



Synthesizing Evolving Symbolic Representations for Autonomous Systems

Gabriele Sartor¹ · Angelo Oddi² · Riccardo Rasconi² · Vieri Giuliano Santucci² · Rosa Meo¹

Received: 25 September 2024 / Accepted: 27 July 2025 / Published online: 21 August 2025
© The Author(s) 2025

Abstract

This work presents a novel architecture for an open-ended learning system that integrates intrinsic motivation (IM) and classical planning to enable agents continuously learn and improve their knowledge over time in an unsupervised fashion. The main goal is to allow the agent to autonomously distill its experience into Probabilistic Planning Domain Definition Language (PPDDL) terms, thereby making causal relationships explicit and supporting automated planning. Starting with a virtually empty set of predefined tasks or goals, the agent harnesses intrinsic motivation to explore the environment autonomously, continuously using and enriching the high-level knowledge acquired through its experience in a virtuous cycle. Experimental evaluation in the Treasure Game domain demonstrates the effectiveness of the proposed approach: starting with only a small set of primitive actions, we show how an agent can autonomously build and refine a high-level representation of the environment. Planning-based strategies grounded in this representation significantly outperform uninformed exploration by reaching intermediate sub-goals more efficiently and substantially reducing the time required to achieve the final objective.

Keywords Artificial intelligence · Abstraction · PPDDL · Open-ended learning

1 Introduction

In recent years, new AI systems have achieved remarkable success in solving complex tasks. These include real-world games such as chess [1] and Go [2–4], video games such as Atari [5], and Dota [6], as well as various robotics tasks [7–10]. These advances have largely been driven by intensive use of Reinforcement Learning (RL, [11]) combined with the resurgence of neural networks and deep learn-

ing technologies [12]. Typically, “standard” RL focuses on learning policies that maximize the completion of fixed, predefined tasks through reward optimization, using a limited set of predefined skills. However, newer approaches inspired by neuroscience and psychology have been proposed to enrich RL by enabling agents to expand their capabilities over time. Studies in animals [13–15] and humans [16–18] have highlighted an intrinsic drive toward novelty, further supported by neuroscientific research [19–21]. The field of intrinsically motivated open-ended learning (IMOL [22]) addresses the challenge of developing agents that aim to enhance their ability to interact with the environment without being assigned specific tasks. More precisely, Intrinsic Motivations (IMs [23, 24]) are a class of self-generated signals used to autonomously guide agents in various processes, from exploring the state and action spaces [25, 26], to autonomously discovering, selecting, and learning multiple goals [27–29]. In general, IMs help agents acquire new knowledge independently, or even in the absence, of any external task, thereby supporting open-ended learning processes [30]. This knowledge can later be leveraged to solve user-assigned tasks [31] or serve as a scaffold for acquiring further knowledge incrementally [32–34], as in curriculum learning [35]).

✉ Gabriele Sartor
gabriele.sartor@unito.it

Angelo Oddi
angelo.odd@istc.cnr.it

Riccardo Rasconi
riccardo.rasconi@istc.cnr.it

Vieri Giuliano Santucci
vieri.santucci@istc.cnr.it

Rosa Meo
rosa.meo@unito.it

¹ Computer Science Department, University of Turin, Via Verdi 8, Turin 10124, Italy

² Institute for Cognitive Sciences and Technologies, Via G. Romagnosi 18A, Rome 00196, Italy

The option framework has been combined with IMs and “curiosity-driven” approaches to support both the learning [32] and discovery of options [36–39]. In the hierarchical RL setting [40], where agents need to sequence various options to accomplish complex tasks, IMs have been used to facilitate the discovery and learning of subtasks [41–43], and improve exploration [26]. Learning and composing different skills autonomously is essential for agents acting in complex environments, where task completion involves achieving several (possibly unknown) interdependent subtasks. An increasing body of research addresses this problem [29, 44, 45], with most work focused on low-level, sub-symbolic policy learning [46], often composed hierarchically through meta-policies [47]. While promising, these approaches inherently suffer from scalability issues, as exploration becomes less efficient with the increasing size of the state-action space.

In contrast, symbolic approaches such as Automated Planning [48, 49] use high-level representations (symbols), allowing faster execution, easier composition of complex task sequences, and improved interpretability of the agent’s knowledge. However, these methods typically require a high-level representation of the planning domain to be specified in advance. Planning usually depends on prior knowledge of the environment, expressed as both preconditions for actions and their corresponding effects. This reliance on hand-crafted symbolic representations limits the applicability of symbolic planning in unknown or highly unstructured environments, where new knowledge and skills are acquired autonomously through exploration. Some studies have attempted to address this limitation by incorporating autonomous model learning [50], human-assisted symbol augmentation [51] or cognitive layers for managing intermediate representations [52].

More recently, the literature has proposed methodologies to integrate sub-symbolic and symbolic approaches, or more broadly, low-level and high-level modules [53]. On one hand, several works have sought to reconcile deep learning with planning [54, 55], goal recognition [56] and symbolic domain synthesis [57, 58]. On the other hand, specific algorithms have been developed to automatically abstract low-level information acquired through exploration into high-level planning representations [59], such as those based on the PDDL formalism [60], which uses symbols to explicitly define the *preconditions* and *effects* of actions. This algorithm has been used as a component in architectures combining abstraction, planning, and intrinsic motivation, such as IMPACT [37, 61]. Although it is widely accepted that multiple levels of abstraction are necessary to approach Artificial General Intelligence (AGI), the most appropriate representation to adopt remains an open question.

Our approach builds upon previous work [50, 51, 54] by synthesizing symbolic abstractions directly from raw sensor data using a dedicated abstraction algorithm [59], rather than relying on predefined structures. While [57, 58] demon-

strated that deep learning can generate planning-suitable representations for limited domains and novel tasks, they do not tackle the challenge of incrementally enriching this knowledge, e.g., by incorporating intrinsic motivation. We also extend the ideas presented in [37–39] by implementing a continuous discover-plan-act loop, allowing us to observe how the high-level representation evolves and how this evolution affects both exploration and goal achievement.

This study proposes an approach for integrating low-level skills with high-level representations, enabling the continuous update of the agent’s capabilities and their abstract representations. The acquisition and expansion of the agent’s knowledge are driven by two forms of intrinsic motivation: (i) one that encourages the agent to learn new policies while exploring the environment, and (ii) another that pushes it to apply its skills to reach under-explored states, on the rationale that unfamiliar states are more likely to yield new learning opportunities. Sensor data collected before and after executing these skills are processed by a specific algorithm that updates the agent’s abstract symbolic representation. This abstraction is then used to plan sequences of low-level actions to achieve increasingly complex goals. The main contribution of this work is a framework that, starting from only a small set of primitive actions and no symbolic knowledge, autonomously builds a symbolic abstraction from raw sensor data. This evolving abstraction supports planning mechanisms that enable the agent to solve progressively more complex tasks.

2 Background

To achieve a high level of autonomy, an agent operating in the low-level space, perceiving the environment through its sensors and interacting with it via its actuators, must construct multiple layers of abstraction over its state and action spaces. Just as humans reason using both simple and complex concepts in their daily activities, robots must also be capable of building abstract representations of the world to manage increasing complexity. These abstractions allow agents to assign labels to actions and events, making them identifiable and subject to reasoning. In this paper, we apply two levels of abstraction: the first maps primitive actions to *options* [11, 62] and the second elevates options to the level of classical planning representations [49].

2.1 From primitives to options

As discussed before, at the lowest level, the agent sees the world with its sensor’s values and changes it through the movement of its actuators. The most common formalism at this stage to deal with this type of representation is the Markov Decision Process (MDP), which models the envi-

ronment as the tuple:

$$(S, A, R, T, \gamma), \quad (1)$$

in which S represents the set of possible high-dimensional states where each $s \in S$ is described by a vector of real values returned by the agent's sensors, A describes the set of low-level actions $a \in A$ in some cases also called *primitives*, R the reward function where $R(s, a, s')$ is a real value returned executing a from state s achieving s' , T the transition function describing for $T(s'|s, a)$ the probability of reaching the state s' executing a from s , and the discount factor $\gamma \in (0, 1]$ describing the agent's preference for immediate over future rewards. Usually, in this setting, the goal is to maximize *return*, defined as

$$\mathcal{R} = \sum_{i=0}^{\infty} \gamma^i R(s_i, a_i, s_{i+1}). \quad (2)$$

However, dealing with the state and action spaces of the formulation (1) is, in certain cases, impractical due to the high dimensional spaces considered. An effective formalism introduced to reduce the complexity of the problems is the *option* framework [62]. The *option* is a temporally-extended action definition which employs the following abstracted representation:

$$o = (I_o, \pi_o, \beta_o), \quad (3)$$

where the option o is defined by an *initiation set* $I_o = \{s | o \in O(s)\}$ representing the set of states in which o can be executed, a *termination condition* $\beta_o(s) \rightarrow [0, 1]$ returning the probability of termination upon reaching s by o , and a policy π_o which can be run from a state $s \in I_o$ and terminated reaching s' such that the probability $\beta_o(s')$ is sufficiently high. A *policy* is a function defining the behavior of the agent, mapping the perceived state of the environment to the action to be taken [11]. It is worth noting that options create a temporally-extended definition of the actions [62]. Indeed, the option is an abstraction defining an action as a repeated execution of policy π from a state $s \in I_o$ to another s' in a maximum amount of time steps τ .

For instance, Figure 1 depicts an environment before and after executing the option “*reach the ball*”, respectively Figure 1a and 1b. When the agent wants to execute this option, it checks whether its current state belongs to the *initiation set* of the option. In our case, and assuming that the option's policy simply consists in moving towards the ball, the agent can execute the option and runs the policy π_o until the *termination condition* returns a sufficiently high probability of success or τ time steps are reached. In the Figure 1b the options successfully terminates getting close to the ball.

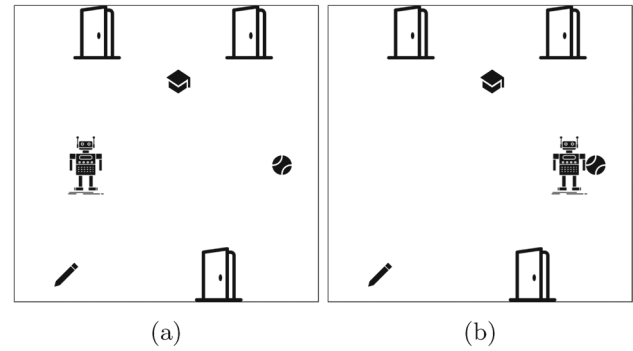


Fig. 1 Representation of the option *reach the ball*

Passing from low-level actions to options reduces the agent's action space. Adopting options in the MDP formalism implies moving to the semi-Markov Decision Process (SMDP):

$$(S, O, R, P, \gamma), \quad (4)$$

where S is the original state space, $O(s)$ is the set of options executable from state s , $R(s', \tau | s, o)$ describes the reward expected executing $o \in O(s)$ from state s reaching s' after τ time steps, $P(s', \tau | s, o)$ returns the probability of reaching state s' after τ time steps executing $o \in O(s)$ from state s , and the discount factor $\gamma \in (0, 1]$. Using options entails moving in the low-level state space S with abstracted actions permitting the agent to perform extended and more complex behaviors, reducing the number of actions to achieve a certain task and simplifying the problem.

2.2 Options and Classical Planning

The option formalism and its way of abstracting the dynamics of an environment share common characteristics with *classical planning*, in which the world is described in a simplified formal description considering only the aspects necessary to solve the agent's task [49]. Planning is the field of research studying formal methods to automatically find solutions, also called plans, to tasks requiring a sequence of actions $[\alpha_1, \dots, \alpha_n]$ to reach a goal state s_g from an initial state s_{init} . The plan ω is obtained by giving in input a model of the environment dynamics and the problem definition to a planner, returning a solution applying general optimization algorithms.

