakes of objects whereas infants exposed to a lower repetition spent more time for turning and rotating objects.

Nagai and Rohlfing³¹ developed a bottom-up architecture of visual saliency that can be used in an infant like learning robot. In this architecture, similar to an infant, the robot tends to find the salient regions in caregiver’s actions, when they are highlighted by instinctive parental strategies. Moreover, the saliency preference of the robot also motivates humans to use motionese as if they were interacting with a human infant, thereby completing the loop.²⁹ Due to the limited attention mechanism of the

robot, human subjects tried to teach the task (cup-stacking) by approaching towards the robot and introducing the object in the proximity of the robot. In general, subjects amplified their movements and made pauses as if to give the robot a chance to understand the scene. Such delimiting pauses and exaggerations are also observed in infant-directed speech,¹⁴ and infant-directed sign language.²⁵

Using human-robot interaction appears to be a fruitful path for robot skill synthesis, and as such it has been employed in a wide variety of contexts with different flavors. On one extreme, the human simply provides a ‘demonstration’ which is adapted to the robot by manual work.⁴¹ In the other extreme, the human operator is considered as part of the robot control system and human sensorimotor learning capacity is exploited to synthesize the robot behaviors.^{2,26} Direct teaching, or kinesthetic teaching^{9,24,38} somewhat takes a path in between and requires the human to specify the motion of the robot by physically guiding it from the start to the end of the task.

In the context of manipulation, parental scaffolding may seem similar to kinesthetic teaching; yet there are several important differences. In kinesthetic teaching, the robot is usually passive and guided by the human continuously. In scaffolding, the caregiver intermittently interferes and corrects an ongoing execution rather than completely specifying it. Parental scaffolding has been adopted in several forms for robotics studies with the goal of, for example, better communication between humans and robots.⁸ It is proposed that a well-founded paradigm of scaffolding will provide the basis for lifelong robot development ‘at home’.³⁹ For instance, domestic service robots can keep learning new tasks from their owners, instead of being restricted to manually designed skills.

In teaching by demonstration or imitation learning studies, a recurring theme is that the robot has no way to infer what to imitate unless specified by the designer. In particular, in goal-oriented tasks where any means to achieve the goal is acceptable, the robot has to decide whether to copy the motion or reproduce the observed effect. In means-oriented tasks, the motion itself is equally or more important as the real goal may not be the physical final state of the environment, but the movement pattern itself.^{6,30,33} This is also a problem faced by human infants until at least age of 18 months,^{10,17} where the parents often provide help in the form of scaffolding. In particular, they signal what to imitate by highlighting important features or parts of the movement.³¹ For example, in a goal-oriented task, initial and final states along with the important sub-goals are emphasized by inserting pauses in the movement sequence. If they want to emphasize the view-points of the object trajectory, they add additional movements to the object (for example by shaking it). Saunders *et al.* utilized the notion that the environment can also be modified to provide a complexity reducing scaffold.³⁹ In their work, human trainers teach robots certain behaviors (such as avoid-light or tidy-up) that are composed of hierarchical sets of reusable primitives. The guidance can be either by direct teaching of primitives, where the human controls the robot’s actions, or via manipulating the environment. For instance, in case the robot is expected to associate certain sensors to certain tasks, a person can arrange the environment accordingly to minimize confusing inputs from irrelevant sensors. Argall *et al.*¹ uses a two staged approach where, first, a policy for the target behavior (positioning hand for a grasp) is synthesized from human demonstration; and then during the execution, the human guides the arm via a tactile interface, effectively correcting and refining the policy. These tactile corrections improve the grasping policies in terms of success and precision.

The studies summarized above focus either on robot imitation learning or modeling parental scaffolding alone. On the other hand, our aim in this paper is to realize parental scaffolding in an integrated developmental framework that starts from motor primitive formation and goes through affordance learning and imitation. Learning through self-exploration and imitation are crucial mechanisms in acquiring sensorimotor skills for human infants. We hold that realization of this developmental progression in robots can give rise to human dexterity and adaptability. In fact, our previous research contributed to this research program by showing that with self-exploration a robot can shape its initially crude motor patterns into well controlled, parameterized behaviors which it can use to understand the world around it.^{47,49} However similar to infants, complex skill acquisition may benefit from external help. In this paper, we extend this framework by enabling the robot to use parental scaffolding in two major robot learning problems

- Section 2 describes our first extension attempt⁴⁶ where a human caregiver speeds up the affordance acquisition of a robot for grasp actions through parental scaffolding. Learning complex object grasping through self-exploration is a slow and expensive process in the high-dimensional behavior

parameter space of dexterous robot hands. A human caregiver can step in for help by physically modifying the robot's built-in *reach-grasp-lift* behavior during execution. While being guided by the human, the robot first detects the 'first-contact' points its fingers make with the objects, and stores the collection of these points as graspable regions, if the object is lifted successfully. Later, it builds up simple classifiers using these experienced contact regions and use these classifiers to detect graspable regions of novel objects. At the end, the robot hand is able to lift an object that may be in different orientations by selecting one of the experienced reach and grasp trajectories.

- In our previous work,⁴⁷ we showed that after learning object manipulation related affordances, the robot can make plans to achieve desired goals and emulate end states of demonstrated actions. In Section 3, we extend this framework and discuss how our robot, by using the learned behaviors and affordance prediction mechanisms, can go beyond simple goal-level imitation and become a better imitator. For this, we develop mechanisms to enable the robot to recognize and segment, with the help of the demonstrator, an ongoing action in terms of its affordance based perception. Once the subgoals are obtained as perceptual states at the end of each segment, the robot imitates the observed action by chaining these sub-goals and satisfying them sequentially. In the experiments, we showed that the robot can better understand and imitate observed complex action executions if the demonstrator modifies his actions by inserting pauses in important points, akin to parents who make similar modifications in infant-directed actions.

2. Scaffolding in Learning Affordances

In this section, we report on parental scaffolding being used to improve the performance of a *reach-grasp-lift* action that was assumed to be acquired in a previous developmental stage. When an object is perceived visually, a reach trajectory towards the center of the object is computed and executed; and in case of any hand-object contact, fingers are flexed in an attempt to grasp the object. However, having no initial knowledge about how to grasp objects, the robot always acts towards the object center. Object graspability on the other hand, depends on the shape and orientation of the objects. This is where a caregiver can interfere and modify the *reach to object center* execution, bringing the robot hand to an appropriate position and orientation with respect to the object, in order to ensure successful grasping. In other words, the robot's default *reach-grasp-lift* trajectory is modified in a semi-kinesthetic way by applying force to the robot arm. The final successful grasp trajectory is a combination of the robot's planned default trajectory and the human's online force-based modification. Although the center of the object was chosen as the target for the default grasp behavior (which is plausible based on³), this is not critical for the operation of the system as the human caregiver modifies the default reach during the execution.

The *first-contact points*, which are the points on the object where the robot fingers made initial touch, are informative about the object graspability; therefore these *first-contact points* are recorded and stored as potential grasp-affording parts. The robot can automatically assess the success of the modified grasp action by lifting the hand after closing it, and then visually checking the table surface to see if the object has been successfully lifted. The robot executes this modified grasp action with different object orientations and learns a model that allows to detect graspable parts of even novel objects.

We defined an affordance as "an acquired relation between a certain effect and a (entity, behavior) tuple, such that when the agent applies the behavior on the entity, the effect is generated".¹¹ In this study, affordance learning is realized by encoding the effect as *grasped*, the behavior as the *reach-grasp-lift* action, and the entity as the *local distance distribution around any object point*. Thus, using the guided interaction experience, the robot will learn grasp affordance by acquiring the ability to predict if a point on the object affords graspability given its local distance distribution. The *first-contact points* are used as positive samples and not-contacted points are used as negative samples during the training. In the following sections, we provide the details of the robot platform, perception, force control feedback and learning algorithms.

