



Emergent imitative behavior on a robotic arm based on visuo-motor associative memories

A. de Rengervé, Sofiane Boucenna, P. Andry, P. Gaussier

► To cite this version:

A. de Rengervé, Sofiane Boucenna, P. Andry, P. Gaussier. Emergent imitative behavior on a robotic arm based on visuo-motor associative memories. International Conference on Intelligent Robots and Systems (IROS), 2010, Oct 2010, Taipei, Taiwan. pp.1754-1759. hal-00522723

HAL Id: hal-00522723

<https://hal.science/hal-00522723>

Submitted on 1 Oct 2010

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

Emergent imitative behavior on a robotic arm based on visuo-motor associative memories

Antoine de Rengervé, Sofiane Boucenna, Pierre Andry and
Philippe Gaussier*

ETIS UMR CNRS 8051, ENSEA, University Cergy Pontoise
F-95000 Cergy Pontoise, France

{antoine.rolland-de-rengerve,sofiane.boucenna,andry,gaussier}@ensea.fr

Abstract—Mimicry and deferred imitation have often been considered as separate kinds of imitation. In this paper, we present a simple architecture for robotic arm control which can be used for both. The model is based on some dynamical equations [1], which provide a motor control with exploring and converging capacities. A visuo-motor map [2] is used to associate positions of the end effector in the visual space with proprioceptive position of the robotic arm. It enables a fast learning of the visuo-motor associations without needing to embed a priori information. The controller can be used both for accurate control and interaction.

It has been implemented on a minimal robotic setup showing some interesting emergent properties. The robot can reproduce simple gestures in mimicry situation and finalized actions in deferred imitation situation. Moreover, it can even show some “intention” recognition abilities. Finally, the experiment of deferred imitation which inherits from learning by demonstration, also provides a good basis for cooperative and interactive experiments.

I. INTRODUCTION

With the aim of discovering and learning new tasks, imitation is an important advantage for autonomous agents. It allows to speed up the learning process by reducing the search space of the learner. It is also a very intuitive and natural way to interact. As a result, in the field of robotics, imitation is very often understood as a powerful behavior corresponding to the ability to learn by observation [3], [4]. A robot can combine known actions to reproduce a behavior that is done by somebody else. Interestingly, this notion of learning by observation is very close to the notion of “deferred imitation” intensively described in developmental psychology. Deferred imitation [5] is the ability of a child (from 18 months old) to reproduce an observed task but in a different context (spatial and especially temporal). According to Piaget, it also corresponds to the ability to memorize and handle symbolic representations of the task. This approach has been reinforced by the discovery of mirror neurons [6] which fires whenever a specific action is observed or done. More interestingly, the last 20 years of debate and research in developmental psychology changed this notion of imitation. Recent findings such as “neonatal imitation” have stressed the importance of imitation for communication, especially pre-verbal communication [7]. Initially described as “mimicry”, or an “automatic response” with limited

interest, this low-level imitation corresponds to a simple, immediate and synchronous imitation of meaningless gestures by the baby (from birth to a few years old). Immediate and spontaneous imitation appears as a precursor of higher level functions such as learning by observation.

Some research explores this bottom-up approach of imitation. To avoid using a full model of other, [8] suggests to use visual segmentation in a humanoid robot for mimicking of a human. The a priori knowledge for the segmentation of human body and the mapping from vision to motor can be replaced by self observation. This principle enables the building of a forward and inverse kinematic model that is not given a priori. This approach was used for imitation of hand posture [9], or movement patterns [10]. It can also help the robot to learn features of an object for grasping perspectives [11].

We also defend that higher level behaviors can emerge from the combination of simple low-level mechanisms coupled with appropriate on-line learning rules and principles like self observation [12]. In this paper, we propose a simple architecture that uses a sensorimotor learning map [2] to learn the backward model and dynamical equations [1] to generate motor commands. It can provide a reliable control for reaching an object using visual cues. It allows low level imitation of meaningless gestures which are done by a human in front of it. It also enables a deferred imitation where the robot sees and then reproduces a finalized action ie. an action that implies goals.

This architecture has been implemented on a minimal robot composed of a robotic arm and a 2D camera (Fig. 1). Visual information is extracted from captured images, using color detection, and projected in a space invariant to the movements of the camera. No depth information is provided. Section II details the learning rules and the dynamic equation that build the kinematics models and enable the control of the arm. In section III, this architecture is used in a setup where immediate imitation can emerge. In section IV, adding a simple sequence learner, we show that our architecture can also reproduce finalized actions in situation of deferred imitation. Finally, our architecture provides basis for guiding the robot through visual interaction. The robot can discover new perceptions that are necessary to build higher level behaviors.