Classical planning is a particular instance of this methodology exploiting a logical language (i.e. propositional logic, first-order logic, etc.) to capture an abstract symbolic description of the world [49]. The *symbol*, which is the core of classical planning, is a name given to a certain set of states $s \in S$ satisfying a specific condition in the sensorimotor space. This *mapping* from the symbol to the real world states

is called *grounding*. Specifically, a symbol $\sigma_Z \in \Sigma$, with Σ set of the available symbols, is the name given to a *test* τ_Z , and the corresponding set of low-level states where the test is satisfied $Z = \{s \in S \mid \tau_Z(s) = \text{True}\}$, with the high-dimensional low-level state space S [59]. To determine whether or not a low-level state s_i belongs to the state set Z , the $\sigma_Z(s_i)$ test is run, which will return either *True* or *False*. Again, both Z and σ_Z provide the *semantics* of the symbol; Z is the symbol's *grounding set*, while σ_Z is its *grounding classifier*. Symbols can be combined using *operators* having the following meaning:

- $\neg\sigma_Z$ corresponds to the *negation* of symbol σ_Z ;
- $\sigma_X \vee \sigma_Y$ corresponds to the *union* of symbols σ_X and σ_Y (i.e., the union of their respective *grounding sets*);
- $\sigma_X \wedge \sigma_Y$ corresponds to the *intersection* of symbols σ_X and σ_Y (i.e., the intersection of their respective *grounding sets*).

In classical planning, operators and symbols are used to describe actions in the following form

$$\alpha_i = (\text{pre}_i, \text{eff}_i^+, \text{eff}_i^-), \quad (5)$$

meaning that the action $\alpha_i \in \mathcal{A}$ can be executed when all the symbols $\{\sigma \mid \sigma \in \text{pre}_i\}$, also called *preconditions*, are *True*, and executing α_i produces the changes of the value of some symbols, also called *effects*, implying that all symbols $\{\sigma \mid \sigma \in \text{eff}_i^+\}$ assume the value *True* and symbols $\{\sigma \mid \sigma \in \text{eff}_i^-\}$ become *False*.

Finally, using symbols Σ and high-level actions $\mathcal{A}$ as building blocks, it is possible to define the model of environment $\mathcal{D}$, also called *domain*, and the *problem* $\mathcal{P}$ to solve. A classical planning domain can be defined as

$$\mathcal{D} = (\Sigma, \mathcal{A}, \Gamma), \quad (6)$$

using a set of symbols Σ , actions $\mathcal{A}$ and a *state-transition function* $\Gamma : \hat{\Sigma} \times \mathcal{A} \rightarrow \hat{\Sigma}$, where $\hat{\Sigma}$ is the set of possible subsets of Σ . The *state-transition function* $\Gamma(\hat{\Sigma}_s, \alpha_i) = (\hat{\Sigma}_s - \text{eff}_i^-) \cup \text{eff}_i^+$, if α_i is applicable to $\hat{\Sigma}_s$, where $\hat{\Sigma}_s$ is the set of symbols whose *grounding set* intersection defines the state $s \in S$. These elements are sufficient to describe the dynamics of the environment, over which the planner can reason and create chains of actions to reach a final goal. The function Γ encapsulates the transition model of the environment in each action model described by (5), defining the possible way to build sequences of them. Instead, the problem can be formalized as

$$\mathcal{P} = (\hat{\Sigma}_{s_{init}}, \hat{\Sigma}_{s_g}), \quad (7)$$

where $\hat{\Sigma}_{s_{init}}$ is the set of symbols whose *grounding set* intersection describes s_{init} and $\hat{\Sigma}_{s_g}$ the set of symbols whose *grounding set* intersection characterizes s_g . Then, the plan solving the problem $\mathcal{P}$ is expressed as the sequence of actions

$$\omega = [\alpha_1, \dots, \alpha_n], \quad (8)$$

resulting in a sequence of state transitions

$$[\hat{\Sigma}_{s_0}, \hat{\Sigma}_{s_1}, \dots, \hat{\Sigma}_{s_n}], \quad (9)$$

such that $\Gamma(\hat{\Sigma}_{s_0}, \alpha_1) = \hat{\Sigma}_{s_1}$, $\Gamma(\hat{\Sigma}_{s_1}, \alpha_2) = \hat{\Sigma}_{s_2}, \dots$, $\Gamma(\hat{\Sigma}_{s_{n-1}}, \alpha_n) = \hat{\Sigma}_{s_n}$ with initial state $\hat{\Sigma}_{s_0} = \hat{\Sigma}_{s_{init}}$ and final state $\hat{\Sigma}_{s_n} = \hat{\Sigma}_{s_g}$. In particular, the formulation presented based on a finite set of symbols Σ and state-transition system $\mathcal{T} = (\Sigma, \mathcal{A}, \Gamma)$ is called a *set-theoretic representation* [49]. In order to use classical planning systems, it is necessary to describe both the domain of the considered world and the tackled problem using a planning definition language. In this work we will use the Probabilistic Planning Domain Definition Language (PPDDL) [63]; once defined, both the domain and the problem definitions will be provided in input to the planner, which will eventually return a solution ω .

It is worth noting that the option o and the planning action α share some similarities. Indeed, o can be converted in its corresponding high-level action α finding the right set of symbols $\text{pre}, \text{eff}^+, \text{eff}^-$ whose grounding sets satisfy the classifiers I_o as preconditions and β_o as effects, for the execution of π_o . The similarity between options and set-theoretic planning actions has been exploited to create automatic abstraction procedures able to convert the options' execution data into a working high-level planning description, thus allowing a solution that integrates low-level and high-level information. In this work we build upon the abstraction procedure implemented by Konidaris et al. [59] to convert sensors raw data into a PPDDL representation.

2.3 Intrinsic Motivation

The impulse to drive the agent away from the monotony of its usual activities, which psychologists and cognitive scientists have studied under the name of *intrinsic motivation*, is one of the most important elements enabling Open Ended Learning (OEL). The research in the field of Intrinsic Motivation (IM) concerns the study of human behaviors not influenced by external factors but characterized by internal stimuli (i.e. curiosity, exploration, novelty, surprise). In the case of artificial agents, we can summarize such aspect as anything that can drive the agent's behavior which is not directly dependent on its assigned task.

The insights provided by the IMs gave the researchers new ideas to model the stimuli of the agent (e.g. curiosity). Indeed, some models have been implemented using the prediction

error (PE) in anticipating the effect of agent's actions (and more precisely the improvement of the prediction error [64, 65]) as an IM signal. A formal definition of agent driven by its curiosity has been formulated by Schmidhuber [64] as simply maximizing its *future success* or *utility*, which is the conditional expectation

$$u(t) = E_{\mu} \left[\sum_{\tau=t+1}^T r(\tau) \middle| h(\leq t) \right], \quad (10)$$

over $t = 1, 2, \dots, T$ time steps, receiving in input a vector $x(t)$, executing the action $y(t)$, returning the reward $r(t)$ at time t , taking into consideration the triplets $h(t) = [x(t), y(t), r(t)]$ as the previous data experienced until time step t (also called *history*). The conditional expectation $E_{\mu}(\cdot|\cdot)$ assume an unknown probability distribution μ from M representing possible probabilistic reactions of the environment. To maximize (10), the agent also has to build a predictor $p(t)$ of the environment to anticipate the effects of its actions. The reward signal is defined as follows

$$r(t) = g(r_{ext}(t), r_{int}(t)), \quad (11)$$

which is a certain combination g of an external reward $r_{ext}(t)$ and an intrinsic reward $r_{int}(t)$. In particular, $r_{int}(t+1)$ is seen as *surprise* or *novelty* in assessing the improvements in the results of p at time $t+1$

$$r_{int}(t+1) = f|C(p(t), h(\leq t+1)), C(p(t+1), h(\leq t+1))|, \quad (12)$$

where $C(p, h)$ is a function evaluating the *performance* of p on a history h and f is a function combining its two parameters (e.g. in this case, it could be simply the improvement $f(a, b) = a - b$). It is important to notice that, as a baby does, an intrinsically motivated agent needs to find regularities in the environment to learn something. Consequently, if something does not present a *pattern*, there is nothing to learn and this becomes boring for both an agent and a human being.

In literature, IMs have also been categorized into different typologies [66–68]. An important discriminant aspect is the kind of signal received by the agent which can be of two types: *knowledge-based* (KB-IMs), which depends on the prediction model of the world (e.g. [64]), and *competence-based* (CB-IMs), which depends on the improvement of the agent's skills (e.g. [27]). In the framework presented in the next section, both these typologies are employed. CB-IM is used at a lower level to learn new skills and KB-IM at higher level to push the system to focus on the frontier of the visited states, from which it is more likely to discover novel information (e.g., find new states and learn new actions).

3 System Overview

This section presents a new framework of an open-ended learning agent which, starting from a set of action primitives, is able to (i) discover options, (ii) explore the environment using them, (iii) create a PPDDL representation of the collected knowledge and (iv) plan to improve its exploration while reaching a high-level objective set by the game.

The aim of this study is to assess the potential of *abstraction* in autonomous systems and propose a new approach for planning systems, extending them with learning capabilities and behaviors driven by IMs. IMs are employed for discovering new options in a *surprise-based* manner at low-level and continuously exploring new states driven by *curiosity* at high-level. By the term abstraction, we basically mean mapping a certain problem space into a simpler one (e.g. converting a continuous domain into a more coarse, discrete domain). In the proposed architecture, the abstraction is applied at two levels: a) passing from the *primitive action space* to the *options action space* and b) converting *low-level data* collected during the exploration into a *high-level domain representation* suitable for high-level planning, thus from raw sensors' data to a PPDDL representation.

Performing the pipeline depicted in Figure 2, the system creates different layers of abstraction, enriching the agent's knowledge with causal correlations between options and enabling more efficient reasoning (i.e. using classical planning). Symbols can be seen as knowledge building blocks that can be used to search for interesting states and find new knowledge in a virtuous loop.

3.1 Architecture description

As depicted in Figure 2, the system can be seen as a three-layered architecture: (i) the higher level contains the explicit agent's knowledge, (ii) the middle layer maps the high-level actions to their controllers and convert the raw sensors data into an explicit representation, and (iii) the lower level containing the components to sense the environment and interact with it. Mainly, the system executes the following pipeline:

1. **Option Discovery:** using a set of primitives $A = \{a_1, a_2, \dots\}$ belonging to the agent, the system combines them to create higher level actions, in this case, the *options*;
2. **Exploration:** the set of options $O = \{o_1, o_2, \dots\}$ discovered in the previous step is used to explore the environment. In the meantime, the visited low-level states data are collected into two datasets: the *initiation data ID* and *transition data TD* containing, respectively, the samples of states in which an option can be run and the transition between states executing it;

no longer be executed or a^t becomes available, (iv) I is the set of states from which a^p can run; and (v) β is the termination condition of the option, corresponding to the availability of the primitive a^t or to the impossibility of further executing a^p . For the sake of simplicity, in the remainder of the paper the option's definition will follow the more compact syntax

$$o(a^p, a^t) \quad (14)$$

meaning that I is the set of states in which a^p can run, β checks the following two conditions: a^t becomes available or a^p is no longer available, and π is the policy corresponding to repeatedly executing a^p until β verifies.

Algorithm 2 describes in further details the option discovery procedure previously described. At the beginning of each discovery episode, the plan ω^{EX} is executed to reach a new area where to start learning new options (line 7–9). Then, for a maximum number of episodes and steps, the agent saves the current state s and randomly selects a primitive a^p which can be executed in s (line 10–12). a^p is repeatedly executed towards reaching the state s' until either a^p is no longer available or new primitives beyond a^p become available (line 13). If $s \neq s'$, the procedure creates a new option o where $o = o(a^p, a^t)$ if a new primitive a^t has become available (line 20), or $o = o(a^p, \{\})$ in the opposite case (line 22). In either way, in case o has not been discovered before, it is added to the other collected options (line 24). It is important to note that the options are independent on the state where the agent is, and are defined by the primitives' availability. This definition makes options reusable on different floors and with different objects, just depending on the agent's abilities. The procedure can discover a subset of the options available in the original implementation¹ sufficient to solve the entire game problem.

