

An Affordance-Based Intersubjectivity Mechanism to Infer the Behaviour of Other Agents

Simon L. Gay¹, François Suro¹, Olivier L. Georgeon² and Jean-Paul Jamont¹

Abstract— We introduce an architecture based on interactionist principles which enables autonomous agents to infer the behavioural patterns of other agents, without prior knowledge.

Previous works have shown that agents can leverage the relations between entities and sensorimotor abilities afforded by them to build a model of the environment which supports behaviours satisfying innate motivational principles. However, efficient collaborative or competitive behaviours relies on the ability to infer hidden motivational principles of other agents by observing their behaviour. In this paper, we introduce a novel architecture able to identify mobile affordances generated by other agents, infer their behavioural preferences and predict their movements. We validate our proposition through experiments showing that a predator agent can learn to infer that its prey is attracted by specific elements of the environment.

I. INTRODUCTION

This paper addresses the issue of predicting the movement of mobile entities driven by an unknown decision mechanism for artificial agents that learn a model of their environment through experience, without *a-priori* knowledge.

This study is related to the domains of artificial constructivist learning [1] and enactive learning [2], where learning occurs through *interactions* which are the enaction of control loops that implement Piagetian *sensorimotor schemes* [3]. The agent has no *a priori* knowledge of its environment. Instead, it is provided with a predefined set of uninterpreted *interactions* associated with predefined numerical *valences*. The agent is driven to enact interactions that have positive valences and avoid interactions that have negative valences. This drive defines a kind of *intrinsic motivation* [4] that we call *interactional motivation* [5] in the framework of *Radical Interactionism* (RI) [6] and *artificial interactionism* [7].

The integration of mobile entities is a step toward the long-term goal of emergent social interactions within groups of artificial agents. This challenge is two-fold:

- 1) learning to define, recognize and localize other agents moving freely in the environment.
- 2) inferring the intentions of these agents based on their own environmental contexts.

The first problem has been addressed in a previous work [8]. The present paper focuses on inferring the intentions of other agents, using the following four step process : 1) localize a mobile entity (other agent) and define the environmental

context from this entity's point of view, 2) keep track of a mobile entity to observe its dynamic properties, 3) learn the behavioural preferences of the mobile entity, 4) predict the future position of the mobile entity from its own context of affordances. This model relies on two assumptions: A) the agent can detect the mobile entity's affordances through its own interactions, B) the decision process of the entity is similar to RI agent (although behavioural preferences can differ).

We implement our proposition in simulations where a marine predator discovers elements of the environment, such as fishes (its prey) or algae, learns to localize and track fishes and finally infers that fishes are attracted to algae, their food source.

This paper is subdivided as follows: Section II summarizes and formalizes the Radical Interactionism model and previously developed models that are exploited by new mechanisms. Sections III to VI present the four processing steps allowing the integration and prediction of mobile entities. Finally, Section VII encompasses some conclusive remarks and future development of the intersubjectivity problem.

II. THE RADICAL INTERACTIONISM MODEL

In contrast with most machine learning approaches, an RI agent cannot directly access states of its environment, but uses outcome of control loops as input data. The agent learns and exploits regularities offered by its coupling with its environment. We do not design RI agents to reach predefined goals or maximize a reward value defined as a function of the system's state. Their purpose is to study the construction of emergent models of the environment, and the generation of behaviours through an open-ended learning process.

Let I be the set of predefined *interactions* (control loops), and $v_i \in \mathbb{R}$ the predefined valence of interaction $i \in I$. Valences define the inborn behavioural preferences of the agent. At the beginning of step t , the agent selects an *intended interaction* $i_t \in I$ to try to enact, and receives, at the end of step t , a set of actually *enacted interaction* $E_t \subset I$. The enaction cycle is a success if $i_t \in E_t$. An example of failure may be when an agent intends to move forward ($i_t = \text{move forward}$), but actually collides with an obstacle ($e_t = \text{collide}$, $e_t \in E_t$).

We distinguish between *primary interactions* that are Piagetian control loops (action, result), and *secondary interactions* that are a couple (primary interaction, additional sensory outcome). These additional outcomes are sensory input that cannot be separated from the movement or environment change produced by the primary interaction. The optical flow is an example of such sensory outcome that must be associated with a movement to characterize a position in space. E_t contains

*This work is supported by the French National Research Agency in the framework of the "Investissements d'avenir" program (ANR-15-IDEX-02).

¹LCIS, Université Grenoble Alpes, Valence, France
simon.gay/francois.suro/Jean-Paul.Jamont@lcis.grenoble-inp.fr

²EA1598, Lyon Catholic University, LIRIS, Claude Bernard University, Lyon, France. ogeorgeon@univ-catholyon.fr

only one enacted primary interaction e_t , other elements are secondary interactions *associated* with this primary interaction.