*Philippe Gaussier is with Institut Universitaire de France,(IUF)

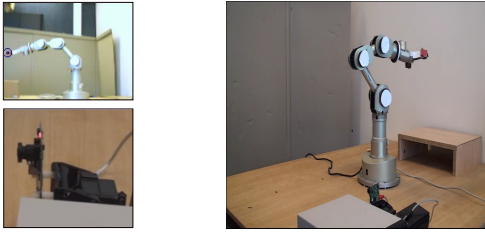


Fig. 1. The robot used for all the experiments is composed of a Neuronics Katana 6M180 with parallel gripper, one camera mounted on two servo-motors controlled by a SSC32 card and a PC for computation. The two servo-motors simulate the movement of a neck with a vertical rotational axis and a horizontal axis. On top left is given a capture from the camera.

II. LEARNING VISUO-MOTOR ASSOCIATIONS

The autonomous control of a multi-DoF robotic arm requires that the architecture copes with both visual and motor information. On a simple embodied robot, vision delivers useful information about the task. The arm controller must be able to link visual information with a command in the motor space. Recently, [1] have proposed a solution (Yuragi/fluctuation method) for arm control in the motor space. This model is based on the use of a limited number of attractors allowing the arm to converge reliably toward one of the motor configurations which correspond to them. Initially, these attractors were randomly distributed. Taking inspiration from this model, our working hypothesis is that proprioceptive configurations associated with the visual positions of the arm end effector can be used as attractors to achieve the visuo-motor control. By doing so, we can obtain a controller able to drive the arm in the working space using visual information to select the motor attractor(s) that will be used to compute the speed of the joints (see Fig. 2).

Following Langevin equation (see (1)) used to describe Brownian movements, [1] proposed that using random configurations expressed in the joint space of the arm ($\mathbf{x}$ is the current proprioception) combined with a noise parameter is enough to move a robotic arm toward any position by controlling the speed $\dot{\mathbf{x}}$.

$$\tau \dot{\mathbf{x}} = A \cdot f(\mathbf{x}) + \epsilon \quad (1)$$

Modulating the gain factor A of the command $\dot{\mathbf{x}}$ or the noise level ϵ enables to switch from converging toward one of the selected motor configurations (Fig. 3(a)) to exploring randomly the working space by “jumping” from an attraction basin to another (Fig. 3(b)) and vice versa. Adapting the initial model described in [1], we propose to take into account the contribution of visual stimuli for generating the motor command with (1) and (2).

$$f(\mathbf{x}) = \sum_{\mathbf{x}_{ij} \text{ attractor}} V_{ij} \cdot N(\Delta \mathbf{x}_{ij}) \cdot \Delta \mathbf{x}_{ij} \quad (2)$$

τ is the time constant of the dynamical system, A is a gain factor and ϵ is the noise level. Each gradient $\Delta \mathbf{x}_{ij}$ to the attractor $\mathbf{x}_{ij}$ is weighted by a normalized Gaussian function

$$N(\Delta \mathbf{x}_{ij}) = \frac{G(|\Delta \mathbf{x}_{ij}|^2)}{\sum_{\mathbf{x}_{i'j'} \text{ attractor}} V_{i'j'} \cdot G(|\Delta \mathbf{x}_{i'j'}|^2)} \text{ which uses a}$$

Gaussian expression $G(x) = \exp(-\beta x^2)$ of the distance $|\Delta \mathbf{x}_{i'j'}|^2$ between the current proprioception $\mathbf{x}$ and the attractor $\mathbf{x}_{i'j'}$. β determines the slope of the attraction basin around the learned attractors. The higher it is, the more attracted by the closest attractor the arm is. Importantly, each gradient $\Delta \mathbf{x}_{ij}$ depends on a weighting by visual neuron¹ V_{ij} . This enables to have a selection of the attractors according to the visual stimuli. Therefore, one or several attractors can be activated according to the implementation of V_{ij} (e.g. hard or soft competition).

Using (1) and (2) enables different control strategies that are shown in Fig. 3 with a simulation of a 3 DoF robotic arm. If A is strong as compared to the noise level ϵ , the arm converges to the nearest attractor. When the noise level is high, the arm explores the motor space. Modulating the gain factor A and the noise level even allows to build virtual attractors i.e. which weren't previously learned.