3.1.2 Exploration

After discovering a set of valid options O as explained in the previous section, the system exploits them to explore the environment, collecting data about the reached low-level states (Algorithm 1, line 10). Considering that the abstracted representation of the world does not change significantly with a small amount of new data, the function *Collect_Data()* is in charge of executing d_steps options for d_eps episodes, in which the agent starts its exploration from the initial configuration of the environment.

At each timestep, the agent attempts to perform an option $o \in O$ from a certain low-level state $s \in S$. The selection of the action o during the exploration can follow different strategies, which are described in the subsection 3.1.4. In case

Algorithm 2 Option Discovery

```

1: procedure OPTION_DISCOVERY( $d\_eps, d\_steps, \omega^{EX}$ )
2:    $O_{new} \leftarrow \{\}$ 
3:    $ep \leftarrow 0$ 
4:   while  $ep < d\_eps$  do //For each episode
5:      $T \leftarrow 0$ 
6:     Reset_Game()
7:     for  $option$  in  $\omega^{EX}$  do // Execute IM plan
8:       Execute(option)
9:     end for
10:    while  $T < d\_steps$  do //For each step
11:       $s \leftarrow Get\_State()$ 
12:       $a^p \leftarrow Get\_Available\_Primitive()$ 
13:      while  $Is\_Available(a^p)$  and not ( $New\_Available\_Prim()$ ) do
14:        Execute( $a^p$ )
15:         $s' \leftarrow Get\_State()$ 
16:      end while
17:      if  $s \neq s'$  then
18:        if  $New\_Available\_Prim()$  then
19:           $a^t \leftarrow Get\_New\_Available\_Prim()$ 
20:           $o \leftarrow Create\_New\_Option(a^p, a^t)$ 
21:        else
22:           $o \leftarrow Create\_New\_Option(a^p, \{\})$ 
23:        end if
24:         $O_{new} \leftarrow O_{new} \cup o$ 
25:      end if
26:       $T \leftarrow T + 1$  //End For each step
27:    end while
28:     $ep \leftarrow ep + 1$  //End For each episode
29:  end while
30:  return  $O_{new}$ 
31: end procedure

```

the execution of the option changes the low-level variables of the state s and, consequently, the *mask*² m is not null, the system registers two types of data tuple (Algorithm 1, line 10): the *initiation data* tuple id and the *transition data* tuple td . The multiple instances of these tuples are stored, respectively, in the datasets ID , for the initiation data, and TD , for transition data. A single *initiation data* tuple id_i has the following structure

$$id_i = (s, o, f(s, o)), \quad (15)$$

where the function $f(s, o)$ returns the feasibility of executing o from s (*True* if $s \in I_o$ and *False* otherwise). The *transition data* tuple td_j takes the following structure

$$td_j = (s, o, r, s', g, m, O'), \quad (16)$$

where s' is the state reached after executing option o from the state s , g is a flag stating whether the final objective of the task has been reached, m is the *mask* of the option and O' is a list defining the options that can be executed from s' . When all the steps of the episode have been executed, the environment is reset and the next episode is started until reaching the maximum number of allowed episodes d_eps , where the *Collect_Data()* procedure terminates and the stored data are added to the existing datasets ID and TD (line 11–12).

¹ Link to the original implementation of [59]: <https://github.com/sd-james/skills-to-symbols>

² The *mask* is the list of all the state variables that are changed by the execution of a specific option. See details in [59].

3.1.3 Abstraction

The datasets collected in the previous step are then used as input for the function *Create_PPDDL()* (Algorithm 1, line 13), returning a symbolic representation $\mathcal{D}$ of the agent's current knowledge expressed in PPDDL formalism (PPDDL domain). The main advantage of the obtained PPDDL representation is that it makes explicit the causal correlations between operators that would have remained implicit at the option level. In the following, we provide a summary description of the abstraction procedure; for further details, the reader is referred to the original article [59].

The abstraction procedure executes the following five steps:

1. **Options partition:** this step is dedicated to partitioning the learned options into *abstract subgoal options*³, a necessary assumption of the abstraction procedure. Abstract subgoal options are characterized by a single precondition and effect set. However, given the uncertainty of the actions' effects in the environment, the operators' effects will be modelled as *mixture distributions* over states. This phase utilizes the transition dataset TD collected before, as it captures the information about the domain segment the option modifies. Basically, the transition dataset is divided into sets of transition tuples presenting the same option o and mask m . Then, for each set, the partitions are ultimately obtained by performing clustering on the final states s' through the DBSCAN algorithm [69]. If some clusters overlap in their initiation set, a unique partition is created with different effects and occurrence probabilities.
2. **Precondition estimation:** this step is dedicated to learning the classifiers that will identify the *preconditions* of the PPDDL operators. In order to have negative examples of the initiation set classifier for each operator, this operation utilizes the initiation dataset ID considering all the samples with option o and $f(s, o) = \text{False}$. The positive examples comprise instead the initial states s taken from TD tuples belonging to the same partition. The initiation set classifier of the option is computed using the Support Vector Machines (SVM) [70]. The output of this phase is the set of all the *initiation set* classifiers of all the operators.
3. **Effect estimation:** analogously, this step is dedicated to learning the symbols that will constitute the *effects* of the PPDDL operators. The effects distribution is modelled

through the Kernel density estimation [71, 72], taking in input the final states s' of each partition.

4. **PPDDL Domain synthesis:** finally, this step is dedicated to synthesizing the PPDDL domain, characterized by the complete definition of all the operators associated with the learned options in terms of preconditions and effect symbols. This step entails the simple mapping of all the data structures generated during the previous steps in terms of symbolic predicates to be used as preconditions and/or effects for every operator.

The produced PPDDL domain can be potentially used to reach any subgoal that can be expressed in terms of the available generated symbols at any point during the Discovery-Plan-Act (DPA) loop. One interesting aspect of the proposed DPA framework is that the semantic precision of the abstract representation and its expressiveness increase as the DPA cycles proceed, as will be described in the experimental section 4.

3.1.4 Goal Selector

This module aims at simulating the *intrinsic motivations* driving the agent towards interesting areas to satisfy its *curiosity* and optimize its exploration. In particular, one of the most fascinating aspects of this system is the capability of setting its high-level goals, which potentially could be a combination of symbols defining a state that the agent has never experienced before. In other words, abstract reasoning could be the driving criterion for *using the imagination to explore an unknown environment*. Despite the previous goal is rather ambitious and still the object of future work, we will demonstrate in this work that the abstract reasoning can indeed be used for the more “down-to-earth” task of devising rational criteria to make more efficient the exploration of unexplored parts of the environment.

The selection of the target state $s_{\text{target}} \in S$ to be reached in the next exploration cycle is performed at line 14 of Algorithm 1, calling the procedure *Get_Target_State()*. The *Goal Selector* suggests such state to the system, following an internal strategy which can be, in this case, *Action Babbling (AB)*, *Goal Babbling (GB)* and *Distance-based Goal Babbling (DB)*.

Action Babbling is the simplest strategy of the system for the exploration, consisting of a pure random walking of the agent. This strategy returns $s_{\text{target}} = \text{NULL}$, so no plan ω^{EX} is generated and executed in the exploration phase on the subsequent cycle. *Goal Babbling* and *Distance-based Goal Babbling* differ in the way they implement IMs (such as curiosity). More specifically, in our system curiosity is formalised as an interest to reach (i) a random goal among those already achieved, and (ii) the border of the already acquired knowledge.

³ An abstract subgoal option o is characterized by a list of indices of the low-level variables, called *mask*, which are changed with the execution of o , without modifying other variables. In addition, the changing variables' values do not depend on their initial value.

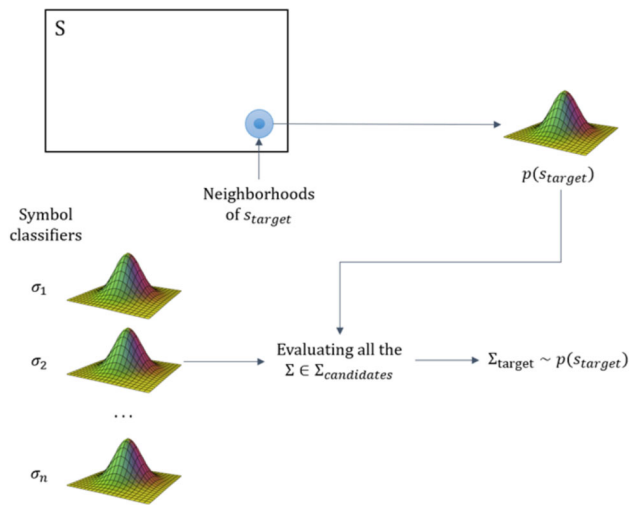


Fig. 3 The conceptual depiction of the selection of the symbols necessary to represent the goal s_{target} . The set of symbols having the distribution most similar to the goal will define the goal of the PPDDL problem file

Goal Babbling consists in randomly selecting a low-level state in the environment and trying to reach it [73]. Usually, the assumption of *Goal Babbling* is that all the goals which can be set belong to the world's low-level states that are reachable; each goal is formalised as the configuration of joints or position to be reached with the robot's actuators [27]. Since in general not all the states $s \in S$ are valid (e.g. the agent cannot find itself inside a wall), this strategy selects a random state s_{target} among the already visited ones (line 14). Subsequently, s_{target} is translated into a set of propositional symbols $\{\sigma_1, \dots, \sigma_k\}$, as described in the following subsection 3.1.5, which represent the high-level goal to be reached using an off-the-shelf PPDDL planner. The capability of translating low-level states into symbols gives the agent a chance to reason on causal dependencies and, consequently, perform planning. It is important to notice that a pure *Goal Propositional Babbling*, consisting of the selection of a random subset of high-level symbols, would not be an effective strategy because only a limited number of combinations of symbols conjunctions are valid goals. Consequently, *Goal Propositional Babbling* is not taken into consideration.

Distance-based Goal Babbling is implemented as a modified version of the *Goal Babbling*, and models the curiosity towards the less explored states as being influenced by the goal's distance from its starting location s_{init} . In this case, the curiosity level for a state is defined as

$$\eta(s) = \|s_{init} - s + \mathcal{Z}\|, \quad (17)$$

consisting of the norm of the difference between the state vector of the agent at the beginning of each episode s_{init} and the visited state s . The low-level s_{target} state is selected

between the farthest visited states, as follows:

$$s_{target} = \max_s \eta(s), \forall s \in S_{visited}. \quad (18)$$

More precisely, the low-level state s , target of the exploration, is the state that maximizes the distance $\eta(s)$. In the computation of s_{target} , a Gaussian noise $\mathcal{Z} \sim \mathcal{N}(0, 1)$ is added, to facilitate the reaching of different states. The *Goal Selector* driven by η continuously pushes the agent towards the border of the already acquired knowledge (similarly to the idea presented in [74, 75], but with a different implementation) and is thus the main responsible for the knowledge increase of the agent. Moving towards the frontier states, the agent is more likely to encounter novel states thus maximizing the probability to acquire new symbolic knowledge, which can in its turn be used to synthesize new (e.g., more expressive) high-level goals to be eventually reached through planning, ultimately creating a virtuous loop in which the agent's capabilities of increasing its knowledge through environment exploration are iteratively enhanced.