A *Signature* mechanism [9] evaluates the possibility of enacting an interaction i_{t+1} in a given context E_t . This mechanism rests upon the assumption that the result of interaction i depends on a limited set of entities in the environment, which defines the context that *afford* [10] i . Since an RI agent can only perceive its environment through E_t , the signature designates one or more sets of interactions whose previous enaction indicates the possibility to subsequently enact i . The signature of interaction i is formalized as a function $S_i : \mathcal{P}(I) \rightarrow [-1; 1]$ giving the estimation of the likelihood of successfully enacting i in a context E_t . $S_i(E_t) = 1$ means success is certain, and -1 means failure is certain. This estimation is improved by adjusting parameters of S_i each time i succeeds or fails. The signature's *pseudo-reverse function* $\hat{S}_i : \{1, -1\} \rightarrow \mathcal{P}(\mathcal{P}(I))$ provides either the minimal context(s) $C_k^i \subset I$ affording ($\hat{S}_i(1)$) or those preventing ($\hat{S}_i(-1)$) interaction i .

This model defines an *affordance* as prerequisites for the success or failure of an interaction. Thus, affordances are conditioned by the agent-environment coupling (i.e. the set I of sensorimotor possibilities). The agent cannot access the physical entity behind an affordance, but learns to detect its presence from the outcome of interactions. Defining entities of the environment through sensorimotor possibilities that they afford is abundant in literature (e.g. [11][12][13]). Because most of these approaches detect affordances using perception, they are limited to the prediction of the next action or require prior knowledge (e.g. [14]) about the environment and/or space. By using interactions, signatures can exploit spatial properties implicitly encoded in interactions, especially the movement that they produce, to detect distant affordances. A signature of an interaction i designates a set of interactions $\{j_k\}_k \subset I$ whose previous enaction signals the possibility to enact i . As interactions j_k , associated to the same primary interaction j , have their own signatures, it is possible to *project* the signature of i backward through j , by combining the signatures S_{j_k} of interactions j_k . The resulting signature $S_i^{(j)}$ designates a context affording i after enacting j , i.e. an affordance of i located a 'step j ahead'. The recursive chaining of primary interactions in a sequence $\sigma = \langle j^1, \dots, j^n \rangle$ thus allows the agent to detect an entity that affords i at the *position* designated by σ .

The agent keeps track of the position of discovered affordances in its Egocentric *Spatial Memory* (ESM) [9]. The ESM assumes that the position of an affordance can be defined by a couple of parameters (i, d) derived from the shortest sequence σ leading to this affordance: $d \in \mathbb{N}$ is σ 's length, and $i \in I$ is σ 's first interaction. A *Place* is a couple (i, d) , covering all sequences σ sharing the same parameters i and d . A place can be preceded by a short sequence of interactions, defining a *Composite Place*. Over time, the ESM learns relations between positions σ and composite places, then between composite places, allowing an accurate tracking of affordances outside perceptual range. The contexts of affordances are thus provided as tuples $(a, (i, d))$, with a the afforded interaction and (i, d) a

place characterizing its position.

A decision process defines a *utility* value of interactions leading to affordances stored in ESM M according to their valence, fostering interactions allowing to move towards affordances with positive valences [9]:

$$u_i = \sum_{(a_k, j_k, d_k) \in M} \nu_{a_k} \times f(d_k) \times id(i, j_k) \quad (1)$$

where $id(i, j) = 1$ when $i = j$ and 0 otherwise, ν_{a_k} is the valence of a_k , and f is a strictly decreasing and positive function characterizing the importance of an affordance according to its distance in the agent's decision.

Learning the signature of an interaction afforded by a mobile entity is challenging because the presence of the affordance at time t may still lead to a failure of the interaction at time $t+1$ if the entity has moved. The extended model proposed in [8] used differences appearing in failure predictions to eliminate false negative. Its implementation was based on a neural architecture where competing neurons define a set of exclusive contexts and their respective probabilities of success. The network can also invert the signature output, allowing to integrate "negative affordance" (i.e. affordances preventing an interaction). The contexts are characterized by weights of their neurons.

III. DEFINING THE CONTEXT FROM ANOTHER AGENT'S PERSPECTIVE

As behaviour is expected to be influenced by the context of surrounding affordances, the first step to infer an entity's behaviour is to estimate such a context in its reference.

A. Principle of Reference Change

A previous work [9] has shown how the ESM learned to update positions of stored affordances, in egocentric reference, as the agent enacted interactions. Let $A^t = \{(a_k, (i_k, d_k))\}_{k \in K^t}$ be the current context of affordances. We propose to apply recursively the interactions of a sequence σ to a copy A' of A^t . Thereby, we simulate the enaction of a sequence σ and obtain the context of affordances A_σ that should be observed when reaching the position defined by σ .