2.1. Robot platform

The aim of the experiments is to provide an environment where a human caregiver naturally interacts with the robot, similar to parents interacting with their children. Additionally, the robot should acquire reaching and grasping capabilities. Therefore we use an anthropomorphic hand-arm robot system,

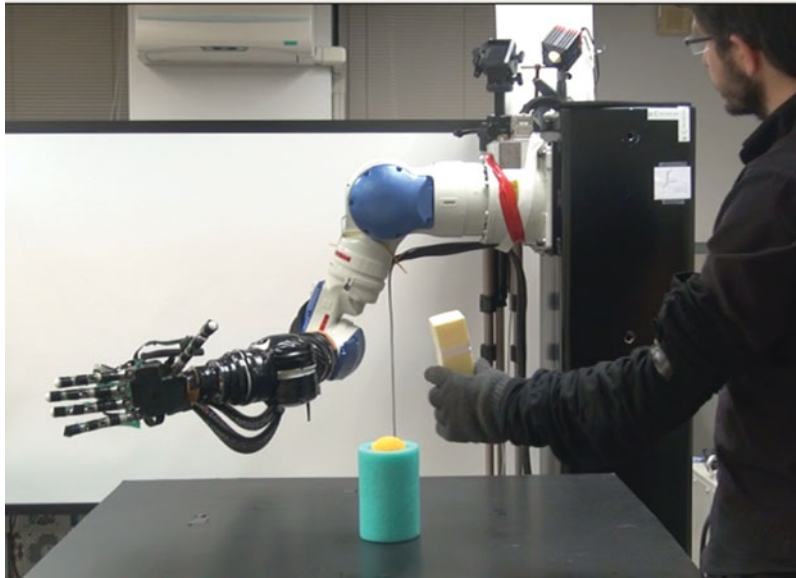


Fig. 2. The robot hand, arm, range camera (top-right), some objects and a tutor constitute the experimental setup. The tutor can change the default trajectory of the robot to enable grasping via the force/torque sensor based control.

where the arm is placed on a vertical stand, similar to a human arm extending from the torso (Fig. 2). The system is composed of a 7 DoF Motoman arm for reaching, and a five fingered 16 DoF Gifu robot hand for object manipulation, with lengths of 123 cm and 23 cm, respectively. In order to sense the force applied by the human caregiver and modify the robot's action execution accordingly, we attached a 6 DoF Nitta force/torque sensor in between the robot hand and the arm, at the wrist position.

For environment perception, an infrared time-of-flight range camera (SwissRanger SR-4000), with 176×144 pixel array, 0.23° angular resolution and 1 cm distance accuracy is used. The range camera is calibrated in the coordinate space of the robot hand by computing the transformation matrix between positions extracted kinematically and perceived from the range camera. The environment is represented as a point cloud in this space.

2.2. Robot perception

The robot uses its range camera to detect the objects on the table. The robot's workspace consists of a black table, a human demonstrator's arm and hand, the robot's own actuator and non-black objects. The demonstrator and the robot body parts are also covered with black material to distinguish the objects easily. The region of interest is defined as the volume over the table, and black pixels are filtered out as the range readings from black surfaces are noisy. As a result, the remaining pixels of the range image are taken as belonging to one or more objects. Since the range camera does not provide reliable color information, the objects are segmented using Connected Component Labeling algorithm¹⁹ based on the depth information. The algorithm differentiates object regions that are spatially separated by a preset threshold value (2 cm in the current implementation). In order to reduce the effect of camera noise, the pixels at the boundary of the object are removed, and median and Gaussian filters with 5×5 window sizes are applied. Finally, various features, such as 3D position of the object center or local distance histogram are computed using the 3D point cloud of the object.

We need to compute the distance between the fingers and the object in order to initiate grasping when a touch occurs. This distance calculation is also necessary to find graspable object regions by marking the object pixels that are first contacted. Finger positions are computed using forward-kinematics based on the robot arm and hand angles. Pixel positions on the other hand are obtained from the object point cloud using the calibrated range camera.

2.3. Force-based human-robot interaction

The force/torque sensor attached to the wrist of the robot outputs two vectors, $\vec{f}_{app}$ and $\vec{m}_{app}$, the current applied force and moment vectors, continuously. Initially, the force vector $\vec{f}_{app}$ in the robot's end-effector frame is mapped to the force vector $\vec{f}_{glob}$ in the global frame. Once mapped, $\Delta \vec{p} \propto \vec{f}_{glob}$ and $\Delta \theta \propto \vec{m}_{app}$ give the desired changes in the position and orientation of the robot hand. Then, the corresponding necessary updates to the joint angles are calculated via inverse kinematics. These desired joint displacements are inserted in robot control loop that may be executing reach-to-grasp action.

Since the force/torque sensor is attached to the wrist, it reads non-negligible measurements due to the hand's inertia, even when there is no human force acting on it. The forces depend on hand orientation and acceleration, which can be used to calculate this systematic errors. However, in our case, there was also significant interference from heavy cables connected to the hand, with resultant complicated dynamical effects, therefore we preferred a learning approach. We collected data by systematically sampling a range of orientations in static configurations. Then, by applying least squares regression on this data, we obtained predictors for every configuration point, which were later used for compensating for the hand's interference.

2.4. Distance histogram based classifier

For learning graspable parts (object handles in our experiments), we need a feature that captures the relative property of the grasped-part with regard to the totality of the object. Inspired from the shape and topological feature detectors in primate brain,²⁷ we define a metric that captures the distribution of three dimensional points (i.e. voxels) that make up a given object. We propose that each voxel is identified by the distribution of its distances from the neighboring voxels that make up the object. This distribution changes smoothly as one moves smoothly on the surface of the object, and is invariant with respect to orientation changes. Our idea was to develop a classifier based on this metric, with the intuition that the handle voxels would have similar distance distributions that are significantly different from the body voxel distributions. The handle voxels that correspond to the automatically detected *first-contact points* are used to construct the handle distance distribution (p_H), and the remaining (not-contacted) voxels are used to compute the body distance distribution (p_B) as follows.

Let H and B be the sets of voxels from the handle and the body of the object, respectively, at a given interaction; and let Z be the union of H and B . We define a δ neighbor distance function to operate on a voxel $\mathbf{x}$ (three dimensional vector) and a set of voxels, Y with

$$\|\mathbf{x}, Y\|_\delta = \{\|\mathbf{x} - \mathbf{y}\| : \mathbf{y} \in Y \wedge \|\mathbf{x} - \mathbf{y}\| < \delta\}$$

Here δ is a constant that defines a voxel neighborhood to span at least the gap between the handle and the body. In the experiments, δ is set to cover the point cloud representation of the objects used. With this we can compactly define these two distance sets for the handle set and the body set as:

$$\Omega_H = \{\|\mathbf{h}_i, Z\|_\delta : \mathbf{h}_i \in H\} \quad (1)$$

and

$$\Omega_B = \{\|\mathbf{b}_i, Z\|_\delta : \mathbf{b}_i \in B\}$$

Taking *hist()* as an operator (with bin size of 20) to return a normalized histogram of a given set we obtain the two probability density estimates for the handle and body voxels:

$$p_H \sim \text{hist}(\Omega_{H_1} \cup \Omega_{H_1} \dots \Omega_{H_N}) \quad (2)$$

and

$$p_B \sim \text{hist}(\Omega_{B_1} \cup \Omega_{B_1} \dots \Omega_{B_N})$$

where N is the number of interactions.

Later, when the robot faces a novel object, it computes a distance distribution for each point on the object, compares them with p_H and p_B , and decides which points can be used as handles.