3.1.5 Translating low-level states into symbols

The peculiarity and, simultaneously, the biggest challenge of this software architecture is thus to use an abstract symbolic representation to manage the evolution of the agent's knowledge. Using a symbolic representation to describe a desired environment configuration is a powerful tool for an efficient exploration of the world.

After selecting s_{target} (line 14), the purpose of the planning module is twofold. On the one hand, it can be employed as usual to verify the reachability of the goal s_g of the environment (i.e. get the treasure and bring it "home" in the Treasure Game, in line 17) and, on the other hand, it can be exploited to generate a plan driving the agent towards states, in our case s_{target} , relevant to extend the knowledge of the system (line 16). In both cases, to take advantage of planning it is necessary to transform a low-level state into a high-level one. This operation requires finding the combination of propositional symbols

$$\hat{\Sigma}_{target} = \{\sigma_1, \dots, \sigma_k\} \quad (19)$$

that best represent the portion of state space including s_{target} . In other words, we look for a subset of symbols $\hat{\Sigma}_{target}$ whose *grounding* is s_{target} (See Figure 3). These symbols make it possible to generate the definition of a planning problem (line 15) which, together with the domain definition, can be used to perform planning and solve the problem (line 16).

In order to select the right symbols conjunction, the system creates the classifier $Cl_{target} \sim p(s_{target})$ approximating a distribution over s_{target} . Cl_{target} is a SVM classifier trained

on the states

$$\{s \mid \|s_{target} - s\| \leq \epsilon\}, \forall s \in S_{visited}, \quad (20)$$

representing, as positive samples, all the neighbours of s_{target} within a maximal distance ϵ and, as negative samples, all the remaining states encountered in the agent's experience. Once the classifier is generated, all the propositional symbols whose *mask* is equal to a *factor*⁴ contained by Cl_{target} are collected as candidates $\hat{\Sigma}_{candidates}$. Then, all the subsets of symbols $\hat{\Sigma}_i \subset \hat{\Sigma}_{candidates}$, whose respective masks do not overlap are evaluated as representations of s_{target} . From each $\hat{\Sigma}_i$, m state samples $S_{\Sigma_i} = \{s_1, \dots, s_m\}$ are generated and the score of the subset $\hat{\Sigma}_i$ is calculated as

$$score(\Sigma_i) = \frac{1}{m} \sum_{s \in S_{\Sigma_i}} Cl_{target}(s) \quad (21)$$

where $Cl_{target}(s)$ returns the probability that s belongs to the positive class of the classifier Cl_{target} . Then, the subset of symbols $\hat{\Sigma}$ maximizing the *score* function is used as a goal in the problem definition $\mathcal{P}_{target}$.

3.1.6 Planning

At the end of each cycle, the planning process generates a plan to reach either s_{target} and s_g . In both cases, in order to create a PDDL problem $\mathcal{P}$, it is necessary to find the set of symbols $\hat{\Sigma}_{init}$, $\hat{\Sigma}_g$ and $\hat{\Sigma}_{target}$, describing the most suitable high-level state representation for s_{init} , s_g and s_{target} respectively, as described in the previous subsection 3.1.5. Indeed, the couples $(\hat{\Sigma}_{init}, \hat{\Sigma}_g)$ and $(\hat{\Sigma}_{init}, \hat{\Sigma}_{target})$ define the problems $\mathcal{P}_g$ and $\mathcal{P}_{target}$. At line 15 of Algorithm 1, $\mathcal{P}_{target}$ is generated as previously discussed and the plan to solve it, ω^{EX} , is generated by the planner (line 16).

Before moving to the next cycle, the system tries to solve also the problem $\mathcal{P}_g$, performing the function *Check_PPDDL_Vailidity* (line 17). The resulting plan ω^g is only used internally by the system to keep track of the success ratio of the planner with the evolution of the synthesized knowledge of the agent.

4 Experiments and Results

This section, dedicated to the experimental analysis, will first describe the dynamics of the environment, followed by an example of the system's cycle execution with its outputs and, finally, the overall results collected over different static and fully observable environment configurations.

⁴ "Sets of low-level state variables that, if changed by an option execution, are always changed simultaneously" [59].

4.1 Environment setup

The implemented system has been tested in the so-called Treasure Game domain [59]. In such an environment, an agent can explore the maze-like space by moving through corridors and doors, climbing stairs, interacting with handles (necessary to open/close the doors), bolts, keys (necessary to unlock the bolts) and a treasure. The agent starts its activity from the ladder on top of the maze (home location) and its overall task is to find the treasure and bring it back to the starting location. In our experimentation, the agent starts endowed with no previous knowledge about the possible actions that can be executed in the environment; the agent is only aware of the basic motion primitives at its disposal $A = \{go_up, go_down, go_left, go_right, interact\}$, respectively used to move the agent up, down, left or right by 2-4 pixels (the exact value is randomly selected with a uniform distribution) and to interact with the closest object. The interaction with a lever changes the state (open/close) of the doors associated with that lever (both on the same floor or on different floors) while the interaction with the key and/or the treasure simply collects the key and/or the treasure inside the agent's bag (in the bottom-right corner of the screen). The interaction with the bolt opens the door next to the treasure and it is feasible only when the agent has the key in its bag. The state $s \in S$ is defined in terms of the following low-level variables:

$$s = (x_{agent}, y_{agent}, \theta_1, \theta_2, \dots, x_{key}, y_{key}, x_{bolt}, x_{treasure}, y_{treasure}) \quad (22)$$

in which x_{agent}, y_{agent} is the (x,y) position of the agent, θ_i is the angle of the lever i , x_{key}, y_{key} is the (x,y) location of the key, x_{bolt} is the state of the bolt (1 if open and 0 if locked) and $x_{treasure}, y_{treasure}$ is the (x,y) location of the treasure.

The system is tested in five different configurations of the environment, of increasing complexity: two small-sized (Figure 4a and 4b), two medium-sized (Figure 4c and 4d) and one complete instance (Figure 4e). For example, in the setting depicted in Figure 4e the agent has to: (i) pull the two levers on the top of the maze to open the doors, (ii) get the key which is used to (iii) unlock the bolt, (iv) get the treasure and (v) bring it back on top of the environment. The obstacles that pertain to the other configurations are shown in Table 1.

4.2 The cycle

For exemplificatory purposes, a complete execution cycle of the system is briefly described in the following, showing the output of each phase of an intermediate cycle (cycle n. 10) performed on *domain3* (see Figure 4c) in the *Goal Babbling* strategy case.

Option Generation

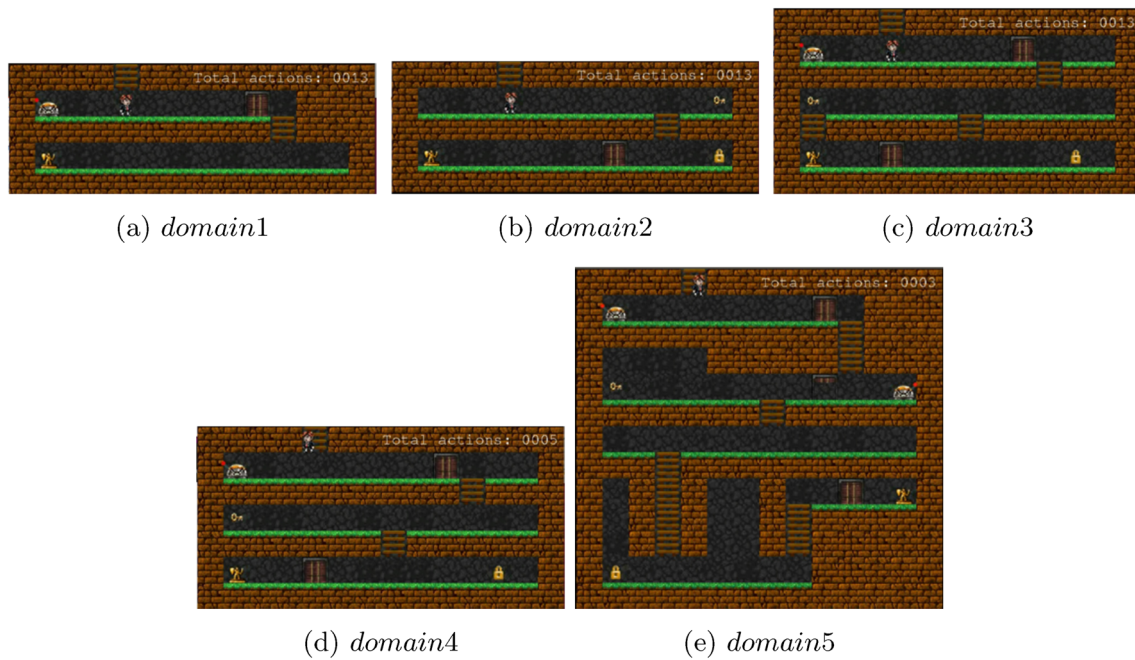


Fig. 4 All the domain configurations used in the experimental analysis. The game's purpose is to get the treasure at the bottom of the maze and bring it back to the top ladder

Table 1 Obstacles that must be overcome to complete the game. Each domain is characterized by a different configuration of levers and bolts that the agent must manipulate to unlock the doors and progress through the environment.

Domain	Levers	Keys	Bolts	Notes
<i>domain1</i>	1	0	0	None.
<i>domain2</i>	0	1	1	None.
<i>domain3</i>	1	1	1	None.
<i>domain4</i>	1	1	1	Shortcut available.
<i>domain5</i>	3	1	1	2 levers going to the treasure, 1 going home.

At the beginning of each cycle, the agent executes Algorithm 2 to collect some options exploiting the agent's primitives, before the exploration can commence. In the Treasure Game environment selected for this work, the agent executes $d_{eps} = 1$ episodes, composed by $d_{steps} = 200$ primitive actions. After the execution of the plan ω^{EX} , the random exploration is reprised using the primitives contained in A (see Section 2.1). The result is the following set of learned options (11 in total) following the formalization (14):

```

O = { (go_up, {}), (go_down, {}),
      (go_left, {}), (go_left, go_up),
      (go_left, go_down),
      (go_left, interact), (go_right, {}),
      (go_right, go_up), (go_right, go_down),
      (go_right, interact),
      (interact, {})}.
```

It is important to note that, in general, the discovered options are not all the options that may be possibly discovered in the environment, but only those experienced by the agent dur-

ing the exploration. This procedure is incremental, adding options to the set O each iteration of the Algorithm 1. As described in section 3.1.1, this procedure leverages IMs *at low level*, capturing the curiosity of the agent when it discovers to have new available primitives to exploit.

Exploration

In this step, the plan ω^{EX} is again executed before starting the random walk over the available options O . In this particular example, 600 data entries have been collected in the *ID* dataset and 6600 in the *TD* dataset. It is important to remember that the number of collected data over the cycles is not constant because when an option does not produce effects on the environment, no data is collected.

Abstraction

Figure 5 shows a part of the output of the PPDDL domain obtained from the abstraction procedure. As visible, the figure does present a valid PPDDL domain description, upon which the planner may reason. Moreover, any subset of the produced domain predicates (symbols) can be used to define high-level goals the planner may try to plan for. In this


```

(define (domain TreasureGame)
  (:requirements :strips :probabilistic-effects :rewards)

  (:predicates
    (notfailed)
    (symbol_0)
    (symbol_1)
    (symbol_2)
    ...
    (symbol_25)
  )

  (:action option-0-partition-0-0
    :parameters ()
    :precondition (and (notfailed) (symbol_5) (symbol_22)
                       (symbol_14))
    :effect (and (symbol_6) (not (symbol_5)) (decrease
                                                (reward) 36.00))
  )

  ...

  (:action option-10-partition-3-715
    :parameters ()
    :precondition (and (notfailed) (symbol_20) (symbol_22)
                       (symbol_14) (symbol_15) (symbol_0))
    :effect (and (symbol_13) (not (symbol_20)) (decrease
                                                (reward) 90.00))
  )

```

Fig. 5 An extract of the PPDDL generated by the abstraction procedure. Notice that `option-0-partition-0-0` can be explained by looking at the symbols of Figure 6. In fact, the effect is to change the x position replacing `symbol_5` with `symbol_6`, maintaining invariant the presence of `symbol_14` and `symbol_22`. Consequently, it is the option moving the agent on the right on the last floor

specific case, the system generated 26 symbols and 716 operators.