However, a sequence σ_i characterizes the position from which it is possible to enact i , and not the *physical* entity affording it. [9] showed the emergence of negative affordances which prevent the enaction of an interaction. For instance, *moving forward* is prevented by the presence of a solid object. The signatures of negative affordance designate the area the agent should occupy after enacting the interaction. This property can be exploited to locate the physical object defining an affordance: for a given affordance a_i at position σ_i , we propose to consider an alternative interaction j of i (i.e. the enaction of i can lead to the enaction of j) that has a negative signature. Then, if there is a part $C \subset E_t$ such that $S_i^{\sigma_i}(C) \approx 1$ and $S_j^{\sigma_j}(C) \approx -1$ (i.e. C affords i at σ_i and prevents j at σ_j), it is possible to assume that the agent can occupy the area of the physical object affording i by moving through sequence $\langle \sigma_j, j \rangle$. We note $\sigma_{a_i} = \langle \sigma_j, j \rangle$ the position of the physical object affording i .

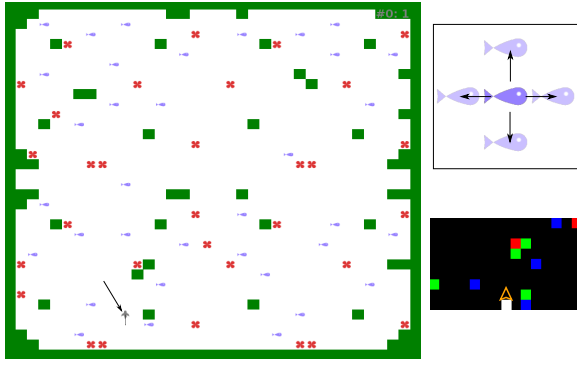


Fig. 1. The test environment. The grey shark (bottom left) is our agent, green blocks are walls, red leaves are algae and blue fish are mobile prey. At each simulation step, the fish can move up, down, left, and right. The bottom-right frame shows the visual system of our agent, which covers an area to the front and sides, and extends one row behind.

Unlike [8], we keep the sets of sequences $\{\sigma_{i,k}\}_k$ (and $\{\sigma_{a_i,k}\}_k$) which describe the paths to the future possible positions of an affordance of i . Each sequence $\sigma_{a_i,k}$ allows the ESM to recall the contexts, noted $A^{\sigma_{a_i,k}}$ that should be observed for each possible position of a_i .

B. Experimental Setup

For our experiments, we use an artificial agent in a discrete¹ 2D environment (Fig. 1). The environment contains three types of objects (entities) that the agent can identify by colour: walls (green) which are obstacles, algae (red) which can be crossed, and fish (blue) which are mobile agents behaving as follows: they are attracted to algae and repulsed by walls. Their perception range is 7 grid units. When a fish *eats* an alga, it is removed and a new one is placed at random in the environment. Fish have a 15% chance of inactivity, and are also inactive when no affordances are detected. A movement vector $\vec{V}$ of a fish is defined from entities o_k within perception range:

$$\vec{V} = \sum_{o_k} v_{o_k} \times e^{-d_{o_k}} \times \vec{p}_{o_k} \quad (2)$$

where v_{o_k} is the valence of o_k (20 for algae, -2 for walls), d_{o_k} the Manhattan distance of o_k , and $\vec{p}_{o_k}$ the vector towards o_k . The fish moves along cardinal directions, according to $\vec{V}$.

The sensorimotor possibilities of the agent define six primary interactions:

- ▷ move forward one step ► bump into an obstacle
- ◁ turn left by 90° ► eat something edible
- ◃ turn right by 90° ► slide on a soft object

The agent is equipped with a visual sensor that detects three colors (red, green, blue) and measure distances. While enacting a primary interaction (except for *bump* that does not produce movement), the agent may *see* red, green or blue entities in positions of its egocentric space. Its field of view, shown in the bottom-right frame of Fig. 1, is discretized as a 15 by 9 regular grid (this distribution of positions in space is unknown to the agent). We thus define a total of 2025 visual