In order to associate each attractor x_{ij} with the position of the arm end effector in the visual space, we introduce a sensorimotor map of cluster of neurons [2]. It is a Neural Network model composed of a 2D map of clusters of neurons which has the same topology as the visual input. Each cluster of this map associates a single connection from one neuron of the visual map with multiples connections from the proprioception of the arm. Visual information (I) controls the learning of particular pattern of proprioceptive input ($Prop$). A cluster (i, j) is composed of:

- one input neuron I_{ij} linked to the visual map. This neuron responds to the visual information and triggers learning.
- One cluster, ie. a small population of Y_{ij}^k neurons ($k \in [1, n]$), which learns the associations between N-

¹For simplicity, we use a generical term because we do not want to refer to a specific cortex area as these visual inputs can basically come from detection of shape, color or movement...

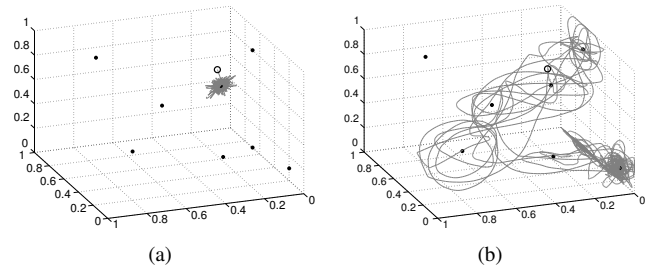


Fig. 3. Simulation of 3 DoF arm proprioception using (1). **a**: The trajectory converges to the nearest attractors. Simulation parameters are following: number of iterations=1000 ; beta (Gaussian parameter)=20 ; noise level $\epsilon_{max}=0$; Number of attractors = 8 ; shading parameter=0.01 ; $A=1$. **b**: When the ratio of A to noise level decreases, noise has a stronger effect on the speed command and allows an exploration of the motor space with jumps from an attractor to another. Simulation parameters are following: number of iteration=about 5000 ; beta (Gaussian parameter)=20 ; noise level $\epsilon_{max}=1$; $A=1$; shading parameter=0.01.