Planning

The produced PPDDL domain is used by the system to (1) solve the entire game, or task, and (2) improve the exploration. Figure 7 illustrates two PDDL problems: the problem on top refers to the general goal of solving the game ($\mathcal{P}_g$) while the problem at the bottom refers to the goal dynamically synthesized by the *Goal Selector* module, which in this

example relates to reaching the left corner of the middle floor ($\mathcal{P}_{target}$). The solution for both problems is presented in Figure 8. The first solution, ω^g , solves the game problem in 21 moves and the second one, ω^{EX} , reaches the *Goal Selector* goal in 6 moves.

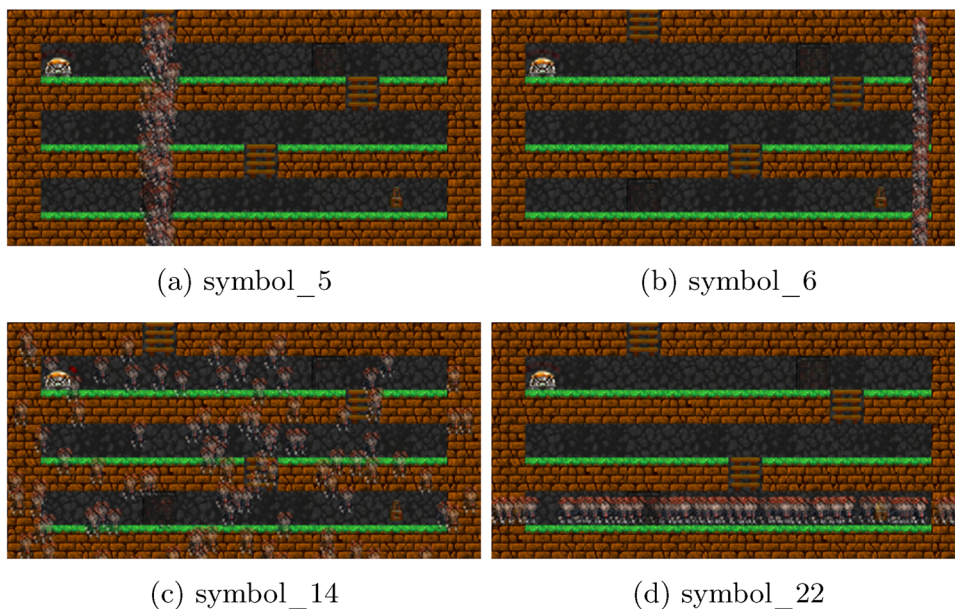
4.3 Results

In this section, the overall results of the system over different domain instances are described. First, the setting of the environment is discussed, then the adopted baseline, as well as the other strategies enabling the planning exploration. Subsequently, some charts are presented, that highlight the system's performance in terms of success ratio on the planning task. Finally, some issues worth being underscored are commented, and the limitations of the employed technologies are discussed.

For our purposes, the Treasure Game⁵ has been used with five different mazes (see Figure 4), to focus on the performances of the system on smaller domains, highlighting pros and cons of the symbolic approach proposed. The system has been executed, following the workflow described in Algorithm 1, in the cited five mazes configurations using parameters suitable to solve the tasks, which are described later. To summarize, the system iteratively (i) looks for new options O_{new} performing d_steps primitives for d_eps episodes, (ii) explores for dpa_eps episodes the maze executing dpa_steps , (iii) creates a symbolic abstraction of the domain $\mathcal{D}$, (iv) creates a plan ω^{EX} to optimize the exploration of the successive cycle $c + 1$. It is important to note that the autonomously discovered options successfully pro-

⁵ Github repository: <https://github.com/sd-james/gym-treasure-game>

Fig. 6 Semantics of the symbols employed by the first operator of the generated PPDDL in Figure 5: (a) and (b) mean being in a certain x position, (c) mean the lever is pulled and (d) being on the bottom floor




```

(define (problem task_goal)
  (:domain TreasureGame)

  (:init (notfailed) (symbol_0) (symbol_1)
         (symbol_2) (symbol_3)
         (symbol_4) (symbol_5) )

  (:goal (and (notfailed) (symbol_4)
              (symbol_17) (symbol_18)) )
)

(define (problem im_goal)
  (:domain TreasureGame)

  (:init (notfailed) (symbol_0)
         (symbol_1) (symbol_2)
         (symbol_3) (symbol_4)
         (symbol_5) )

  (:goal (and (symbol_13) (symbol_24)
              (symbol_14) (symbol_2)
              (notfailed)) )
)

```

Fig. 7 On the top, the problem *task_goal* of solving the task is synthesized by the system $\mathcal{P}_g$ and, on the bottom, the problem *im_goal* generated by the *Goal Selector* module $\mathcal{P}_{target}$

PLAN TASK GOAL:

```

[ 1:(go_down,{}),      ; climb down the stairs
  2:(go_left,go_down), ; go left until it can go down
  3:(interact,{}),     ; pull the lever
  4:(go_right,go_down), ; go right up to the stairs to go down
  5:(go_down,{}),      ; climb down the stairs
  6:(go_right,{}),     ; go right up to the end of the corridor
  7:(go_left,{}),     ; go left up to the end of the corridor
  8:(interact,{}),     ; take the key
  9:(go_right,go_down), ; go right up to the stairs
  10:(go_down,{}),     ; climb down the stairs
  11:(go_right,interact), ; go right up to the bolt
  12:(interact,{}),    ; unlock the bolt
  13:(go_left,go_down), ; go left until it is possible
  14:(interact,{}),    ; get the treasure
  15:(go_right,go_up), ; go right up to the stairs
  16:(go_up,{}),       ; go upstairs
  17:(go_right,{}),    ; go right up to the end of the corridor
  18:(go_left,go_up), ; go left up to the stairs
  19:(go_up,{}),       ; go upstairs
  20:(go_left,go_up), ; go left up to the stairs
  21:(go_up,{}),      ; go up to home ]

```

PLAN IM GOAL:

```

[ 1:(go_down,{}),      ; climb down the stairs
  2:(go_left,go_down), ; go left up to the lever
  3:(interact,{}),     ; pull the lever
  4:(go_right,go_down), ; go right until it can go down
  5:(go_down,{}),      ; climb down the stairs
  6:(go_left,interact) ] ; go left up to the key

```

Fig. 8 The plans generated ω^g and ω^{EX}

duced a valid high-level model, usable to build correct plans. To assess the accuracy of the current abstraction $\mathcal{D}$, we evaluate its ability to reach the goal s_g at the end of each cycle. This is done by requesting the planner to solve the problem $\mathcal{P}_g$ using the currently produced symbolic domain description $\mathcal{D}$.

The first strategy used in the experiments is the random walk, which is called *Action Babbling* [50] (see Section 3.1.4). This strategy simply executes the Algorithm 1

Table 2 Main parameters used in the different environment configurations. For each domain, the number of episodes (*dpa_eps*) and steps (*dpa_steps*) per cycle employed in the exploration phase is declared. The last column shows the number of steps required by the optimal plan to solve the entire game.

Domain	<i>dpa_eps</i>	<i>dpa_steps</i>	minimal solution steps
<i>domain1</i>	4	50	11
<i>domain2</i>	4	50	13
<i>domain3</i>	4	150	19
<i>domain4</i>	4	200	15
<i>domain5</i>	4	800	31

without using planning because the *Goal Selector* returns $s_{target} = NULL$ and no plan ω^{EX} is generated. This simple strategy is used as a baseline for the experimental analysis, and it is necessary to fully appreciate the advantages of using the symbolic approach. The exploration is equivalent to the one used by Konidaris [59], and we use it to observe the development of the symbolic model over time.

To support the use of symbolic planning, two other strategies have been considered: *Goal Babbling* and *Distance-based Goal Babbling* (see Section 3.1.4). These three strategies could be seen as having an increasing complexity and effectiveness in the exploration:

- *Action Babbling* selects completely random actions;
- *Goal Babbling* selects random goals, reaching them with a planned sequence of actions;
- *Distance-based Goal Babbling* selects the farthest goals and reaches them by using planning (see subsection 3.1.4).

It is reasonable to conjecture that in smaller domains where the reasoning is less necessary and purely random exploration may suffice, we do not expect to observe much difference among the previous strategies; while in the bigger domains the time to reach the goal may significantly change, depending on the strategy employed.

The five configurations considered are depicted in Figure 4, and the parameter values used in the exploration function *Collect_Data* for each configuration are shown in Table 2. Smaller domains *domain1* and *domain2* required the execution of 50 options per episode, *domain3* required 150, *domain4* required 200, and finally 800 options were executed in *domain5*.

4.3.1 Planning success rate

The main results of the system are depicted in Figure 9. Precisely, in the graphs, it is shown the probability of success in solving the game over time using different strategies. Such

success entails the generation of a sufficiently mature domain $\mathcal{D}$ and a correct problem formalization of $\mathcal{P}_g$, resulting in a correct plan ω^g .

The system has been run with $\text{cycles} = 15$, assessing at the end of each cycle whether ω^g was able to solve the game using the current synthesized PPDDL domain representation $\mathcal{D}$. This mechanism has been performed for ten trials. Consequently, in the charts of Figure 9 are depicted cycles on the x -axis, and the percentage of trials that have succeeded in solving the game in the specific cycle on the y -axis.

In the smaller domains *domain1*, *domain2*, *domain3* and *domain4*, the planning contribution is limited and sometimes not visible at all. Especially in the simplest domain *domain1* the three strategies present similar performances, solving in the last cycles 80-90% of the trials (Figure 9a). In domains *domain2* and *domain3*, presenting slightly higher complexities, the advantages of executing plans start being visible (Figure 9b and 9c). It is evident that *Goal Babbling* behaves almost perfectly in the last cycles, demonstrating the usefulness of collecting transition data with reasonable sequences of actions. The combination of randomness in selecting the goal to be reached and the reasoned transitions speed up the exploration and the consequent maturity of the PPDDL representation of the environment. Instead, the *Distance-based Goal Babbling* seems better than a completely random search but less effective than the precedent. This result entails that applying an exploration focused on the frontier's goals is not convenient in smaller domains. In fact, in the conclusive cycles of *domain2* and *domain3*, *Goal Babbling* maintains 90-100% of success ratio, *Distance-based Goal Babbling* around 80% and *Action Babbling* between 60-70%.

Although the problem faced by *domain3* seems easier than *domain4*, presenting respectively minimal solutions of 19 and 15 steps (see Table 2), the system needs more steps per episode to solve *domain4* using the same amount of episodes used by *domain3*. The reason of this behaviour is due to the higher complexity of synthesizing a PPDDL domain $\mathcal{D}$ for this scenario. In fact, the scenario *domain4* presents a sort of "shortcut" to reach the treasure, which does not require collecting the key and opening the door. From the point of view of the abstraction procedure, the shortcut represents a branch in the possible choices of the agent, thus presenting the agent with additional concepts to be abstracted in order to synthesize a complete representation. Consequently, the system generates additional symbols and operators, requiring more experience to strengthen the transition model captured by the PPDDL and provide satisfying plans.