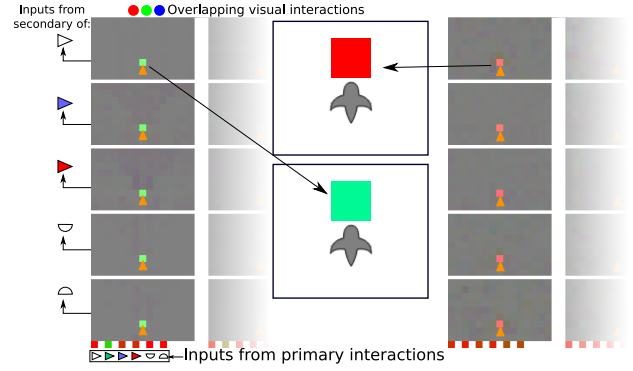


Fig. 2. Signatures of interaction *bump* (left) and *slide* (right), recorded after 100 000 simulation steps. A signature is characterized by the weights of 7 formal neurons, each neuron being represented by a column (in the same way than in [8]). As these signatures identified a unique context (static object), we only represent weights of one neuron. As external observers, we can organize weights to make signatures more readable: first, weights related to primary interactions are represented with six squares below (green for a positive weight, red for a negative weight). Weights associated with secondary interaction are grouped according to their primary interaction, forming the five groups (from top to bottom: forward, eat, slide, turn left, turn right; bump does not produce visual interactions). Each group is organized to place visual interaction with their associated position in space, relative to the agent (orange triangle). Colours associated with visual interactions are overlapped to generate signatures under the form of an RGB image. Signature of bump identified a context that consist of seeing a green element in front of the agent, and slide to a red element in front of the agent. *Bump* is also related to the success of *bump*, since this interaction can be enacted repeatedly.

interactions from the grid positions (9x15), colours (3) and primary interactions (5).

Signatures and distant affordance localization use the model and implementation proposed in [8], with a sequence length limit of 8 interactions. The agent is provided with a learning mechanism that fosters interactions with low certainty of success or failure (low $|S_i(E_t)|$).

We use a hard-coded ESM implementing properties shown in [9] to avoid observation bias induced by an incompletely learned structure. This ESM can *add* an affordance a detected at σ_a , *update* the context with an enacted interaction e , and *simulate* a sequence σ on a copy of the context. The ESM provides the current and simulated contexts of affordances as a set of tuples $\{(a_k, (i_k, d_k))\}_k$ where a_k is the afforded interaction and (i_k, d_k) its current *place*. The horizon of the ESM is set to 10 interactions, and will not register longer sequences nor keep older affordances. As the ESM was not tested with mobile entities, affordances of *eat* will not be stored.

Simulations on the acquisition of signatures have given results similar to [8]: signatures of interactions related to static affordances (bump and slide) emerged within 5000 simulation steps, while signatures of interactions related to prey (eat and moving forward unobstructed) required around 50 000 steps to get accurate contexts and probabilities. Signatures of visual interactions related to seeing blue (fish) also required more time, as their enaction were less frequent. These signatures successfully integrated the movement of primary interactions and the probabilities of each contexts making possible to detect mobile affordances. In the subsequent experiments, we used signatures obtained after 350 000 simulation steps.

¹For applications of the signature mechanism, detection mechanism and ESM in a continuous environment, see [9].

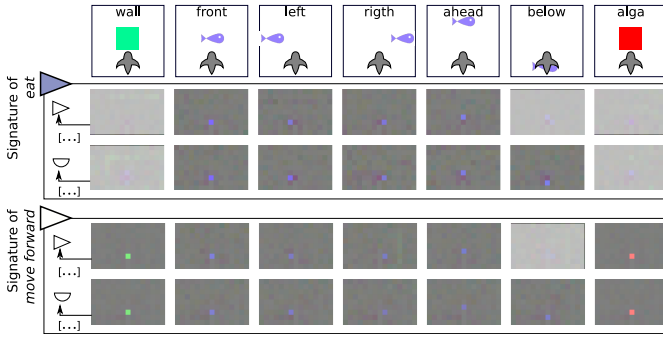


Fig. 3. Signature of *eat* and *move forward*, recorded after 100 000 simulation steps. Each column represents a neuron of the signature. We only display weights related to secondary interactions associated to move forward (first line) and turn right (fourth line), as other contexts are similar, as shown in Fig. 2. Greyed contexts have low weights and are thus unused by the signature. The signatures successfully integrated the possible contexts affording them (in this Figure, neurons are sorted to match the above situations). The second layer weight of *move forward* is negative: the signature thus represents contexts *preventing* moving forward. As it is an alternative of *eat*, it can be used to detect the physical position of the affordance of *eat*.

C. Defining the Context of an Affordance

After the agent enacts an interaction, the detection mechanism uses the new environmental context E_t to detect distant affordances. Fig. 4 shows that the *eat* interaction is seen at several positions corresponding to the possible moves of the fish and that all entities, including the fish, are identified as affordances preventing *move forward*. It is thus possible to correlate this data to determine the position of the *eat* affordance and perform simulations in the Space Memory.

IV. TRACKING THE POSITION OF A MOBILE AFFORDANCE

Observing the movement of a mobile affordance is required to infer its behaviour. The agent must therefore have the ability to keep track of the affordance over multiple steps.