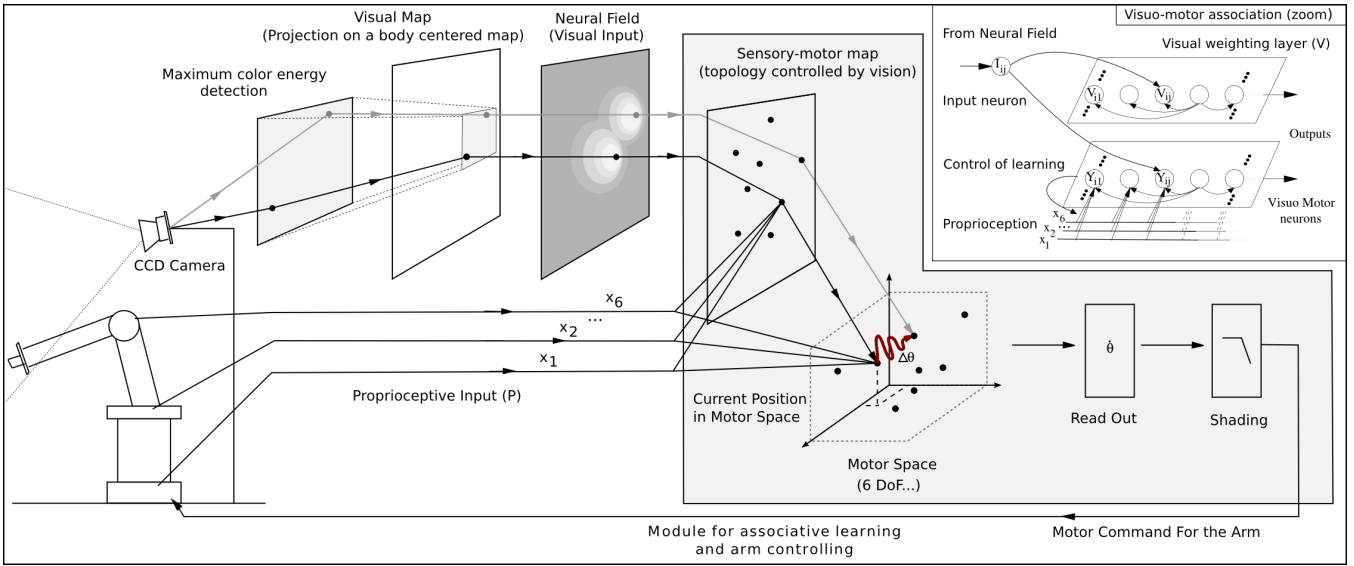


Fig. 2. Model of the arm controller. The sensorimotor map can learn to associate visual stimulus and proprioceptive information of the arm. A competition between visuo-motor neurons enable to associate current proprioception with the most activated visual input neuron (I_k). Visual weighting layer can support a different processing. Thus, neurons on this layer can activate one or several attractors (constructed from visuo-motor neurons) in the motor space. There is no constraints on the number of Degrees Of Freedom. If the current position of the arm is different of the generated attractor, a non-null motor command is read out by Fukuyori's adapted dynamical equations (1),(2) and given to the robotic arm.

Dimension proprioceptive vectors and 2D vision position. This population is a small Self Organized Map (SOM) [13].

In each cluster, the potentials of the neurons depends on the distance between its proprioceptive weight connections and the proprioceptive inputs. If the distance is null, the potential is maximal and equal to 1. The greater is the distance, the lower is the potential.

The learning of a cluster is dependent on the activation of the maximal trigger I_{ij}^M from visual input (3). Thus, each cluster learns the different proprioceptive configurations independently.

$$I_{ij}^M = \begin{cases} 1 & \text{if } I_{ij} = \max_{i,j} (I_{ij}) \\ 0 & \text{otherwise} \end{cases} \quad (3)$$

When a cluster receives learning signal, each of its neurons Y_{ij}^k performs the learning rule :

$$\Delta W_{ij}^{Prop} = \alpha \cdot (\mathbf{x}_m - W_{ij}^{Prop}) \cdot \text{dist}_{\text{weight}}, \quad (4)$$

where α is a learn factor, $\mathbf{x}_m$ is the m^{th} component of the proprioception vector and W_{ij}^{Prop} is the weight of the link between the proprioception neuron and the sensorimotor neuron. $\text{dist}_{\text{weight}}$ is looked up from a Difference of Gaussian table according to the topological distance between the k^{th} neuron Y_{ij}^k and the neuron with the maximal potential (the topology of the SOM is used). Therefore in a visually activated cluster, a sensorimotor neuron will encode on its weights the associated proprioception.

In an extreme case, each cluster can be reduced to a single neuron if the task to be learned does not need the learning of two positions that could only be separated according to the depth information. Positions that have close coordinates

x, y in camera image activates to the same cluster in the sensorimotor map and the same proprioceptive configuration.

The architecture has been tested on our robotic setup. To test the precision of the controller, a point is selected in the visual space (160x120 pixels). This point simulates a target to be reached by the robot.

The visual space is uniformly divided into few areas so that selected proprioceptive configurations can be associated with them. Only the 120x120 pixels on right side of the visual field are reachable by the robotic arm. Nine proprioceptive attractors are learned on this area. Each one of them is associated with a 40x40 pixels areas.

During the test, visual neurons V_{ij} are set to 1 if the distance between visual target and the center of the visual area associated to V_{ij} is lower than 30 pixels. Otherwise, V_{ij} is set to 0. This enables to select few attractors around the target. The distance d between the target point and the robot hand is also measured. The gain factor A is chosen proportional to this distance d : $A = K \cdot d$ if the distance decrease ($\frac{\delta d}{\delta t} < 0$) and null otherwise. ϵ is defined by $\epsilon = \epsilon_N \cdot d$. The modulation of the speed $\dot{\mathbf{x}}$ by the distance to the target introduces a bias in the movement which generates a virtual attractor at the target position. According to the constants K and ϵ_N the arm will be able or not to move from one attractor to another in order to get closer to the target. If ϵ_N is too low as compared to K the arm can be trapped in a distant attractor. The value of β is chosen low to make the random exploration between the learned attractors easier.

After verifying the convergence of the arm to a position learned during the learning phase, we tested the convergence to a visual position spotted between four learned attractors (Fig. 4 and 5). If the arm/target distance is lower than 3 pixels then the movement is stopped, the target is reached.