The significance of the frontier exploration emerges in the *domain5* configuration, where planning evidently boosts the results of the agent. Indeed, being bigger than other domains, *domain5* is more difficult to be solved and planning results to be efficient in driving the exploration of the agent immedi-

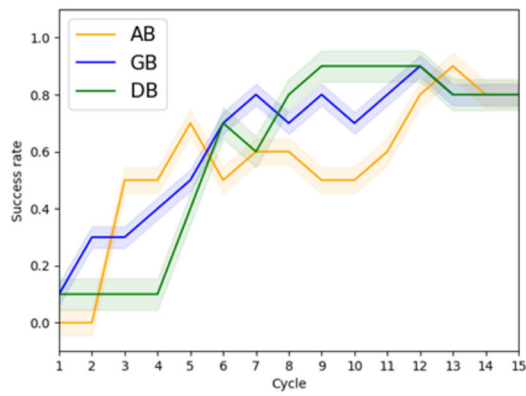
ately towards the borders of its knowledge. The main result highlighted by the charts is that in bigger domains (Figure 9e), the impact of using planning is evident. Planning is able to easily drive the agent towards interesting visited states, where it can continue exploring. Statistically, after collecting the transition data for 15 cycles, the system struggles to solve the problem, never exceeding 10% of success. Similar behaviour for *Goal Babbling*, reaching most 20% of success. Instead, *Distance-based Goal Babbling* results being extremely effective, reaching 90% of success. Conceptually, the strategy continuously pushes the agent towards the frontiers of the visited states, increasing the probability of encountering new unexplored areas and reducing redundant data.

4.3.2 Spatial and Temporal exploration

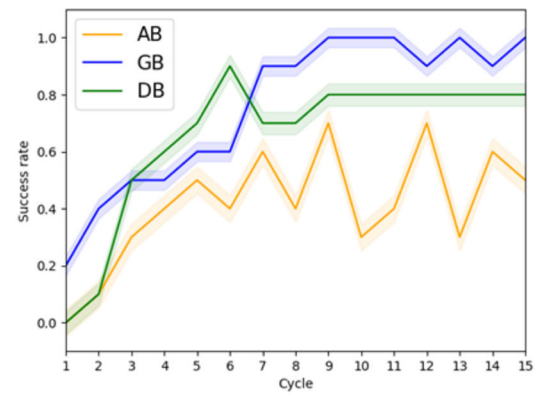
In addition to the planning success rate analysis presented in Section 4.3.1, we further evaluate the system's ability to explore and reach intermediate subgoals (i.e., unlocking the bolt and retrieving the treasure) from both spatial and temporal perspectives. This analysis focuses on *domain5*, selected for its higher dimensionality and extended temporal dynamics.

First, we examine the temporal evolution of two low-level variables: x_{bolt} (Figure 10a) and $y_{treasure}$ (Figure 10b), respectively representing the x -coordinate of the bolt object and the y -coordinate of the treasure in the environment (see equation 22). These plots show the mean value on all the episodes of both variables, computed over time. They illustrate how quickly each strategy can alter the environment's state. In our case, a transition of x_{bolt} from 1 to 0 indicates successful unlocking of the bolt; thus, lower mean values over time reflect faster achievement of this subgoal. As Figure 10a shows, the Distance-based Goal Babbling (DB) strategy consistently reaches this subgoal faster than both Goal Babbling (GB) and Action Babbling (AB) across all timesteps. A similar trend is observed in Figure 10b, which tracks changes in $y_{treasure}$ from 0.62 to 0.92, signaling the agent's success in grabbing the treasure. Again, the DB strategy outperforms the others, more rapidly shifting the variable toward its target value. It is important to note that the curves in Figure 10 do not reach their minimum (for x_{bolt}) or maximum (for $y_{treasure}$) values even at the final timestep. This occurs because, during the initial learning cycles, the agent may fail to achieve these subgoals, since even planning-based strategies require time to synthesize a usable high-level representation.

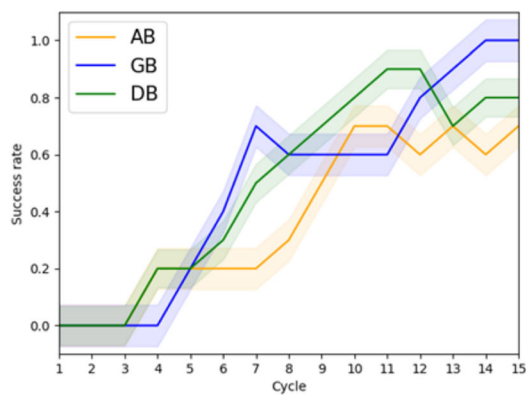
Figures 11 and 12 provide a spatial analysis of the agent's exploration patterns, focusing on the x and y coordinates of the agent's location before and after the treasure is obtained, respectively. Each figure depicts heatmaps showing all visited states across all episodes, organized by strategy (columns) and different timestep limits (100, 200, 400, and



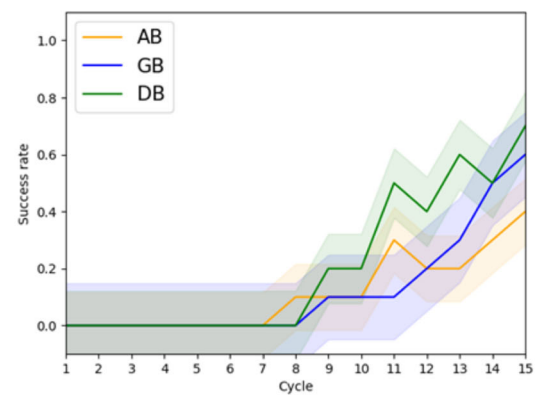
(a) domain 1



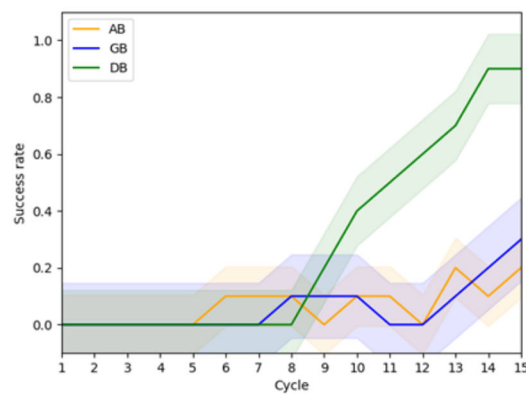
(b) domain 2



(c) domain 3



(d) domain 4

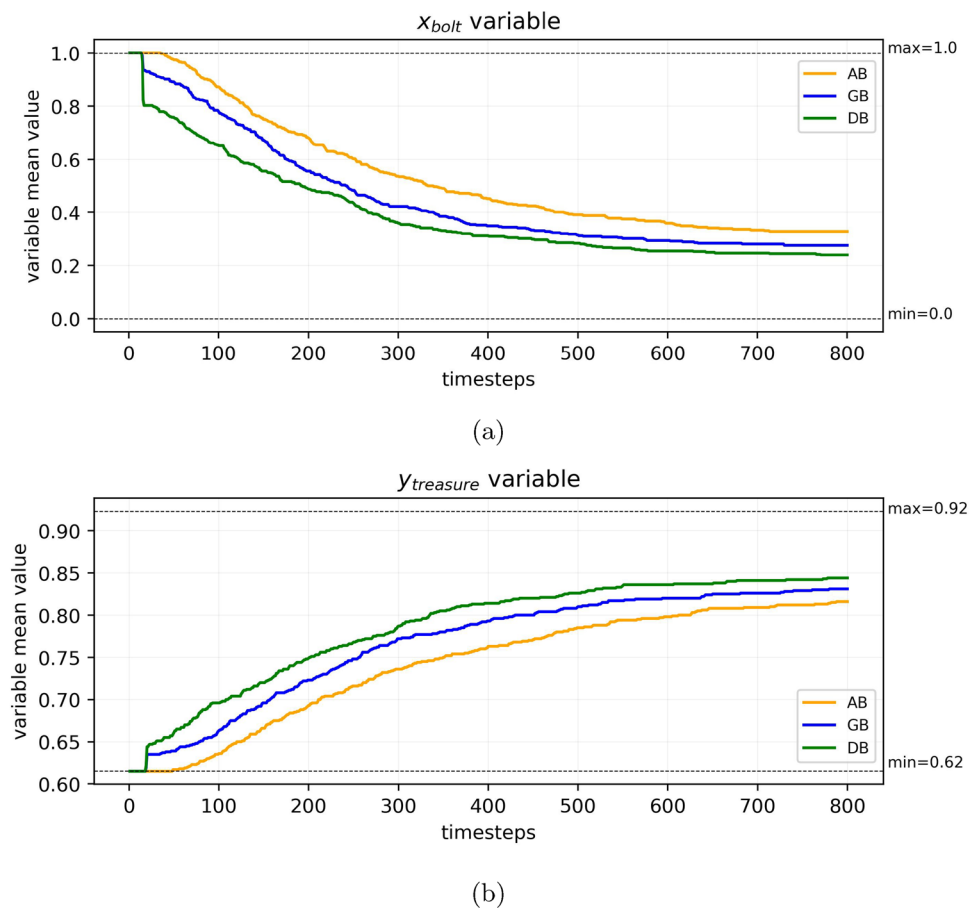


(e) domain 5

Fig. 9 The figure depicts the success rate of ω^g in solving the game using the exploration strategies Action Babbling (AB), Goal Babbling (GB) and Distance-based Goal Babbling (DB) over time. Figures 9a,

9b, 9c, 9d, and 9e, respectively show the results of the domains depicted in figures 4a, 4b, 4c, 4d, and 4e, over 15 cycles of exploration

Fig. 10 The mean value over all episodes of the variables x_{bolt} and $y_{treasure}$ in *domain5*. In subfigure (a) it is depicted the achievement of unlocking the bolt, detected by the change of state from value 1 to 0. In subfigure (b) it is shown the mean value of the y-axis position of treasure varying from value 0.62 to 0.92, determining the achievement of the subgoal “get_treasure”, respectively. It is clear that Distance-based Goal Babbling (DB) reaches faster both the subgoals compared to Goal Babbling (GB) and Action Babbling (AB) strategies



800 steps; rows). These heatmaps reveal how each strategy allocates the exploration effort before and after accomplishing the subgoal (i.e., obtaining the treasure). Intuitively, more efficient strategies visit fewer states before grabbing the treasure and more states afterward. This is evident in Figure 11: within the first 100 timesteps (first row), AB displays a high density of visits to the upper portion of the environment, particularly around the top ladder, indicating prolonged exploration in this area. In contrast, both GB and DB show a more directed movement toward the treasure, with fewer samples in the upper region early on. By the 800th timestep (bottom row), the heatmaps shift downward across all strategies, reflecting progress toward deeper levels of the environment. However, the density of visited states remains noticeably higher for AB, while GB and especially DB demonstrate more focused and efficient exploration patterns.

Finally, Figure 12 presents heatmaps of the states visited after the agent successfully grabs the treasure, varying across timesteps (rows) and exploration strategies (columns). In the first row, corresponding to the initial 100 timesteps, it is evident that the Action Babbling (AB) strategy rarely reaches the treasure, resulting in an almost empty heatmap. In contrast, both Goal Babbling (GB) and, more prominently, Distance-

based Goal Babbling (DB) display a significantly higher density of post-treasure samples. Notably, the DB strategy occasionally manages to reach the top ladder (represented by the fourth pixel from the left on top of the heatmap), as indicated by a slightly darker red area not present in the other strategies. This suggests that, even within a limited number of steps after acquiring the treasure, DB is sometimes able to reach the final goal. In comparison, GB and AB exhibit the first signs of reaching the top ladder only in the heatmaps corresponding to 200 and 400 timesteps, respectively. These observations are further supported by the heatmaps at later timesteps, which show a consistent increase in sample density across all strategies. However, DB maintains a clear advantage, with a higher number of visited states concentrated around relevant goal areas. This emphasizes the effectiveness of planning-based exploration in accelerating progress toward complex goals.