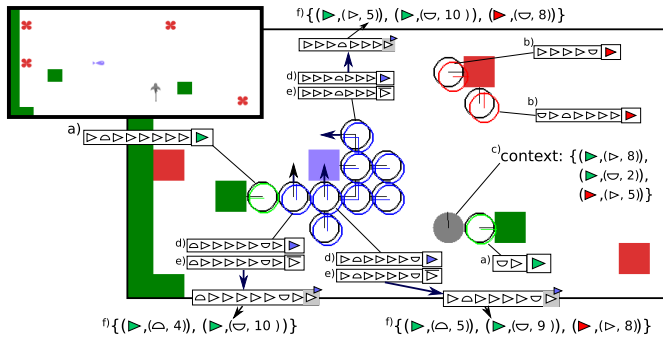


Fig. 4. Detection of distant affordances: the agent enacts an interaction, and perceives its environment (top left). The main representation shows the detected sequences σ as circles showing the position and orientation that would be obtained by enacting the sequence (the agent cannot access this geometric representation). A small offset is applied to observe overlapping circles. The agent detects two instances of *bump* (a) and of *slide* (b). These static affordances are stored in the Space Memory and localized with *places*, forming the context (c). *Eat* is detected through a set of sequences giving multiple positions (we only represent three sequences (d) among 28 detected). The agent also detects that these positions prevent *move forward* (black circles) (e), that has a negative signature. Then, by using the simulation possibilities of the Space Memory, affordance contexts (f) are defined from these positions.

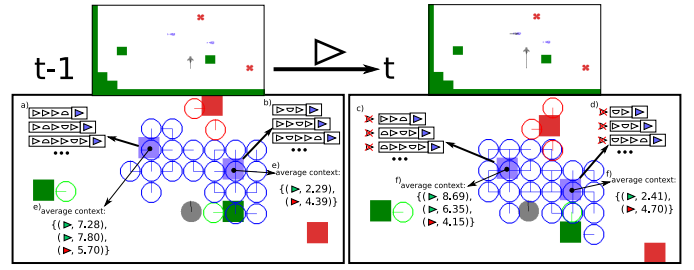


Fig. 5. Tracking and recognizing mobile entities. At step $t - 1$, the agent detected two instances of *eat*, respectively localized with 29 (a) and 11 (b) sequences, then enacted an interaction (here, *move forward*). The agent detects two instances of *eat*. Here, the recognition mechanism detected 7 common sequences for the left affordance (c) and 10 for the right affordance (d). Then, the agent computes the average distance of surrounding affordances from each mobile affordance's position. By comparing the distance at $t - 1$ (e) and t (f), the agent detected that left fish moved while right fish remains immobile.

A. Recognizing a Mobile Affordance

We can assume that we are considering the same static entity when affordances of the same interaction j are found at the same position σ over two steps. After the agent enacts an interaction i , the situated interactions σ_j^t and σ_j^{t-1} are considered related when $\sigma_j^{t-1} = \langle i, \sigma_j^t \rangle$.

The possible positions of a mobile affordance a_j are defined as a set of sequences $\Xi_{a_j} = \{\sigma_{a_j,k}\}_k$. Thus, after enacting i , the agent is expected to find at least a sequence of Ξ_{a_j} such as $\sigma_{a_j}^t = \langle i, \sigma_{a_j}^{t+1} \rangle$. Formally, we note the subset of sequences starting with i : $\Xi_{a_j}^i$. We can consider that two affordances of j detected at step $t - 1$ and t are the same if $\Xi_{a_j}^{t-1,i} \cap \Xi_{a_j}^t \neq \emptyset$. To associate multiple instance of the same affordance detected at different steps, a matching algorithm is defined, based on $\text{Card}(\Xi_{a_j}^{t-1,i} \cap \Xi_{a_j}^t)$.

B. Detecting Movement

We propose to use the variation in estimated distances of surrounding affordances: formally, a mobile affordance of i is localized through a set of sequences $\Xi_{a_i}^t = \{\sigma_{a_i,k}\}_{k \in K^t}$, and the space memory provides, for each sequence $\sigma_{a_i,k}$, a predicted context $A^{\sigma_k} = \{(a_m, (i_m^{\sigma_k}, d_m^{\sigma_k}))\}_m$, with $(i_m^{\sigma_k}, d_m^{\sigma_k})$ the simulated place of a_m through $\sigma_{a_i,k}$. Then, an average context can be defined using the average distance for each surrounding affordance. Since interactions $i_m^{\sigma_k}$ are here unnecessary, we define the average context as $\bar{A}_{a_i} = \{(a_m, \bar{d}_m)\}$, with $\bar{d}_m = \text{average}_k(\{d_m^{\sigma_k}\})$.