To determine whether a given voxel v belongs to the handle or not, we compute the distance distribution $p_{\{v\}} \sim \text{hist}(\Omega_{\{v\}})$, and assess its similarity to p_H and p_B as follows:

$$\text{sim}_H(v) = \max \left(\int_0^s p_H - p_{\{v\}} \right) \quad \text{sim}_B(v) = \max \left(\int_0^s p_B - p_{\{v\}} \right)$$

where the max runs over $s < 1$. The voxel v is marked as a handle voxel if

$$\text{sim}_H(v) - \text{sim}_B(v) < \tau \quad (3)$$

otherwise it is deemed as a body voxel. The τ parameter is used to shift the decision boundary and minimize false positives.

2.5. Grasping experiments

2.5.1. Detecting touch regions in guided grasps. The robot begins with an initial rudimentary reach and grasp primitive: On perceiving an object, it automatically tries to reach for its center. If the reach results in object contact, the robot will close its fingers in a grasp attempt. The human caregiver intervenes here: During the motion s/he holds the robot arm, and by applying force, corrects the automated trajectory. Typically the aim of this interference is to correct the grasping target, since humans intuitively try to grasp an object by its handle rather than its geometric center. Yet, this is still a collaborative effort, since the caregiver tends to correct the robot's trajectory minimally, and only when necessary: If s/he believes the robot's movement is optimal, s/he may not interfere at all. Also critical is his/her understanding of the robot's control system and the properties of the robot hand. S/he will develop a sense of the physical system in time, resulting in more effective movements. Once the grasp is achieved, the robot proceeds to identify the contact points on the object, utilizing its camera and proprioceptive sensors. Using the distance histogram based classifier explained above, "handle-detectors" are then built on top of this information, so that the robot can perform efficient grasps on novel objects without guidance in the future.

The robot used 6 successful grasps guided towards the handle of the object placed in different orientations to train the "handle-detectors". The positive samples are obtained from the *first contact points* which correspond to parts of the handle, whereas the untouched object body points provide the negative samples. This approach is different from self-exploratory affordance learning where the robot's exploratory interactions were labeled as success or failure depending on the grasped/not-grasped effect, and graspability was predicted based on the global shape and size features.⁴⁷

2.5.2. Generalizing grasp detection. In this section, we verify the generalization performance of the "handle-detectors" as follows. The first-contact points, which are used as the ground truth in training the "handle-detectors", do not cover all the graspable regions of the object as fingers may not contact with the complete handle. This can be seen in the top row of Fig. 3, where illustrated first-contact points correspond to only a small part of the handles. We first test whether the trained "handle-detectors" can detect the graspable regions that were not marked as the handle during training. The middle row of Fig. 3 shows the distance histograms that are used for building handle classifiers. As shown in the bottom row of Fig. 3, most of the handle voxels can be classified correctly, with few false positives. Note that the tutor guided the hand of the robot only towards the object handles, thus the robot learned that only handle-like structures afford graspability. However, by guiding the robot hand towards other parts of other objects, various grasp affordances could have been learned. Next, we proceed to predicting handle points of completely novel objects. Six novel objects are utilized for this experiment (Fig. 4(a)). Figure 4(b) shows that, without any optimization, despite the occasional false matches, most of the predicted voxels indeed either belong to designated handles, or regions that can be used as a handle. Furthermore, with a hand tuned threshold parameter ($\tau = 0.13$), false positives can be minimized while still extracting grasp points (see Fig. 4(c)).

2.5.3. Executing grasps autonomously. The main focus of this experiment is to identify candidate grasping points of objects using the parental scaffolding framework. Therefore, once the target

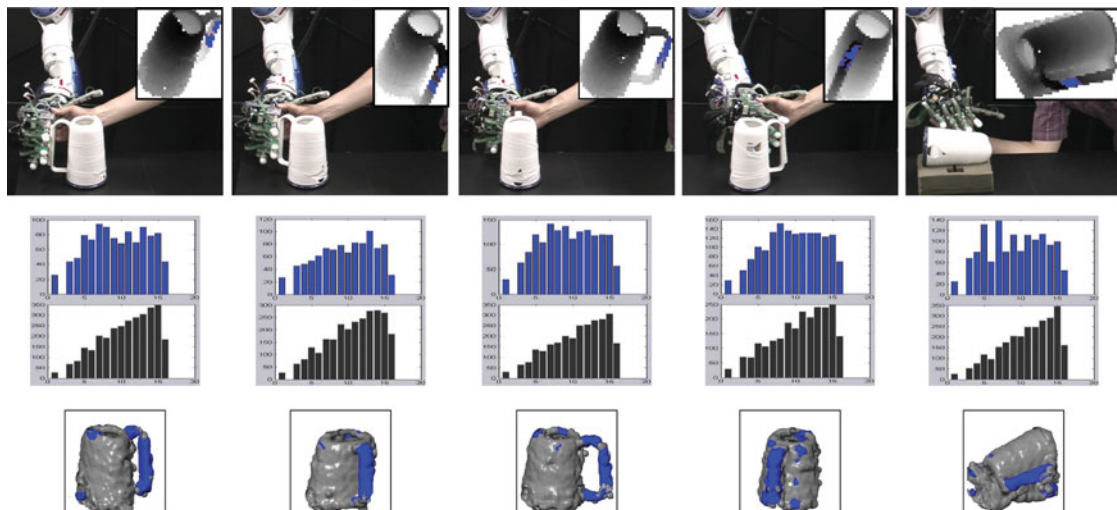


Fig. 3. Top row: Guided grasp experience. The first-contacts are overlaid on the snapshots of the guided grasp executions. Middle row: Distance histograms of the guided grasp executions. Blue and gray histograms correspond to distance distributions of the graspable and the remaining pixels, respectively. Bottom row: The results obtained from grasp classification of each voxel on the training objects.

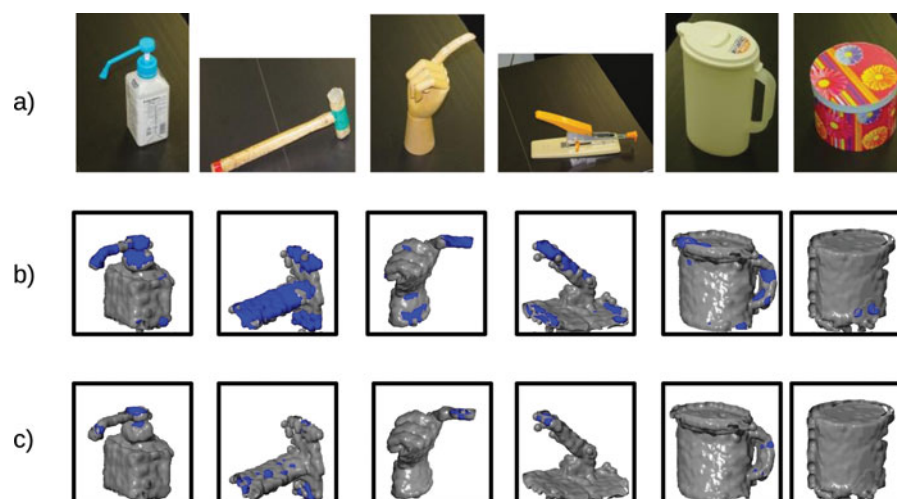


Fig. 4. Grasp classification is performed for the novel objects shown in (a). The results of the grasp region detection are shown (b) and (c) where the τ threshold is adjusted to minimize the false positives in the latter case.

grasping location is calculated, exactly how to execute these grasps is not of primary interest. Yet, it is also possible to extend the learned information in order to enable an automated control of the arm for reaching and lifting the object. We design a simple lookup-table based mechanism to select a *reach-grasp-lift* execution trajectory, based on the handle orientation with respect to the object center.

During training, we store three kinds of information for each object: (1) the set of object voxels, (2) the set of touch voxels, and (3) the hand-arm angle trajectory (the trajectory that is guided by the caregiver). We then compute the position of the largest touch region relative to the object center. The lookup-table is constructed such that it will take this relative position as input, and return the associated hand-arm trajectory as learned from this experience. After training, when the robot perceives an object, it first predicts the grasp regions as detailed above, then computes the relative position of the largest grasp region with respect to the object center, and uses this relative position to retrieve a previously learned hand-arm trajectory from the lookup table.