4.4 Discussion and Future Work

As observable from the data plots, the results are not deterministic and change over time. On the one hand, the system is continuously evolving in terms of knowledge and, on the other hand, ML techniques introduce further stochasticity

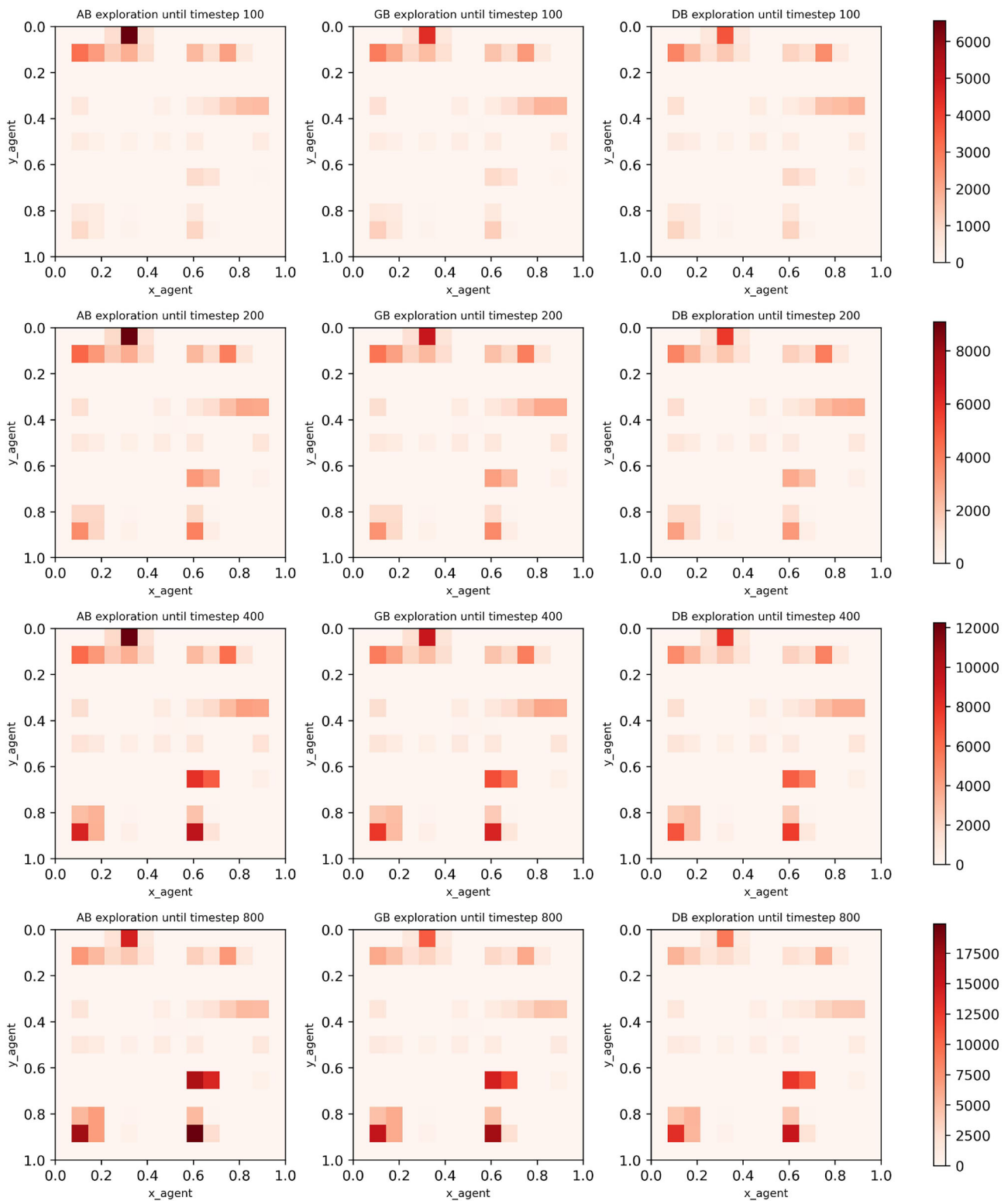


Fig. 11 In the figure, it is depicted the heatmap of the coordinates (x,y) of all the states visited by the agent before getting treasure until a certain timestep. Each column presents the heatmaps for a certain strategy (from the left Action Babbling, Goal Babbling and Distance-based Goal Babbling), and each row defines the number of considered timesteps (from

top 100, 200, 400 and 800). In these heatmaps it is clear that the more the strategy is guided by planning the less presents samples, because it reaches faster the subgoal of getting the treasure, from which the states are not included in the heatmap

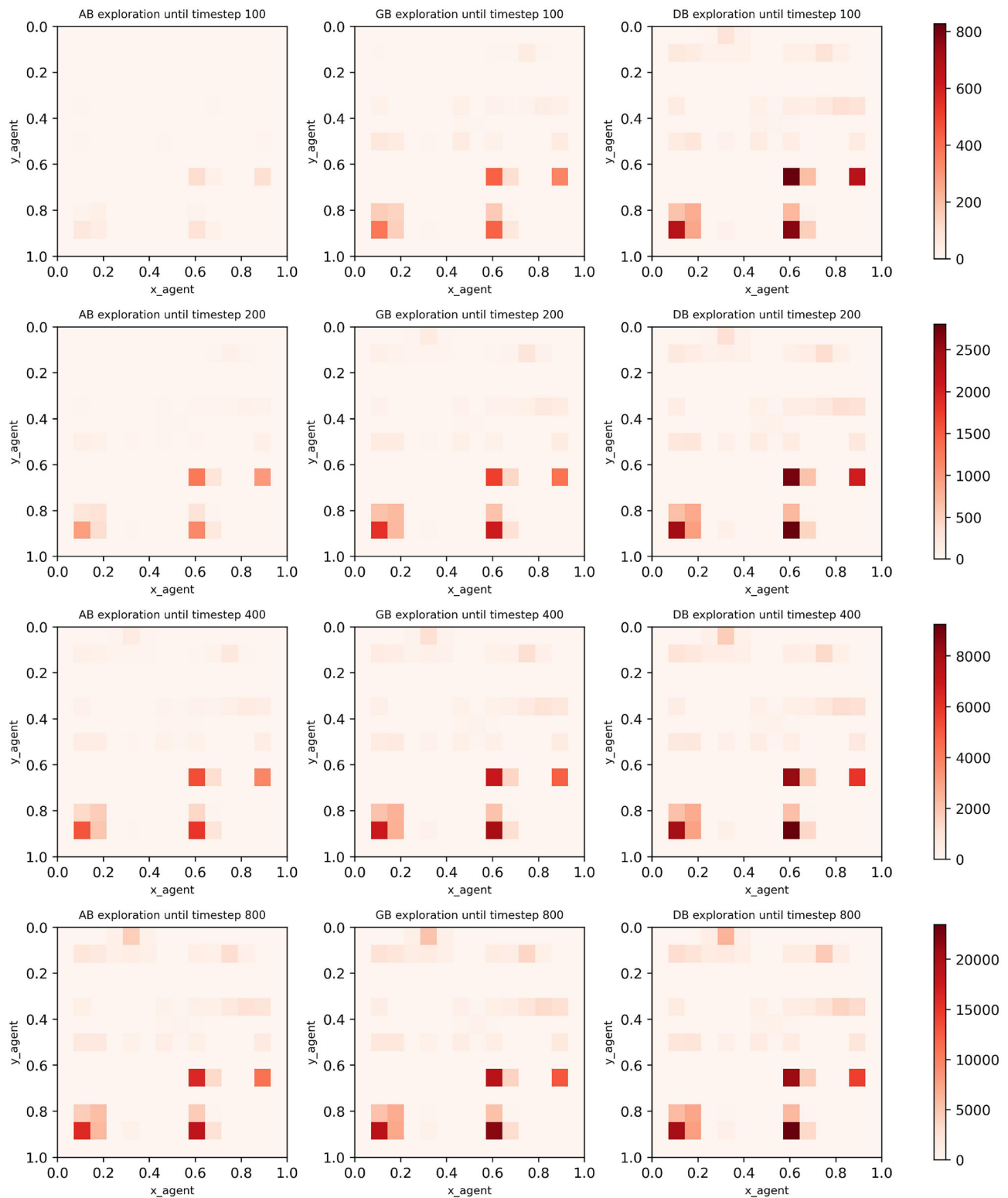


Fig. 12 In the figure, it is depicted the heatmap of the coordinates (x, y) of all the states visited by the agent after getting treasure until a certain timestep. Each column presents the heatmaps for a certain strategy (from the left Action Babbling, Goal Babbling and Distance-based Goal Babbling), and each row defines the number of considered timesteps (from

top 100, 200, 400 and 800). In these heatmaps it is clear that the more the strategy is guided by planning the more presents samples, because it reaches faster the subgoal of getting the treasure, from which the states are included in the heatmap.



Fig. 13 The graphical representation of the symbol “on the higher y-axis position”. It can be seen that it is significantly noisy because some samples of the agent are depicted on the top y-axis and other samples almost on the lower floor

to the final representation generated. During the abstraction procedure, described in Section 3.1.3, some statistical tools are employed to create data structures to represent preconditions, effects and symbols.

For instance, an erroneous⁶ clustering phase could generate unexpected effects on symbols and operators. An example highlighting this fact is the *noisiness of some symbols*, interfering with the generation of a correct formalization of $\mathcal{D}$, $\mathcal{P}$ and, consequently, the resulting plan ω . In Figure 13 it can be seen the graphical representation of the symbol “on the highest y-axis position”, meaning “on the top ladder” because it is the only possible state with such y-axis value (it is not allowed to move inside the walls).

In some cases, it could happen that the initial state of the game is interpreted as being at the top floor under the ladder and, consequently, it is not necessary to climb down the ladder to execute the agent’s task. Then, although almost the whole plan ω is correct to complete the game, without the option of climbing down the ladder as the first action, the plan is incomplete, and the trial is considered unsuccessful. On the one hand, this is a drawback of using ML tools, which can introduce noise in the symbolic representation. On the other hand, it is the strong point of the probabilistic approach, always proposing solutions, even though presenting a certain degree of error. In any case, the hyperparameter values of the employed machine learning components could

be estimated using regression models to further reduce this noise and improve reliability [76, 77].

In addition, the system proposed is limited by the *propositional logic* representation, which does not permit the inclusion of arguments in operators and symbols. Such limitation could be overtaken by developing a new abstraction procedure following other logics (e.g. First Order Logic (FOL)). Recently, an implementation of an object-centric abstraction procedure has been proposed [78], demonstrating that it is possible to create a lifted representation suitable for transfer the knowledge to new tasks.

Another aspect to be analyzed deeper is the *management of the knowledge data*. Unfortunately, the abstraction module has no real incremental nature, but the abstraction procedure is executed over all the data, each cycle from scratch. This way of operating is not computationally efficient and does not reflect the learning process of the human being. However, given that in our domains all the necessary knowledge is discovered in a certain amount of time and does not change, a first improvement to be applied to the system could be maintaining a maximum of transition and initiation data discarding over-sampled transitions. Therefore, a sort of filter in the knowledge acquisition could stop the growth in terms of the abstraction procedure’s time in stationary and limited domains but also reduce it for all the other cases.

Then, the abstraction time seems to be linear with respect to the number transition tuples (in the worst-case exponential according to Konidaris [59]). Consequently, it would be fundamental to find solutions that mitigate this aspect. Possible actions could be to filter the acquired data, or structure the collected knowledge in a more efficient form for the abstraction process and, eventually, produce different complementary representations which are used according to the task to be tackled.

Finally, a natural evolution of this work lies in the neuro-symbolic approach. The integration of explicit knowledge inside sub-symbolic systems makes the learning process more effective and efficient, taking advantage of both the approaches. One of most promising extension of this work should embrace this new techniques, as some other preliminary works have already done [54, 55, 57].