The average context can then be compared for two consecutive steps $t - 1$ and t . We detect that a mobile affordance actually moved through changes in relative distances of at least one surrounding (static) affordances: $\exists a_m \mid (a_m, d_m^{t-1}) \in \bar{A}_{a_i}^{t-1}, (a_m, d_m^t) \in \bar{A}_{a_i}^t, |d_m^{t-1} - d_m^t| > \text{threshold}$. To make the system more tolerant to imprecision in distance estimations, a threshold $\in [0, 1]$ must be defined.

Fig. 5 shows an agent identifying two separate affordances of *eat* over two consecutive steps. The matching mechanism associates the two occurrences of step t with the corresponding occurrences of step $t - 1$. Then, the average contexts $\bar{A}$ of the two occurrences are obtained over the two steps, showing that a fish moved while the other remained static.

An interesting aspect of this mechanism is that the movement is not defined in the reference of the agent (egocentric) but instead in the reference of static affordances, making the agent able to generate allocentric references.

V. DEFINING THE BEHAVIOURAL PREFERENCES OF MOBILE AFFORDANCES

Once the agent can define the context of a mobile entity and observe its movement, it becomes possible to infer the influence of affordances, and thus, the behavioural preferences of this entity.

A. Finding the Valence of a Mobile Affordance

We assume that the agent “projects” its own decision model on mobile affordances, that is, closing in on affordances (Eq. 1). However, it is not possible for the agent to observe the afforded interactions of an entity whose sensorimotor abilities may be entirely different. We thus rely on the differences between positions σ_1 and σ_2 to define the utility value of a movement. If these two positions have contexts with the same affordance, Eq. 1 can be adapted :

$$u_{\sigma_1, \sigma_2} = \sum_{\substack{(a_k, (j_{k,1}, d_{k,1})) \in A_{\sigma_1} \\ (a_k, (j_{k,2}, d_{k,2})) \in A_{\sigma_2}}} \nu_{a_k}^i \times (f(d_{k,2}) - f(d_{k,1}))$$

The agent does not know the valence of affordances used by the mobile entities and assumes it always follows the maximum utility value. Therefore, the comparison of two average contexts $\bar{A}_{a_i}^{t-1}$ and $\bar{A}_{a_i}^t$ is expected to yield a positive utility value, and a negative value is assumed to be an error in the estimated valences. In this case, the valence ν_j^i of an interaction j , for a mobile affordance designated by i , is increased when affordances of j move closer to the mobile affordance of i and decreased for those moving away from the mobile affordance. The learning rate decreases over time to allow a stabilization of the valences and also takes into account the distance between the agent and the mobile affordance a_i , which mitigates the impact of decisions of the mobile entity based on affordances beyond the visual range of the agent. The adjustment value is defined as $\alpha/n \times f(d_{\sigma_k})$, where α is the learning rate, f a strictly positive and decreasing function, d_{a_i} the distance of a mobile affordance, and n the number of updates.

B. Implementation

The valence learning system uses the following adjustment value: $10/n \times \exp(d_{a_i})$. To maximize the probability of identifying a prey and estimate valences, the agent’s behaviour is modified to follow the closest perception of a prey.

Results show that, in less than 500 prey observations, the estimation of valences converges toward a positive value for affordances of *slide* (algae) and negative for affordances of *bump* (walls), with a greater absolute value for *slide*. Table I shows a sample of valences obtained after 10 000 observations on five different runs. Since the behaviour model of the fish is exclusively based on walls and alga, only the sign and ratio

TABLE I
ESTIMATED VALENCES OF *bump* (WALLS) AND *slide* (ALGAE) FOR MOBILE AFFORDANCES OF *eat* (FISH), IN FIVE DIFFERENT RUNS.

Run	#1	#2	#3	#4	#5
alga	6.681	2.994	0.597	2.650	11.404
wall	-0.711	-0.170	-0.078	-0.178	-0.408
ratio	-9.397	-17.612	-7.654	-14.888	-27.950

between their related utility values are significant. For given values of 20 for algae and -2 for walls, we expect a ratio of -10 . This ratio can however vary around this value due to the discrete nature of movements.

Table I shows that the agent inferred that affordances of *eat* are attracted by affordances of *slide* and repulsed by affordances of *bump*, and the slide affordances have the most influence on the decision. The majority of runs result in a ratio between -5 and -20 . Some runs, such as #5, are biased by updates based on incomplete observations happening early in the training process when the learning rate is high. Other mechanisms will be investigated to mitigate such effects and improve the agent’s tolerance to environment changes.

VI. PREDICTING THE MOVEMENTS OF A MOBILE AFFORDANCE

As the agent can obtain the future positions and behavioural preferences of a mobile affordance, it becomes possible to define its next moves. The future positions of a mobile affordance a_i is given by the set of sequences $\Xi_{a_i}^t = \{\sigma_{a_i, k}\}_k^t$. For each sequence, it is possible to define a context of affordances $A^{\sigma_{a_i, k}, t}$, and define the utility value of each position from the current average context $\bar{A}_{a_i}^t$, noted $u_{\sigma_{a_i, k}} = u_{\bar{A}_{a_i}^t, A^{\sigma_{a_i, k}, t}}$. The expected next positions of the mobile affordance is then defined as sequences $\{\sigma_{a_i, k} | u_{\sigma_{a_i, k}} > 0\}$.