We demonstrate the efficiency of this simple approach with the mug shaped object placed in 5 different orientations. Four of these executions proved to be successful, since the object was placed



Fig. 5. Robot grasps objects using the trajectories learned during scaffolding. Each row corresponds to a different grasp execution.

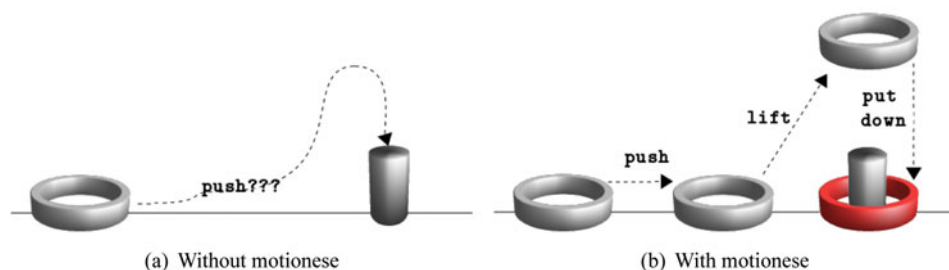


Fig. 6. An example scenario where demonstration of enclosing a cylinder with a ring is continuous in (a) and exaggerated in (b). The goal configuration of the ring is indicated by the red colored ring enclosing the cylinder. In (a), if the robot cannot capture the important features of the demonstration, it may attempt to bring the ring to the goal position by simply pushing it to the right, where the ring will push the cylinder away rather than enclosing it. On the other hand, when important steps are highlighted by for example pauses as in (b), the robot can extract sub-goals represented in its perceptual space and find a behavior sequence from its behavior repertoire to imitate the action correctly.

in similar orientations with the training instances. The final execution, which required reaching for a handle behind the object that the robot did not encounter before, failed as expected. Figure 5 depicts snapshots of the trajectories of successful grasp executions. (One of the successful grasps was very similar to the first one so it is excluded to save space.)

3. Scaffolding in Imitation Learning

In this section, we explore how scaffolding can be used to transform our goal-emulation-able robot to a better imitator as a given demonstration may not be replicated with simple goal emulation. In fact learning higher level skills based on previously learned simpler ones is more economical and usually easier for building a complex sensorimotor system.²² Therefore, here we propose to use sequential goal emulation (that is based on learned affordance prediction in our framework) to achieve complex imitation.

Extracting sub-goals or important features from a demonstration is not straightforward, as demonstrated action trajectory may not correspond to any robot behavior developed so far. For example when the robot is asked to imitate a demonstration shown in Fig. 6(a), as the observed trajectory is not represented in the robot's sensorimotor space, executing the behavior that seemingly achieves the goal would not satisfy the imitation criteria (for example a right push of the ring would tip over the cylinder). Young infants also have similar difficulties in mapping certain observed actions to their own action repertoire, so they fail imitating these actions³⁵ (also see⁵⁰ for neurophysiological

evidence on how infants' own action repertoire affects the understanding of observed actions of others.)

As explained in the Introduction Section, to overcome this difficulty, parents are known to make modifications in infant-directed actions, i.e. use "motionese". Our fine-grained analysis using a computational attention model also reveals the role of motionese in action learning.²⁸ Longer pauses before and after the action demonstration usually underline the initial and final states of the action (i.e. the goal of the action) whereas shorter but more frequent pauses between movements highlight the sub-goals.

Inspired from this, we implement perception mechanisms in our robot so that it can make use of motionese (when available) to identify important steps and boundaries in the otherwise complex stream of motion. A human tutor can exaggerate the relevant features in his demonstration as in Fig. 6(b), and enable the robot to map the movement segments into its own behavior repertoire, and imitate the action sequence successfully.

3.1. Affordances and effect prediction

In our previous work,^{47,49} the affordances were defined as (object, behavior, effect) relations, and with this we have shown that affordance relations can be learned through interaction without any supervision. As learning the prediction ability is not focus of this paper, we skip the details and shortly present how prediction operator works. This prediction operator can predict the continuous change in features given object feature vector (f^0), behavior type (b_j) and behavior parameters (ρ_f):

$$(f^0, b_j, \rho_f) \rightarrow f_{\text{effect}}^{b_j} \quad (4)$$

where $f_{\text{effect}}^{b_j}$ denotes the effect predicted to be observed after execution of behavior b_j .

3.2. State transition

The state corresponds to the list of feature vectors obtained from the objects in the environment:

$$S_0 = [f_{o_0}^0, f_{o_1}^0, \dots, f_{o_m}^0]$$

where () denotes the zero length behavior sequence executed on the objects, and m is the maximum number of objects. If the actual number of objects is less than m , the *visibility* features of non-existing objects are set to 0.

State transition occurs when the robot executes one of its behaviors on an object. Only one object is assumed to be affected at a time during the execution of a single behavior, i.e. only the features of the corresponding object is changed during a state transition. Thus, the next state can be predicted for any behavior using the prediction scheme given in Eq. (4) as follows:

$$S'_{t+1} = S_t + [\dots 0, f_o^{b_j \text{ effect}}, 0, \dots] \quad (5)$$

where behavior b_j is executed on object o and features of this object change by the summation operator.

Using an iterative search in the behavior parameter space, the robot can also find the best behavior (b^*) and its parameters that is predicted to generate a desired (des) effect given any object:

$$b^*(f^0, f_{\text{effect}}^{\text{des}}) = \arg \min_{b_j, \rho_f} (f_{\text{effect}}^{\text{des}} - f_{\text{effect}}^{b_j}) \quad (6)$$

3.3. Goal-emulation

Goal-emulation refers to achieving a goal represented as desired world state ($S^* = [f_{o_1}, f_{o_2}] \dots$) by generating a plan. The plan consists of a behavior sequence ($b_0, b_1, \dots$) that is predicted to transform the given state into the goal state ($S_0 \xrightarrow{b_0} S_1 \xrightarrow{b_1} S_2 \dots \rightarrow S^*$). Because prediction is based on vector summation Eq. (5), the robot can estimate the total effect that a sequence of behaviors will create by simply summing up all the effect vectors, and thus can use this for multi-step prediction.

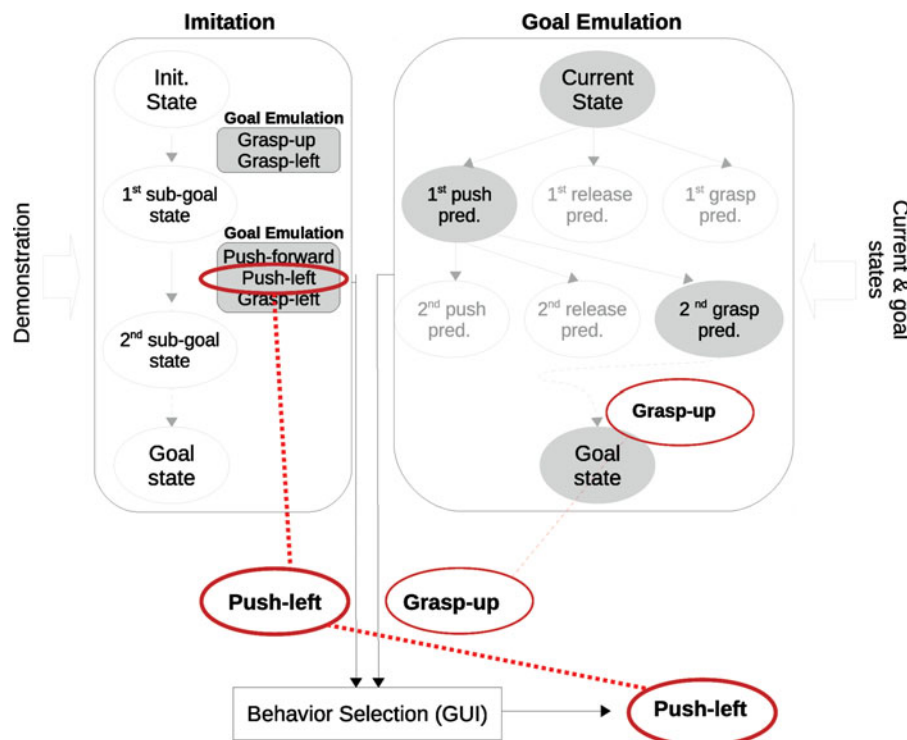


Fig. 7. The robot can choose to follow the demonstration by executing the behaviors from Imitation module (left panel) or can find the sequence of behavior using Goal Emulation module (right panel) in any step to reach to the desired state. Note that Imitation Module also uses Goal Emulation Module inside to find behavior sequence between each sub-goal state pair. In this example, *push-left* and *grasp-up* behaviors are selected by the Imitation and Goal-Emulation modules, respectively; and *push-left* is selected by the user through Graphical User Interface.