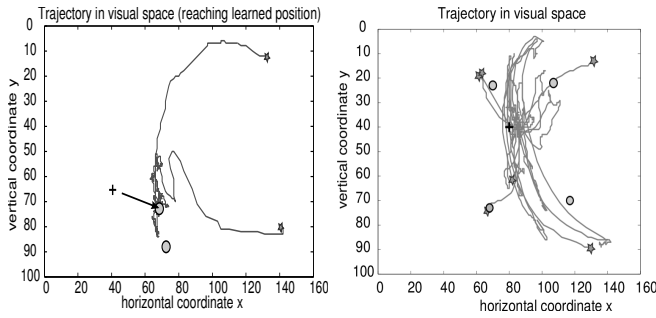


Fig. 4. Trajectory of the robot arm end effector in visual space. The black circles correspond to the learned attractors and the black cross is the visual target to be reached. The stars are the starting positions for each trial. Left: reaching a learned attractor, 2 attractors activated. Right: reaching a not previously learned position, 4 attractors activated.

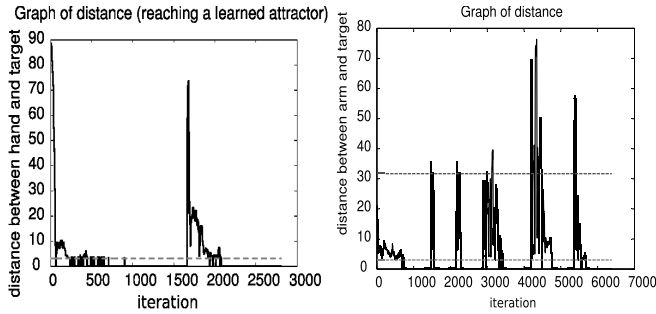


Fig. 5. Reaching a target. The figures show the result of experiments where the robot arm had to reach a target in the visual space. Experiments are made of several trials. For each trial the arm is initialized at a different position. We record the distance between the arm end effector and the target in the visual space (number of pixels). Dark gray dash-line shows the average distance to the attractors. The light gray line shows the threshold under which the target is reached. Left: Reaching a learned position, 2 trials. Right: Reaching a not previously learned position, 6 trials.

Fig. 5 shows the results when the robot has to reach a not learned position. The target position is at a distance of about 25 pixels of the visual location of the four activated attractors. By using a single learned attractor, the precision in visual space would be limited to 20 pixels (visual distance to the closest attractor). In this experiment, the visual distance between the position of the arm after convergence and the target is under the precision of visual localization (3 pixels). This shows that the cooperation of the sensorimotor map and the Yuragi method (1) offers an interesting basis for the control of a robotic arm with a self learning of the associations between visual and motor spaces.

III. EMERGENCE OF LOW LEVEL IMITATION

We consider a system that is not able to discriminate visual information about self movements from the movements of others (perception ambiguity). In such a system, imitative behavior can emerge easily. Whereas this hypothesis can be seen as very speculative, numerous psychological works show comparable human behaviors when visual perception is ambiguous. In 1963, Nielsen proposed an experiment in which subjects were placed in front of a semi-reflecting mirror [14]. In a first condition the mirror is transparent, and the subject sees his own hand placed on a table under

the mirror. In a second condition the mirror reflects, and the subject sees another hand (the demonstrator's hand) that he will mismatch for his own hand. Because a black glove has been put on both hands, the subject has the feeling to see his own hand and does not imagine there is another hand in the experiment. During each trial, the subject has to draw a straight line with a pen in the direction of his own body axis. When the perceived hand is his own hand, the performance is perfect. When the mirror reflects, the subjects "imitate" the other hand movements and do not perceive any difference if both hands are almost synchronous. If the perceived hand moves in a quite different direction, the subjects tend to correct the error by drawing in the opposite direction but they never suspect the presence of another arm (they believe that the "wrong" trajectory is due to their own mistake). In our robotic transposition of this experiment, a robot is able to imitate simple movements performed by a human demonstrator by using directly our visuo-motor controller.

Since the visual processing of our robot is based on color detection, its vision system can't differentiate its extremity (identified by a red tag) from another red target like a human hand (because of the skin color). As a consequence, moving the hand in front of the robot will induce visual changes that the robot interprets as an unforeseen self movement. Besides, the architecture which uses a simple perception-action control loop is designed with respect to the *homeostatic* principle. The system tends to maintain the equilibrium between vision and proprioceptive information. If a difference is perceived then the system tries to act in order to reach an equilibrium state. To do so, the robot moves its arm so that its proprioceptive configuration corresponds to the perceived visual stimuli according to its sensorimotor learning. This correction induces the following of the demonstrator's gesture.