A. Predicting the Future Position of Mobile Affordance

A sequence $\sigma_{a_i}^t$ gives the expected position of the physical affordance a_i of i , but not the position σ_i allowing to enact i , and cannot be used recursively predict next positions. $\sigma_{a_i}^t$ must be converted into a set of interactions $C_{a_i} \subset I$ allowing to detect it. As the affordance is present in the environment, a single interaction in C_{a_i} is sufficient to characterize it. We thus propose to define an interaction $i \in I$ to characterize the most probable next position of a_i . We use the whole set of sequence $\Xi_{a_i}^t = \{\sigma_{a_i, k}\}_k^t$, from which two subsets are defined: $\Xi_{a_i}^{t,+}$ the subset of sequences for which $u_{\sigma_{a_i, k}} > 0$ and $\Xi_{a_i}^{t,-}$ the subset of sequences for which $u_{\sigma_{a_i, k}} < 0$.

First, for each sequence $\sigma_{a_i, k} \in \Xi_{a_i}^t$ we get the position $\sigma_{j, k}$ of j , alternative of i , that was used to define the position of the physical affordance of i (c.f. Sec. III-A). The position of j is defined as $\sigma_j | \sigma_{a_i} = \langle \sigma_j, j \rangle$. We then define the two sets $\Xi_j^{t,+}$ and $\Xi_j^{t,-}$ respectively from sequences of $\Xi_{a_i}^{t,+}$ and $\Xi_{a_i}^{t,-}$.

Next, we define the set of interactions designated by $\hat{S}_j^\sigma$ that can characterize the affordance of i . This set is defined as $\{j_k \in \hat{S}_j^\sigma | \exists \sigma_i, j_k \in \hat{S}_i^{\sigma_i}\}$. A probability value is defined for each interaction $i \in I$, by adding/subtracting probabilities defined by projected signatures Ξ_j^σ :

$$p_i = \sum_{\sigma \in \Xi_j^{t,+} | i \in \hat{S}_j^{\sigma}} p_{S_j^{\sigma}} - \sum_{\sigma \in \Xi_j^{t,-} | i \in \hat{S}_j^{\sigma}} p_{S_j^{\sigma}}$$

Finally, we get the interaction $i_{a_i}^{t+1} | p_{i_{a_i}^{t+1}} = \max_i(p_i)$ as the perception of a_i in its new position. As a_i is considered present, it becomes possible to locate it at a sequence σ when $i_{a_i}^{t+1} \in \hat{S}_j^{\sigma}$. Then, the prediction process can be repeated to define a sequence of future positions S_i^{t+k} . An interesting consequence of this principle is that the prediction of movements of a mobile affordance is independent of the movements of the agent: the successive positions S_i^{t+k} are defined in reference to other affordances, which contribute to the emergence of an allocentric reference.

B. Implementation

The prediction model was implemented in the agent and tested in different configurations. We used a valence value of 2 for *slide* interactions and -0.2 for *bump* interaction. Fig. 6 shows an example of configuration. The agent then enacts an interaction to perceive its environment, and generates predictions on the detected fish (*eat*). The recursive prediction is limited to ten steps. At each step, we get the interactions $i_{a_i}^{t+k}$ and the position of the affordance Ξ_i^{t+k} , allowing to observe the evolution of the position considered by the agent. In the configuration shown in Fig 6, the fish is located at a position implying a *turn left* interaction, then, after three simulation steps, in front of the agent, then, at a position implying a *turn right*. The sequence of positions is close to the expected movements of the fish: the agent successfully predicted that the affordance of *eat* will move toward the affordance of *slide*. When the predicted position of the fish arrives next to the alga, prediction stops or becomes inaccurate as the current prediction model still cannot detect the enaction of an interaction by other entities, and thus their consequences.

This evolution of the considered position will be used by future decisional mechanisms to generate interception behaviours instead of following behaviours.