Given a goal state, the robot can use any state space search method in finding this behavior sequence. In this paper, we use forward chaining which uses a tree structure with nodes holding the perceptual states and edges connecting states to the predicted next states. Each edge corresponds to the prediction of execution of a different behavior on a different object, and transfers the state to a different state based on Eq. (5). Starting from the initial state encoded in the root node, the next states for different behavior-object pairs are predicted for each state. A simple example where forward chaining is applied to achieve the goal state is provided in Fig. 7 right panel.

3.4. Imitation

Imitation (as opposed to goal-emulation where demonstration details are not important) refers to finding the behavior sequence that enables the robot to follow a similar trajectory with the demonstration.⁴⁴ As our focus is object manipulation, the robot needs to detect the important subgoals in the demonstration which can be used to generate a similar object trajectory by achieving those subgoals successively in discrete steps.

The robot observes the demonstration and extracts the initial and goal states, as well as the intermediate states (encoded as sub-goals) by detecting pauses which may be introduced by a motionese engaged tutor. If no pause can be detected, then a random intermediate state will be picked up as the sub-goal state, which may lead to a failed imitation attempt¹.

Imitation Module (Fig. 7, left panel) finds the behavior sequence that brings the initial state (S_0) to the goal state (S^*) following the detected sub-goal state sequence. Assuming three pauses were

¹ The aim of introducing an intermediate state is to provide feedback to the tutor about the failed imitation attempt of the robot, and provide some indirect hint about the observation mechanisms of the robot. In the current implementation, this state is selected randomly. However the selection of a sub-state that reflects the current sensorimotor capabilities of the robot, and propagating this selection to the tutor will have a significant impact on the tutor's comprehension of the robot skills.

detected along with their states (S_0^* , S_1^* , S_2^*), Imitation Module needs to find four behavior sequences that successively transfer the initial state to the goal state through the sub-goal states:

$$S_0 \xrightarrow{\text{beh-seq-0}} S_0^* \xrightarrow{\text{beh-seq-1}} S_1^* \xrightarrow{\text{beh-seq-2}} S_2^* \xrightarrow{\text{beh-seq-3}} S^* \quad (7)$$

In the experiments reported in this article, each arrow happened to correspond to a single-affordance perception, i.e. each state transition was achieved by one behavior. However, our framework is not limited to this, and in fact it can form multi-step plans for reaching individual sub-goals.⁴⁷

3.5. Imitation experiments

3.5.1. Behaviors. The robot interacts with the objects using three behaviors, namely *grasp*, *release*, and *push* which were learned before.⁴⁹ The focus of this experiment is the realization of the imitation rather than a complete verification of the previously learned behaviors using a simple grasp scenario. For this purpose, we used simple convex objects without handles which could be grasped using the transferred *grasp* behavior.

For any behavior, *initial*, *target* and *final* hand positions are computed as offsets from the object center, and a trajectory that passes through these via points are executed as follows:

- *Initial* position is the offset from the object where the robot places its hand prior to interaction. This parameter is fixed and same for all behaviors, and places the robot hand in between the object and the robot torso towards right in lateral direction. If the object to be interacted is already in the robot's hand, the initial position is set as the current position of the robot hand as there is no need to re-position the hand.
- *Target* is the offset from the object-center that determines which part of the robot's hand makes contact with the object. Using this parameter which is fixed and previously learned for each behavior, the robot touches to the object with its palm in *grasp* and *release* behaviors, and with its fingers in *push* behavior.
- *Final* position is the offset from the object where the robot brings its hand at the end of the behavior execution. Final position can be set to any arbitrary point in the robot's workspace unless it is too close to the table surface proximity to avoid any collision.
- *Hand-close* and *hand-open* positions are again given as offsets from the object. The hand encloses into a fist with *grasp* and *release* behaviors when it is close to the object center, and the fingers fully extend with *release* behavior at the end of action execution. *Push* behavior does not change hand-state unless the object is already in the robot's hand. In this case the fingers fully extend in the beginning of the hand movement.

With the appropriate ranges of action parameters, target objects can be grasped, released or pushed to different locations depending on the *final* position. As the *final* position appears to be the primary determinant of the effect of an applied action, it is used in searching for the best behavior parameter (ρ_f) to generate a desired effect (Eq. (6)) and imitate the demonstration (Eq. (7)).

3.5.2. Control architecture. From a robotic architectural perspective, imitation and goal-emulation based control can co-exist and work complementarily. Selecting behaviors based on Imitation Module (see Fig. 7) results in following the exact trajectory of the demonstrator, to the extent that it is decimated by the pauses inserted by the tutor. On the other hand, when Goal Emulation Module (right panel) is selected, then a behavior sequence that brings the current state to the goal state is found independent of the intermediate states. In the current implementation, we followed a simple approach where the Behavior Selection Module simply reflects the choice of the experimenter that is conveyed through a Graphical User Interface. For an autonomous robotic agent, an arbitration mechanism to choose an appropriate module depending on the context and demonstration needs to be implemented. Although, work is underway in this direction, we leave this topic for an architecture focused publication as it is out of the scope of the current paper.

3.5.3. Imitation performance. This section presents the experimental results of the developed imitation system. In this system the robot uses the discovered behavior primitives and learned affordances to imitate demonstrations which include deliberate motions. As described in Section 3.4, the robot observes the demonstration and extracts the initial and goal states, as well

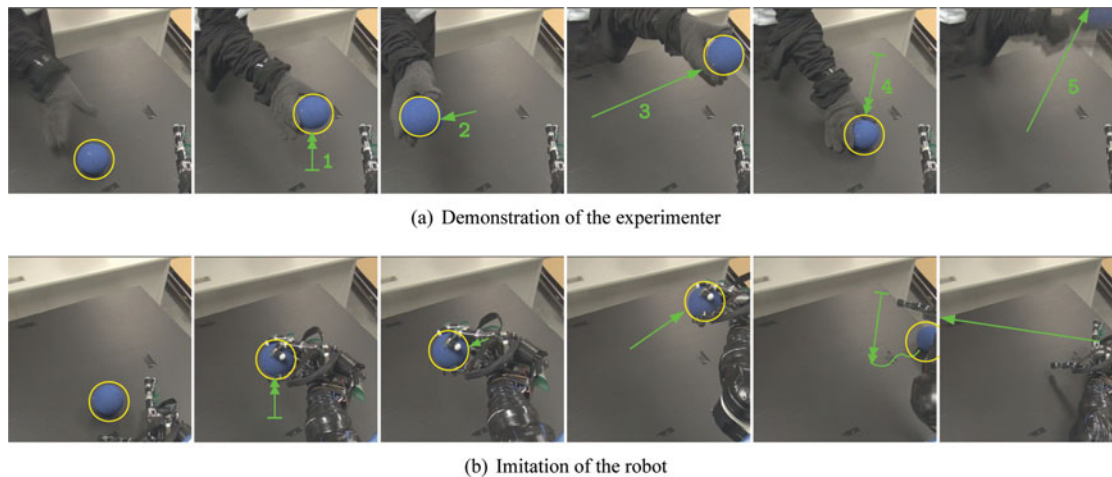


Fig. 8. Snapshots where the robot imitates the demonstration performed by an experienced tutor. In (b) from left to right, grasp-up, grasp-left, grasp-right, release, and push actions are planned and executed. The object paths are shown with arrows.

as the intermediate states (encoded as sub-goals) by detecting pauses which may be introduced by a motionese-engaged tutor.