To perform the homeostatic control, the gain factor A is set to 1 and the noise ϵ is set to 0. The visual weighting neurons V_{ij} now perform a hard competition (Winner-Take-All) so that $V_{ij} = 1$ if $I_{ij} = \max_{i',j'}(I_{i'j'})$ and 0 otherwise. Therefore, only the most visually activated attractor is used to control the arm. We tested the architecture with the basis of 16 learned attractors: the robot moves its arm in all the workspace and learns the visuo-motor associations corresponding to 16 different visual positions of the arm's end effector. The positions are chosen with respect to the visual sampling. This procedure takes approximately two minutes which corresponds to the time necessary for the end effector to visit the 16 positions, one by one. Then, the camera is shifted so that the robot only see the human's hand. Color detection² activates a visual neuron which selects a proprioceptive attractor. The arm moves to reach the learned configuration. Thus, the ambiguity of perception of the system starts to drive the hand along a trajectory similar to the human's hand (Fig. 6). The gesture imitative behavior emerges as a side effect of the perception limitation and the homeostatic property of the controller [2].

²In previous work the movement was used in the same way.

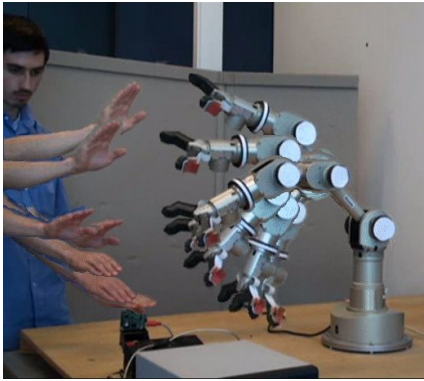


Fig. 6. Trajectories during low level imitation.

IV. DEFERRED IMITATION OF FINALIZED ACTIONS

Fig. 7 presents an experiment where the same architecture as in previous section is used for catching and manipulating a can. As in low-level imitation, color detection activates visual neurons. As each neuron encodes for a specific area of the visual field, the activated neurons can be used as visual states. Thus, the robot translates an observed action into a sequence of such visual states. A simple temporal sequence learning module based on discrete transition learning [12] enables the robot to memorize the action.

The learned task is composed of only four visual states. The first state is the visual position of the hand of the professor before he catches an object, the second one is the visual position of the hand when the professor catches the object, the third one is an intermediary point on the trajectory and last state is the position of the hand when the experimenter releases the object.

The temporal sequence learner can output a predicted visual state thanks to the learned transition. This prediction is timing dependent. The output is feedback to the sequence learner. Therefore, it can replay internally a complete action with its learned timing. Once the temporal sequence has been learned, seeing the hand at a learned visual state (eg. the first one) induces the prediction of the next visual state. The internal loop enables the robot to continue this replay until the end. Besides, the output visual state from the sequence learner can activate an attractor just like if the visual information comes directly from vision processing. The same homeostatic principles are respected. Given a visual state, the arm will try to reach the proprioceptive configuration which best corresponds according to the sensorimotor learning. Thus, the robot, stimulated by the human hand at the first position, reproduces the reaching and grasping (reflex³) of the object. It then moves it to the final position following approximately the same trajectory as in demonstration.

This setup is very similar to the one used in learning from demonstration paradigm [15], [16] and [17]. First, during a demonstration phase, the architecture learns the action to be

³Grasping issues were strongly simplified, using an ad-hoc procedure driving the robot to close its gripper whenever an object is detected by the proximetric sensors of the fingers.

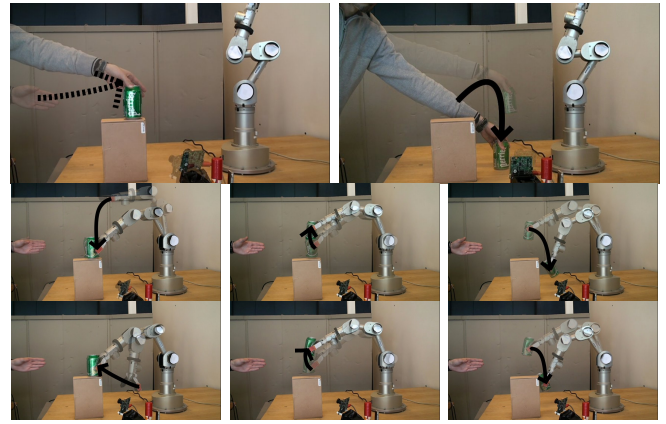


Fig. 7. Deferred imitation based on visual sequential learning : an emergent cooperative behavior. The robot is shown the action : grasping and moving an object to a given location. In each picture, two successive positions of the arm are superposed with the black arrow showing the performed movement. Up : Learning phase. The robot sees the teacher doing the action. Middle and Bottom : Reproduction of the action. The robot sees the visual trigger from the hand on left side of the visual field. It then completes the action previously demonstrated according to its motor capacities. An accompanying video to this paper details the experiment.