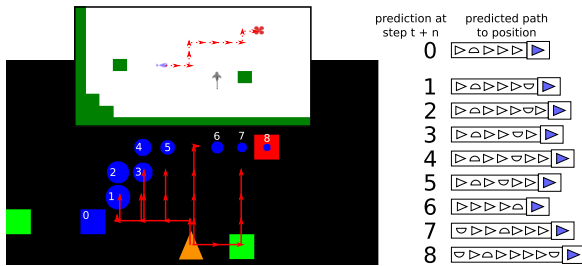


Fig. 6. Prediction of future positions of an affordance of *eat*. The blue circles represent the visual interactions characterizing the affordance at each prediction cycle. On the left, for readability reasons, we only represent, at each step, a single sequence leading in front of the prey, also represented with red arrows in the middle representation. The affordance, initially located through sequences including *turn left* (steps 0 to 5), is then located in front (step 6), then with sequences including *turn right*. The errors (steps 0 to 1 and 5 to 6) are mainly due to remaining errors in signatures. Although the estimated path slightly differs from the real path (dotted red line), the agent successfully predicted that the affordance of *eat* will move towards the affordance of *slide*.

VII. CONCLUSION

This work introduces a new mechanism enabling an agent to integrate mobile entities in its emergent model of the environment, and infer its decisional properties to enable the prediction of its future positions. This mechanism is spatial/temporal independent, and learns through interaction [9], and can thus be applied to vastly different contexts based on different modes of interaction. However, the predictions made by the agent are biased by its own decisional model, a limit that we will address in future works to improve the intersubjectivity possibilities between agents. Our next step is to study how an agent can employ the estimation of the behavioural model and predict the movements of a mobile entity and increase the complexity of its behaviours. Also, we will investigate methods enabling the integration of the agent itself in the estimated context of mobile entities, taking into account the influence of the agent on others. We intend to implement these mechanisms in multi-agent contexts to study the emergence of collaborative behaviours, such as coordinate hunting of large prey.

REFERENCES

- [1] K. R. Thorisson, A new constructivist ai: From manual methods to selfconstructive systems, in *Theoretical Foundations of Artificial General Intelligence*, P. Wang and B. Goertzel, Eds., vol. 4, pp. 145–171. Atlantis Press, 2012, Series Title: Atlantis Thinking Machines.
- [2] T. Froese and T. Ziemke, Enactive artificial intelligence: Investigating the systemic organization of life and mind, *Artificial Intelligence*, vol. 173, no. 3–4, pp. 466–500, January 2009.
- [3] J. Piaget, The construction of reality in the child, New York: Basic Books, 1954.
- [4] F. Kaplan P.-Y. Oudeyer and V. Hafner, Intrinsic motivation systems for autonomous mental development, in *IEEE Transactions on Evolutionary Computation*, vol. 11, no. 2, pp. 265–286, April 2007.
- [5] O. L. Georgeon, J. B. Marshall, and S. Gay, Interactional motivation in artificial systems: Between extrinsic and intrinsic motivation, in *IEEE Int. Conference on Development and Learning and Epigenetic Robotics (ICDL)*, pp. 1–2, 2012.
- [6] O. L. Georgeon and D. Aha, The radical interactionism conceptual commitment, *Journal of Artificial General Intelligence*, vol. 4, no. 2, pp. 31–36, December 2013.
- [7] M. Guillermin and O. L. Georgeon, Artificial interactionism: Avoiding isolating perception from cognition in ai, *Frontiers in Artificial Intelligence*, vol. 5, February 2022.
- [8] S. L. Gay, J. -P. Jamont, and O. L. Georgeon, Identifying and localizing dynamic affordances to improve interactions with other agents, *IEEE Int. Conf. on Development and Learning (ICDL)*, pp. 42–47, 2022.
- [9] S. L. Gay, O. L. Georgeon, A. Mille and A. Dutech, Autonomous construction and exploitation of a spatial memory by a self-motivated agent, *Cognitive Systems Research*, vol. 41, pp. 1–35, August 2016.
- [10] J. J. Gibson, The theory of affordances, Perceiving, Acting, and Knowing: Toward an Ecological Psychology, Hilldale, USA, vol. 1, no. 2, pp. 67–82, 1977.
- [11] A. Paola, P. Eric, L. S. Katrin, R. Subramanian, and P. P. A. Ronald, Building affordance relations for robotic agents - a review, in *International Joint Conference on Artificial Intelligence*, pp. 4302–4311, 2021.
- [12] P. Santana J. Baleia and J. Barata, On exploiting haptic cues for self-supervised learning of depth-based robot navigation affordances, *Journal of Intelligent and Robotic Systems*, vol. 80, February 2015.
- [13] E. Uğur and E. Şahin, Traversability: A case study for learning and perceiving affordances in robots, *Adaptive Behavior*, vol. 18, no. 3–4, pp. 258–284, 2010.
- [14] S. M. Nguyen A. Manoury and C. Buche, Hierarchical affordance discovery using intrinsic motivation, in *7th International Conference on Human-Agent Interaction*, AMC, pp. 186–193, 2019.
- [15] A. Netanyahu, T. Shu, B. Katz, A. Barbu, and J. Tenenbaum, Phase: Physically-grounded abstract social events for machine social perception, in *35th AAAI Conference on Artificial Intelligence*, March 2021.