We tested the Imitation Module expecting the robot to follow the demonstration of a subject who deliberately inserts pauses into his trajectory in the ‘right’ moments. The demonstration was performed by an experienced tutor who knows the working principles of the system. The aim of this experiment was to validate the robot’s ability to find sub-goals which include lift-object, move-object-around or disappear-object-from-the-view by using affordance prediction capability for graspability, rollability, etc.

To test the imitation capability of the robot, the tutor was asked to grasp the object on the table and move it in space as shown in Fig. 8. The tutor used his own action repertoire to move the object around, while the robot monitored the object, detected the pauses, and encoded these pauses as intermediate states. Following the demonstration, the robot found the sequence of behaviors to obtain the intermediate states in sequence using the imitation procedure summarized in Fig. 7 and the behavior selection procedure given in Eq. (6).

The human’s demonstration and imitation of the robot are shown in Fig. 8. From the robot’s point of view, the object was (1) lifted up, (2) moved to left, (3) moved to right in the air, (4) put on the table, and (5) removed from the table. The robot, after detecting the intermediate states and encoding them as sub-goals, computed a behavior sequence provided in Table I. As mentioned before, the behaviors used by the robot were autonomously discovered in the previous developmental stage, labeled as *grasp*, *release* and *push* to improve readability of the text, and not used by the robot. Behavior parameters (ρ_f) in this experiment correspond to *final* position that was explained in Section 3.5.1. As the object was in the robot’s hand prior to behavior executions for 2nd, 3rd and 4th steps, the hand was not re-positioned to *initial* position.

Execution of these behaviors resulted in a trajectory that is similar to the demonstration with the following exceptions: First of all, the object was not moved exactly to the same positions in intermediate steps because of the noise in the perception and due to kinematic constraints of the robot. Second, after the object was released, it rolled over the table (to the right) and did not end up exactly below the hand as the robot predicted. Third, push behavior could not roll the object off the table without (the experimenter) bending the table as the push was not strong enough to cause a high-speed roll. Still, the robot was able to accomplish the task by achieving the observed sub-goal changes using its own behaviors. For example, the tutor’s ‘removal’ of the object from the table (by bringing it outside camera view) was mapped to *push* behavior of the robot as the object was rollable, and *push* applied to a rollable object was predicted to make the object disappear. As another example, ‘putting on table’ action performed by the tutor was mapped to the robot’s *release* behavior which had similar effect.

Table I. Sequence of behaviors that Imitation Module generated to imitate the demonstration shown in Fig. 8(a). The behavior primitives were discovered in the previous behavior formation phase and each id was assigned automatically. Step corresponds to the order in the sequence. Labels are provided for readers' better understanding of the behaviors. The first, second and third parameters correspond to movement in lateral, vertical and frontal axis with respect to the robot body, after object contact, respectively. The last column includes explanations of the parameters with large effects only.

Step	Beh. Id	Beh. Label	Beh. Param.	Explanation
1	0	<i>grasp</i>	(−3, +21, −1)	Grasped the object and lifted up 21 cm.
2	0	<i>grasp</i>	(+9, +13, −3)	Lifted it up further 13 cm and moved 9 cm to the left.
3	0	<i>grasp</i>	(−21, +5, +4)	Moved object 21 cm to the right.
4	2	<i>release</i>	(+4, 0, +4)	Released the object.
5	1	<i>push</i>	(+10, +4, 0)	Pushed the object towards right for 10 cm.

This experiment shows that using discovered behaviors and learned affordances enabled the robot to seamlessly embody the observation in its sensorimotor space and flexibly imitate the demonstration based on the learned relations between own capabilities and the object properties. If the object was a non-rollable one, then the robot would probably push it several times in order to drop it from the edge of the table as shown by Ugur *et al.*⁴⁸

3.5.4. Imitation versus goal-emulation. In the Introduction Section, we discussed that understanding “what to imitate” is a nontrivial problem both in infants and robots; and is an active research in robotics and developmental science. It is not clear whether humans have multiple mechanisms for imitation or how they modulate or arbitrate their imitation modes. From a robotics point of view, hosting multiple mechanisms (modules) may be necessary. However, employing multiple mechanisms (modules) to imitate an observed action has its own advantages and disadvantages. If the tutor has engaged in a motionese based interaction with the robot, and provides sufficient cues to the robot, the Imitation Module makes complex imitation possible. However this requires keeping all sub-goals in the memory of the cognitive system of the robot, and executing all corresponding actions, which might not be practical if the tutor makes a large number of pauses. Furthermore, the Imitation Module needs additional mechanisms to deal with failures during execution, and to take corrective actions. On the other hand, Goal-emulation needs no additional mechanisms as it can automatically recover by simply reassessing the current state and re-planning. However it may fail in complex environments as predictions are made based on basic object affordances. An autonomous mechanism that is guided by a caregiver's signals should be formalized and implemented based on insights obtained from developmental psychology.

4. Conclusion

In this paper, we discussed how we can extend our previously developed unsupervised affordance learning framework so that the robot can benefit from parental scaffolding for affordance learning and action imitation. For the former, we used physical parental scaffolding to shape the premature reach and grasp attempts of the robot, where the caregiver intervened the ongoing executions to transform them into successful grasps. This resulted in (1) the discovery of new affordances, i.e. handles and (2) the formation of new motor skills, i.e. handle-directed grasping. For the latter, we utilized the capability of goal-emulation, which had been obtained in the previous developmental stages, to replicate a demonstrated action. An important point here is that once a demonstrated action can be segmented as sub-goals in the perceptual space of the robot, satisfying them sequentially would reproduce the observed action. Therefore, motionese provided cues to the robot for parsing a demonstrated action into performable chunks. Our experiments showed that complex actions that were not directly represented in the sensorimotor space of the robot, and hence could not be simply reproduced by goal-emulation, can be imitated when motionese was adopted by the demonstrators.

The scientific contribution of this work is twofold.

- First, with the addition of imitation with motionese, we have obtained an integrated developmental system that starts off with simple motor acts that are turned into well defined motor behaviors through development, which are then used by the robot to explore its action capabilities. With

this, the robot represents the world in terms of perceptual changes it can induce through its behaviors, and uses this information to make plans to achieve goals. Current work showed how this can be enhanced via scaffolding, and how the imitation skill can be built upon capabilities and structures acquired through development. In particular, we studied two of the five attention-directing gestures,⁵³ namely ‘embody’ and ‘demonstrate’, in our robot-caregiver system. We believe that our system parallels infant development and has the potential to provide insights into mechanisms of mental development in biological systems.

- Second, when looked at from a robotics perspective we have developed a learning robotic platform that has the ability to self-learn object affordances and actions applicable to them, and can make use of human teaching by interpreting the demonstrations in terms of the structures developed in its perceptual space. This is the first robotics work where motionese based robot learning and control is realized in a truly developmental setting. The system can sustain both goal-emulation and imitation guided by motionese when it is available.