done. Then, an operator switches the robot in reproduction mode. In [16] and [17], the task learning is a behavior model training which needs several sets of data. Besides, this training is based on the proprioception. A backward model must be provided a priori when using vision to recognize and monitor the task. On the contrary, our architecture uses one-shot learning. It needs to learn previously the backward model but it can then learn the action directly in the visual space. No model was given a priori by a human. So our robot can perform more autonomously deferred imitation based on observation.

Our robot reproduces a sequence of known gestures without specific meaning. This sequence drives the robot to catch an object and move it to another place but the robot has no representation of the task. Even, the robot discovers that there is an object only when it closes its gripper on it and starts to move it. Interestingly, the robot is not trying to reach the experimenter's hand. This latter is outside of the working space of the robot and no visuo-motor attractor were learned there. For an exterior observer, the robot performs deferred imitation. An observer can make the link between the human's intentions (here the apparition of the hand toward the object) and the robot's behavior (responding by grasping the object) and he can interpret it as some intention recognition and cooperation. This is an emergent property that comes from the interplay between what the robot can and cannot do and the ambiguity of its perception.

V. CONCLUSION

In this paper, we proposed a simple architecture which can deal with both low-level and deferred imitation. When considering imitation from a classical point of view, the following important issues are usually raised:

- 1) the issue of the elementary motor repertory, and how to combine actions to form new behaviors. If the

actions are defined at a too high level (for example using symbols), there is a risk of a symbol grounding problem [18], otherwise it remains a complex issue of motor control.

- 2) the issue of being able to link actions of self, with those of others. This issue is known as the correspondence problem [19]. The morphologies and dynamics of the imitator and imitated can be slightly different.
- 3) the issue of using notions such as “self” and “other” in the control architecture of a robot.

In this paper, we have adopted a developmental approach of robotics, with the goal of understanding the mechanisms underlying the various forms of imitations, and being able to obtain imitation as an emergent behavior from the interplay of low-level mechanisms. Adopting such a bottom-up approach have allowed us to avoid to have a top down analysis and to cope with these issues. With ambiguity of perception, our robot has no idea of “self” or “other” (3). Even, this ambiguity is necessary for the emergent imitation. The correspondence problem (2) is completely left aside as the robot will imitate only according to its own capacities. The robot has no special symbols used to pre-categorize and combine motor actions (1). However, these actions could help it to build progressively some representation of this world according to the bottom-up approach.

We tested this architecture on a simplified robot (one arm, one 2D camera) and showed the emergent imitating abilities. Of course, a unique 2D camera can not deal with positions visually similar but different in depth. In the case of tasks in a horizontal plane (over a table) with the 2D camera placed on top of the robot (human like configuration), this issue is not a major problem. Though, future work should develop the capacities of the robot and its multimodal perception. The robot should be able to memorize positions at different depth at the same visual position. With several configurations possible for a given visual target, the robot should be able to rely on other modalities (touching, another camera,...) to disambiguate its visual information and select the right action. As in [11] and [20], new perceptions can raise interest of the robot and influence its actions. Discovering an object by touching it should raise the curiosity of the robot as a new range of possible actions is discovered. In [21], one of the caregiver roles is to provide a guidance to infants so that they can feel new perceptions that require to learn new behaviors. Imitation, like a complement to passive manipulation, allows to help a robot to explore its environment and find new perceptions. How new perceptions should modify the internal dynamics of the robot to generate learning of new abilities is a major and challenging issue.

Finally, it is very interesting to use a visual task space for memorizing the task while the control is performed in the motor space. Playing with the limitation between what can be seen and what can be done, the robot can get for free some intention recognition ability. When it sees the human hand at a specific position, it can trigger the action of catching and moving the object as observed in the demonstration. This intention recognition can easily be used for cooperative tasks

like completing an action which is initiated by the human.

VI. ACKNOWLEDGMENTS

This work was supported by the French Region Ile de France, the Institut Universitaire de France (IUF), the FEELIX GROWING european project and the INTERACT french project referenced ANR-09-CORD-014.