5. Discussion

The two robotics problems that were tackled in this paper, namely grasp learning and learning from demonstration, have been intensively studied in recent years. Our main focus was to use parental scaffolding as a bootstrapping mechanism from a developmental perspective, yet our approach can be discussed in relation to other robotic methods that deal with these two problems. Regarding learning from demonstration studies, the focus is generally on how to represent the demonstrated actions in a generic way, and how to generalize them in new contexts. For example, Calinon *et al.*⁹ used Principal Component Analysis and Gaussian Mixture models to find the relevant features of the actions, and to encode them in a compact way, respectively. Recently, Dynamical Motor Primitives (DMP) framework⁴² has been shown to be very effective in learning of discrete³² and periodic movements,¹⁶ and in generalizing multiple movements to new situations.⁴⁵ Pastor *et al.*³⁴ showed that the sensory feedback obtained from object-robot interactions can be used to adapt DMP-based movements in environments with uncertainties. We believe that such powerful motion representation frameworks are necessary in imitation of complex actions where the focus is on replicating end-effector or joint trajectories respecting kinematic constraints and velocity profiles. Our object-based imitation, which uses learned effect prediction and sub-goal emulation mechanisms, different from the above studies, resides between full trajectory level imitation and goal emulation; and should be coupled with such approaches in order to exhibit full-range of imitation capabilities. Even in the context of goal-emulation based imitation, such a capability becomes crucial in case the tutors are not cooperative in adapting their demonstrations. In such cases, the robot can use a combination of cues from end-effectors and object movements; or can fall back to self-learning mode (e.g. by using Reinforcement Learning³⁶).

Regarding grasp learning, various methods that use force analysis,⁵ local image features,⁴⁰ global object models,¹³ or object parts²⁰ have been developed and applied in real-world settings successfully. Particularly, Saxena *et al.*⁴⁰ trained the robot on synthetic 2D images marking good grasping points. After the training, the robot could identify good grasping points on a novel object, given only its 2D images. Once the candidate grasping points were identified in 2D, several images were combined to triangulate the points and estimate their 3D positions. Detry *et al.*¹² represented grasping affordances for a gripper as a continuous distribution in a 6D pose space where the distribution was initialized by predefined visual descriptors or sample grips executed by a human (not the robot), and refined through the robot’s self-exploration. While these works aim to provide the robot with initially perfect or almost-perfect information, our focus has been the online guidance of a human teacher to correct the robot’s naive movements on-the-fly. Although the results demonstrate that the robot was able to efficiently learn graspable regions through parental scaffolding, and was able to detect graspable parts of novel objects, there is room for improvement which we discuss in the next section.

6. Future Work

In this study, we focused on parental scaffolding from the perspective of the robot. As in infant-parent interaction, we required guidance of the tutor (i.e. scaffolding) to enable the robot to acquire manipulation skills. For this, the tutor’s understanding of the perception and action capabilities of

the robot is critical. Generalizing from infant-parent interactions, we expect that even tutors without much knowledge about the capabilities of the robot will be able to “learn how the robot learns” by observing the failed attempts of the robot. In the current work, in terms of “embodied scaffolding”, the naive tutors need to adapt to the force-based human-robot interaction setup, and learn the grasp learning mechanism of the robot. Regarding imitation learning, naive subjects need to understand the imitation capability of the robot, and act accordingly. Indeed, our preliminary experiments with naive demonstrators, though not reported here, indicate that the system engages the naive tutors to use motionese perceivable to the robot until the demonstrated action is imitated by the robot accurately. In these experiments, the only observable feedback to the naive tutors was the failed imitation attempts of the robot. However, as shown by Fischer *et al.*¹⁵, how the tutors adapt their teaching behavior for the requirements of the learning robot largely depends on the type of the feedback obtained from the system. Thus, suitable feedback mechanisms are crucial to achieve more effective human-robot learning systems. How this feedback can be learned by the robot or encoded in the parental scaffolding framework are interesting open research questions that we plan to study next.

Parental scaffolding has been used to extend the basic sensorimotor capabilities of the robot in this work. Detecting progressively more complex affordances (such as handle graspability in the first case study and learning of complex actions in the second one) requires better learning methods and a large amount of experience obtained through robot-environment interactions. The complex skills learned through embodied scaffolding and motionese in general require affordance detection and behavior execution on multiple objects as our motivating examples illustrated in Figs. 1 and 6. Our system in its current form cannot deal with such situations, and should be extended to account for multi-object interactions. To extend this framework for a large set of objects, actions and affordances, we believe that using previously learned affordances, which partially capture robot-object dynamics, will provide bootstrapping in learning more complex affordances, reducing the exploration and training time significantly.

Acknowledgments

This research was partially supported by European Community’s Seventh Framework Programme FP7/2007-2013 (Specific Programme Cooperation, Theme 3, Information and Communication Technologies) under grant agreement no. 270273, Xperience; by European Community’s Seventh Framework Programme FP7/2007-2013 under the grant agreement no. 321700, Converge; and by JSPS/MEXT Grants-in-Aid for Scientific Research (Research Project Number: 24000012, 24119003). It was also partially funded by a contract in H23 with the Ministry of Internal Affairs and Communications, Japan, entitled ‘Novel and innovative R&D making use of brain structures’.

References

1. B. D. Argall, E. L. Sauser and A. G. Billard, “Tactile Guidance for Policy Refinement and Reuse,” *Proceedings of the 9th IEEE International Conference on Development and Learning* (2010) pp. 7–12.
2. J. Babic, J. Hale and E. Oztop, “Human sensorimotor learning for humanoid robot skill synthesis,” *Adapt. Behav.* **19**, 250–263 (2011).
3. T. M. Barrett and A. Needham, “Developmental differences in infants’ use of an object’s shape to grasp it securely,” *Developmental Psychobiology* **50**(1), 97–106 (2008).
4. L. E. Berk and A. Winsler, “Scaffolding Children’s Learning: Vygotsky and Early Childhood Education,” *National Assoc. for Education* (1995).
5. A. Bicchi and V. Kumar, “Robotic Grasping and Contact: A Review,” *Proceedings of the IEEE International Conference on Robotics and Automation* (2000) pp. 348–353.
6. A. Billard, “Learning motor skills by imitation: A biologically inspired robotic model,” *Cybern. Syst.* **32**, 155–193 (2000).
7. R. J. Brand, D. A. Baldwin and L. A. Ashburn, “Evidence for ‘motionese’: Modifications in mothers’ infant-directed action,” *Developmental Sci.* **5**(1), 72–83 (2002).
8. C. Breazeal, Learning by Scaffolding *PhD Thesis*, Elec. Eng. Comp. Sci. (MIT, Cambridge, MA, 1999).
9. S. Calinon, F. Guenter and A. Billard, “On learning, representing, and generalizing a task in a humanoid robot,” *IEEE Trans. Syst. Man Cybern.* **37**(2), 286–298 (2007).
10. M. Carpenter, J. Call and M. Tomasello, “Understanding ‘prior intentions’ enables two-year-olds to imitatively learn a complex task,” *Child Dev.* **73**(5), 1431–1441 (2002).
11. E. Şahin, M. Çakmak, M. R. Doğan, E. Uğur and G. Üçoluk, “To afford or not to afford: A new formalization of affordances toward affordance-based robot control,” *Adapt. Behav.* **15**(4), 447–472 (2007).

12. R. Detry, E. Başeski, M. Popović, Y. Touati, N. Krüger, O. Kroemer, J. Peters and J. Piater, "Learning Continuous Grasp Affordances by Sensorimotor Exploration," *In: From Motor Learning to Interaction Learning in Robots* (Springer, Berlin, 2010) pp. 451–465.
13. R. Detry, D. Kraft, O. Kroemer, L. Bodenhagen, J. Peters, N. Krüger and J. Piater, "Learning grasp affordance densities," *Paladyn*, **2**(1), 1–17 (2011).
14. A. Fernald, "Four-month-old infants prefer to listen to motherese," *Infant Behav. Dev.* **8**, 181–195 (1985).
15. K. Fischer, K. Foth, K. J. Rohlfing and B. Wrede, "Mindful tutors: Linguistic choice and action demonstration in speech to infants and a simulated robot," *Interact. Stud.* **12**(1), 134–161 (2011).
16. A. Gams, A. J. Ijspeert, S. Schaal and J. Lenarcic, "On-line learning and modulation of periodic movements with nonlinear dynamical systems," *Auton. Robots* **27**(1), 3–23 (2009).
17. G. Gergely, H. Bekkering and I. Kiraly, "Rational imitation in preverbal infants," *Nature* **415**, 755 (2002).
18. N. Goubeta, P. Rochat, C. Maire-Leblond and S. Poss, "Learning from others in 9–18-month-old infants," *Infant Child Dev.* **15**, 161–177 (2006).
19. R. M. Haralick and L. G. Shapiro, *Computer and Robot Vision, Volume I* (Addison-Wesley, New York, 1992).
20. A. Herzog, P. Pastor, M. Kalakrishnan, L. Righetti, J. Bohg, T. Asfour and S. Schaal, "Learning of grasp selection based on shape-templates," *Auton. Robots* **36**(1-2), 51–65 (2014).
21. R. M. Hodapp, E. C. Goldfield and C. J. Boyatzis, "The use and effectiveness of maternal scaffolding in mother-infant games," *Child Dev.* **55**(3), 772–781 (1984).
22. M. Kawato and K. Samejima, "Efficient reinforcement learning: Computational theories, neuroscience and robotics," *Curr. Opin. Neurobiology* **17**, 205–212 (2007).
23. E. A. Koterba and J. M. Iverson, "Investigating motionese: The effect of infant-directed action on infants' attention and object exploration," *Infant Behav. Dev.* **32**(4), 437–444 (2009).
24. D. Kushida, M. Nakamura, S. Goto and N. Kyura, "Human direct teaching of industrial articulated robot arms based on force-free control," *Artif. Life Robot.* **5**(1), 26–32 (2001).
25. N. Masakata, "Motherese in a signed language," *Infant Behav. Dev.* **15**, 453–460 (1992).
26. B. Moore and E. Oztot, "Robotic grasping and manipulation through human visuomotor learning," *Robot. Auton. Syst.* **60**, 441–451 (2012).
27. A. Murata, V. Gallese, G. Luppino, M. Kaseda and H. Sakata, "Selectivity for the shape, size, and orientation of objects for grasping in neurons of monkey parietal are AIP," *J. Neurophysiology* **83**(5), 2580–2601, (2000).
28. Y. Nagai and K. J. Rohlfing "Computational analysis of motionese toward scaffolding robot action learning," *IEEE Trans. Auton. Mental Dev.* **1**(1), 44–54 (2009).
29. Y. Nagai, A. Nakatani and M. Asada, "How a Robot's Attention Shapes the Way People Teach," *Proceedings of the 10th International Conference on Epigenetic Robotics* (2010), pp. 81–88.
30. Y. Nagai and K. J. Rohlfing, "Can Motionese Tell Infants and Robots: What to Imitate?," *Proceedings of the 4th International Symposium on Imitation in Animals and Artifacts, AISB* (2007) pp. 299–306.
31. Y. Nagai and K. J. Rohlfing, "Parental Action Modification Highlighting the Goal versus the Means," *Proceedings of the IEEE 7th International Conference on Development and Learning* (2008).
32. J. Nakanishi, J. Morimoto, G. Endo, G. Cheng, S. Schaal and M. Kawato, "Learning from demonstration and adaptation of biped locomotion," *Robot. Auton. Syst.* **47**(2), 79–91 (2004).
33. C. L. Nehaniv and D. Dautenhahn, "Like me? Measures of correspondence and imitation," *Cybern. Syst.* **32**, 11–51 (2011).
34. P. Pastor, L. Righetti, M. Kalakrishnan and S. Schaal, "Online Movement Adaptation Based on Previous Sensor Experiences," *Proceedings of the 2011 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*, IEEE (2011) pp. 365–371.
35. M. Paulus, S. Hunnius, M. Vissers and H. Bekkering, "Imitation in infancy: Rational or motor resonance?," *Child Dev.* **82**(4), 1047–1057 (2011).
36. J. Peters, S. Vijayakumar and S. Schaal, "Reinforcement Learning for Humanoid Robotics," *Proceedings of the Third IEEE-RAS International Conference on Humanoid Robots* (2003) pp. 1–20.
37. K. J. Rohlfing, J. Fritsch, B. Wrede and T. Jungmann, "How Can Multimodal Cues from Child-Directed Interaction Reduce Learning Complexity in Robots?," *Adv. Robot.* **20**(10), 1183–1199 (2006).
38. J. Saunders, C. Nehaniv, K. Dautenhahn and A. Alissandrakis, "Self-imitation and environmental scaffolding for robot teaching," *Int. J. Adv. Robot. Syst.* **4**(1), 109–124 (2007).
39. J. Saunders, C. L. Nehaniv and K. Dautenhahn, "Teaching Robots by Moulding Behavior and Scaffolding the Environment," *Proceedings of the ACM SIGCHI/SIGART Conference on Human-robot Interaction* (2006) pp. 118–125.
40. A. Saxena, J. Driemeyer and A. Y. Ng, "Robotic grasping of novel objects using vision," *Int. J. Robot. Res.* **27**(2), 157–173 (2008).
41. S. Schaal, "Is imitation learning the route to humanoid robots?," *Trends Cogn. Sci.* **3**(6), 233–242 (1999).
42. S. Schaal, "Dynamic Movement Primitives-a Framework for Motor Control in Humans and Humanoid Robotics," *In: Adaptive Motion of Animals and Machines*. (Springer, 2006) pp. 261–280.
43. C. S. Tamis-LeMonda, Y. Kuchirko and L. Tafuro, "From action to interaction: Infant object exploration and mothers' contingent responsiveness (june 2013)," *IEEE Trans. Auton. Mental Dev.* **5**(3) (2013).
44. M. Tomasello, "Do Apes Ape?," *In: Social Learning in Animals: The Roots of Culture* (C. M. Heyes and B. G. Galef, eds.) (Academic Press, Inc., San Diego, CA, 1996) pp. 319–346.

45. A. Ude, A. Gams, T. Asfour and J. Morimoto, "Task-specific generalization of discrete and periodic dynamic movement primitives," *IEEE Trans. Robot.* **26**(5), 800–815 (2010).
46. E. Ugur, H. Celikkanat, Y. Nagai and E. Oztop, "Learning to Grasp with Parental Scaffolding," *Proceedings of the IEEE International Conference on Humanoid Robotics* (2011a) pp. 480–486.
47. E. Ugur, E. Oztop and E. Sahin, "Goal emulation and planning in perceptual space using learned affordances," *Robot. Auton. Syst.* **59**(7–8), 580–595 (2011b).
48. E. Ugur, E. Sahin and E. Oztop, "Affordance Learning from Range Data for Multi-Step Planning," *Proceedings of the 9th International Conference on Epigenetic Robotics* (2009) pp. 177–184.
49. E. Ugur, E. Sahin and E. Oztop, "Self-discovery of Motor Primitives and Learning Grasp Affordances," *IEEE/RSJ International Conference on Intelligent Robots and Systems* (2012) pp. 3260–3267.
50. M. van Elk, H. T. van Schie, S. Hunnius, C. Vesper and H. Bekkering, "You'll never crawl alone: Neurophysiological evidence for experience-dependent motor resonance in infancy," *Neuroimage*, **43**(4), 808–814 (2008).
51. A. Vandermeer, F. Vanderweel and D. Lee, "The functional-significance of arm movements in neonates," *Science* **267**, 693–695 (1995).
52. D. Wood, J. S. Bruner and G. Ross, "The role of tutoring in problem-solving," *J. Child Psychol. Psychiatry* **17**, 89–100 (1976).
53. P. Zukow-Goldring and M. A. Arbib, "Affordances, effectivities, and assisted imitation: Caregivers and the directing of attention," *Neurocomputing* **70**, 2181–2193 (2007).