REFERENCES

- [1] I. Fukuyori, Y. Nakamura, Y. Matsumoto, and H. Ishiguro. Flexible control mechanism for multi-dof robotic arm based on biological fluctuation. In *Proceedings of the Tenth International Conference on the simulation of adaptive behavior (SAB'08)*, pages 22–31, 2008.
- [2] P. Andry, P. Gaussier, and J. Nadel and. B.Hirsbrunner. Learning invariant sensory-motor behaviors: A developmental approach of imitation mechanisms. *Adaptive behavior*, 12(2), 2004.
- [3] Y. Kuniyoshi. The science of imitation - towards physically and socially grounded intelligence -. Special Issue TR-94001, Real World Computing Project Joint Symposium, Tsukuba-shi, Ibaraki-ken, 1994.
- [4] S. Schaal. Is imitation learning the route to humanoid robots ? *Trends in cognitive sciences*, 3(6):232–242, 1999.
- [5] J. Piaget. La formalisation du symbole chez l'enfant. imitation, jeu et rêve. image et représentation. 1945.
- [6] V. Gallese, L. Fadiga, L. Fogassi, and G. Rizzolatti. Action recognition in the premotor cortex. *Brain*, 119:593–609, 1996.
- [7] J. Nadel. *The evolving nature of imitation as a format for communication.*, pages 209–235. Cambridge University Press, 1999.
- [8] G. Cheng and Y. Kuniyoshi. Real-time mimicking of human body motion by a humanoid robot. In *in: Proceedings of the Sixth International Conference on Intelligent Autonomous Systems (IAS2000)*, pages 273–280, 2000.
- [9] T. Chaminade, E. Oztop, G. Cheng, and M. Kawato. From self-observation to imitation: visuomotor association on a robotic hand. *Brain research bulletin*, 75(6):775–784, April 2008.
- [10] Y. Kuniyoshi, Y. Yorozu, S. Suzuki, S. Sangawa, Y. Ohmura, K. Terada, and A. Nagakubo. *Emergence and development of embodied cognition: a constructivist approach using robots*, volume 164, pages 425–445. Elsevier, 2007.
- [11] L. Natale, F. Orabona, F. Berton, G. Metta, and G. Sandini. From sensorimotor development to object perception. In *5th IEEE-RAS International Conference on Humanoid Robots, 2005.*, pages 226–231. IEEE, 2005.
- [12] M. Lagarde, P. Andry, P. Gaussier, S. Boucenna, and L. Hafemeister. Proprioception and imitation: On the road to agent individuation. In Olivier Sigaud and Jan Peters, editors, *From Motor Learning to Interaction Learning in Robots*, volume 264, chapter 3, pages 43–63. Springer Berlin Heidelberg, Berlin, Heidelberg, 2010.
- [13] T. Kohonen. Analysis of a simple self-organizing process. *Biological Cybernetics*, 44(2):135–140, juillet 1982.
- [14] T.I. Nielsen. Volition: A new experimental approach. *Scandinavian Journal of Psychology*, 4:225–230, 1963.
- [15] B. D. Argall, S. Chernova, M. Veloso, and B. Browning. A survey of robot learning from demonstration. *Robot. Auton. Syst.*, 57(5):469–483, 2009.
- [16] S. Calinon and A. Billard. *Learning of Gestures by Imitation in a Humanoid Robot*. Cambridge University Press, k. dautenhahn and c.l. nehaniv edition, 2006. in press.
- [17] S. Schaal. Learning from demonstration. In *advances in neural information processing systems 9*, pages 1040–1046. mit press, 1997.
- [18] S. Harnad. The symbol grounding problem. *Physica D*, 42:335–346, 1990.
- [19] C. L. Nehaniv and K. Dautenhahn. Mapping between dissimilar bodies: Affordances and the algebraic foundations of imitation. In J. Demiris and A. Birk, editors, *Learning Robots: An Interdisciplinary Approach*. World Scientific Press, 1999.
- [20] P.-Y. Oudeyer, F. Kaplan, V. V. Hafner, and A. Whyte. The playground experiment: Task-independent development of a curious robot. In *Proceedings of the AAI Spring Symposium on Developmental Robotics, 2005*, Pages, pages 42–47, 2005.
- [21] P. Zukowgoldring and M. Arbib. Affordances, effectivities, and assisted imitation: Caregivers and the directing of attention. *Neurocomputing*, 70(13-15):2181–2193, August 2007.