5 Conclusions

In this paper, we present a novel approach to open-ended learning in autonomous agents based on intrinsically motivated planning. This approach combines two powerful paradigms, intrinsic motivation and classical planning, to enable agents to continuously acquire and refine knowledge and skills without relying on external supervision or predefined rewards.

⁶ Statistical methods are not inherently flawed, as they simply generate representations based on the data provided. However, for our purposes, certain model instances may hinder the achievement of our objective.

The proposed method offers an alternative, or a valuable complement, to widely adopted sub-symbolic approaches, and exhibits several compelling advantages. First, it allows agents to explore and learn in a self-directed, open-ended fashion, free from the constraints of a fixed set of tasks or goals. Second, it enables agents to represent and reason about their knowledge and skills in a structured, formal way, facilitating planning and generalization to novel situations. Third, by integrating intrinsic motivations, it encourages exploration and learning beyond extrinsic objectives, enhancing the agent's adaptability, robustness, and creativity. Our experimental analysis demonstrates the feasibility of this approach: starting with no prior knowledge, an agent can autonomously explore its environment and acquire useful abstractions. Performance improves significantly when planning mechanisms are used to select intermediate subgoals, effectively guiding and accelerating the exploration process.

Nevertheless, several challenges and opportunities remain for future research aimed at enabling such systems to perform complex tasks in realistic settings. Notably, the expressiveness of the high-level representations is currently limited by the symbolic abstraction method adopted, which relies on propositional logic, and by the lack of human-interpretable symbols and operator names in the generated PPDDL models. Progress in this area could be achieved by integrating a layer powered by a Large Language Model (LLM), capable of translating the automatically synthesized symbolic representations into more comprehensible, human-aligned forms. This direction is particularly promising in robotics, where agents could acquire new skills (e.g., learning new options via Imitation Learning) and autonomously construct their own high-level representations. Ultimately, this would allow agents not only to plan more effectively but also to interpret and respond to goals expressed in natural language, bridging the gap between autonomous learning and human-robot interaction. Overall, we believe that our approach represents a promising step toward the development of more autonomous, intelligent agents capable of open-ended, self-directed learning and continual improvement.

Funding Open access funding provided by Università degli Studi di Torino within the CRUI-CARE Agreement. This work was partially funded by the European Union's Horizon 2020, research and innovation programme under GA 101070381 ('PILLAR-Robots - Purposeful Intrinsically-motivated Lifelong Learning Autonomous Robots') and PNRR MUR project PE0000013-FAIR.

Data Availability The implementation of the system is available at https://github.com/gabrielesartor/discover_plan_act.

Declarations

Conflicts of Interest The authors declare that they have no conflicts of interest.

Open Access This article is licensed under a Creative Commons Attribution 4.0 International License, which permits use, sharing, adaptation, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons licence, and indicate if changes were made. The images or other third party material in this article are included in the article's Creative Commons licence, unless indicated otherwise in a credit line to the material. If material is not included in the article's Creative Commons licence and your intended use is not permitted by statutory regulation or exceeds the permitted use, you will need to obtain permission directly from the copyright holder. To view a copy of this licence, visit <http://creativecommons.org/licenses/by/4.0/>.

References

1. Campbell, M., Hoane, A.J., Hsu, F.H.: Deep Blue. *Artif. Intell.* **134**(1), 57–83 (2002). [https://doi.org/10.1016/S0004-3702\(01\)00129-1](https://doi.org/10.1016/S0004-3702(01)00129-1)
2. Silver, D., Hubert, T., Schrittwieser, J., Antonoglou, I., Lai, M., Guez, A., et al.: A general reinforcement learning algorithm that masters chess, shogi, and Go through self-play. *Science* **362**(6419), 1140–1144 (2018)
3. Silver, D., Huang, A., Maddison, C.J., Guez, A., Sifre, L., Van Den Driessche, G., et al.: Mastering the game of Go with deep neural networks and tree search. *Nature* **529**(7587), 484–489 (2016)
4. Silver, D., Schrittwieser, J., Simonyan, K., Antonoglou, I., Huang, A., Guez, A., et al.: Mastering the game of go without human knowledge. *Nature* **550**(7676), 354–359 (2017)
5. Oh, J., Guo, X., Lee, H., Lewis, R.L., Singh, S.: Action-conditional video prediction using deep networks in atari games. *Adv. Neural. Inf. Process. Syst.* **28**, (2015)
6. Berner, C., Brockman, G., Chan, B., Cheung, V., Debiak, P., Dennison, C., et al. Dota 2 with large scale deep reinforcement learning. *arXiv preprint arXiv:1912.06680*. 2019
7. Koch, W., Mancuso, R., West, R., Bestavros, A.: Reinforcement learning for UAV attitude control. *ACM Transactions on Cyber-Physical Systems*. **3**(2), 1–21 (2019)
8. Nguyen, H., La, H.: Review of deep reinforcement learning for robot manipulation. In: 2019 Third IEEE International Conference on Robotic Computing (IRC). IEEE; p. 590–595 (2019)
9. Ibarz, J., Tan, J., Finn, C., Kalakrishnan, M., Pastor, P., Levine, S.: How to train your robot with deep reinforcement learning: lessons we have learned. *The International Journal of Robotics Research*. **40**(4–5), 698–721 (2021)
10. Bellemare, M.G., Candido, S., Castro, P.S., Gong, J., Machado, M.C., Moitra, S., et al.: Autonomous navigation of stratospheric balloons using reinforcement learning. *Nature* **588**(7836), 77–82 (2020)
11. Sutton, R.S., Barto, A.G.: Reinforcement learning: An introduction. MIT press (2018)
12. Bengio, Y., Goodfellow, I., Courville, A.: Deep learning, vol. 1. MIT press Cambridge, MA, USA (2017)
13. Butler, R.A.: Discrimination learning by rhesus monkeys to visual-exploration motivation. *J. Comp. Physiol. Psychol.* **46**(2), 95 (1953)
14. Harlow, H.F.: Learning and satiation of response in intrinsically motivated complex puzzle performance by monkeys. *J. Comp. Physiol. Psychol.* **43**(4), 289 (1950)
15. Montgomery, K.C.: The role of the exploratory drive in learning. *J. Comp. Physiol. Psychol.* **47**(1), 60 (1954)
16. Berlyne, D.E.: Novelty and curiosity as determinants of exploratory behaviour. *Br. J. Psychol.* **41**(1), 68 (1950)

17. Berlyne, D.E.: Curiosity and Exploration: Animals spend much of their time seeking stimuli whose significance raises problems for psychology. *Science* **153**(3731), 25–33 (1966)
18. Ryan, R.M., Deci, E.L.: Intrinsic and Extrinsic Motivations: Classic Definitions and New Directions. *Contemp. Educ. Psychol.* **25**(1), 54–6 (2000). <https://doi.org/10.1006/ceps.1999.1020>
19. Düzel, E., Bunzeck, N., Guitart-Masip, M., Düzel, S.: NOvelty-related motivation of anticipation and exploration by dopamine (NOMAD): implications for healthy aging. *Neuroscience & Biobehavioral Reviews*. **34**(5), 660–669 (2010)
20. Ranganath, C., Rainer, G.: Neural mechanisms for detecting and remembering novel events. *Nat. Rev. Neurosci.* **4**(3), 193–202 (2003)
21. Redgrave, P., Gurney, K.: The short-latency dopamine signal: a role in discovering novel actions? *Nat. Rev. Neurosci.* **7**(12), 967–975 (2006)
22. Santucci, V.G., Oudeyer, P.Y., Barto, A., Baldassarre, G.: Intrinsically motivated open-ended learning in autonomous robots. *Front. Neurobot.* **13**, 115 (2020)
23. Oudeyer, P.Y., Kaplan, F., Hafner, V.: Intrinsic motivation systems for autonomous mental development. *IEEE Trans. Evol. Comput.* **11**(2), 265–286 (2007)
24. Baldassarre, G., Mirolli, M.: *Intrinsically Motivated Learning in Natural and Artificial Systems*. Springer Science & Business Media (2013)
25. Frank, M., Leitner, J., Stollenga, M., Förster, A., Schmidhuber, J.: Curiosity driven reinforcement learning for motion planning on humanoids. *Front. Neurobot.* **7**, 25 (2014)
26. Bellemare, M., Srinivasan, S., Ostrovski, G., Schaul, T., Saxton, D., Munos, R.: Unifying count-based exploration and intrinsic motivation. *Advances in neural information processing systems*. **29** (2016)
27. Santucci, V.G., Baldassarre, G., Mirolli, M.: GRAIL: A Goal-Discovering Robotic Architecture for Intrinsically-Motivated Learning. *IEEE Transactions on Cognitive and Developmental Systems*. **8**(3), 214–23 (2016). <https://doi.org/10.1109/TCDS.2016.2538961>
28. Colas, C., Fournier, P., Chetouani, M., Sigaud, O., Oudeyer, P. Y.: CURIOUS: intrinsically motivated modular multi-goal reinforcement learning. In: *International conference on machine learning*. PMLR; p. 1331–1340 (2019)
29. Blaes, S., Vlastelica Pogančić, M., Zhu, J., Martius, G.: Control what you can: Intrinsically motivated task-planning agent. *Advances in Neural Information Processing Systems*. **32** (2019)
30. Sigaud, O., Baldassarre, G., Colas, C., Doncieux, S., Duro, R., Perrin-Gilbert, N., et al. A Definition of Open-Ended Learning Problems for Goal-Conditioned Agents. *arXiv preprint arXiv:2311.00344*. (2023)
31. Baldassarre, G., Duro, R. J., Cartoni, E., Khamassi, M., Romero, A., Santucci, V. G.: Purpose for Open-Ended Learning Robots: A Computational Taxonomy, Definition, and Operationalisation. *arXiv preprint arXiv:2403.02514*. (2024)
32. Singh, S., Barto, A. G., Chentanez, N.: Intrinsically Motivated Reinforcement Learning. In: *Proceedings of the 17th International Conference on Neural Information Processing Systems*. NIPS'04. Cambridge, MA, USA: MIT Press; p. 1281–1288 (2004)
33. Forestier, S., Portelas, R., Mollard, Y., Oudeyer, P. Y.: Intrinsically motivated goal exploration processes with automatic curriculum learning. *arXiv preprint arXiv:1708.02190*. (2017);
34. Romero, A., Baldassarre, G., Duro, R. J., Santucci, V. G.: Autonomous learning of multiple curricula with non-stationary interdependencies. In: *2022 IEEE International Conference on Development and Learning (ICDL)*. IEEE; p. 272–279 (2022)
35. Bengio, Y., Louradour, J., Collobert, R., Weston, J.: Curriculum learning. In: *Proceedings of the 26th annual international conference on machine learning*; p. 41–48 (2009)
36. Machado, M. C., Bellemare, M. G., Bowling, M.: A laplacian framework for option discovery in reinforcement learning. *arXiv preprint arXiv:1703.00956*. (2017)
37. Oddi, A., Rasconi, R., Santucci, V. G., Sartor, G., Cartoni, E., Mannela, F., et al. Integrating open-ended learning in the sense-plan-act robot control paradigm. In: *ECAI 2020*. IOS Press; p. 2417–2424 (2020)
38. Sartor, G., Zollo, D., Mayer, M. C., Oddi, A., Rasconi, R., Santucci, V. G.: Autonomous Generation of Symbolic Knowledge via Option Discovery. In: *Proceedings of the 9th Italian workshop on Planning and Scheduling (IPS'21) and the 28th International Workshop on "Experimental Evaluation of Algorithms for Solving Problems with Combinatorial Explosion" (RCRA'21)*. vol. 3065 of CEUR Workshop Proceedings. CEUR-WS.org; (2021)
39. Sartor, G., Zollo, D., Cialdea Mayer, M., Oddi, A., Rasconi, R., Santucci, V. G.: Option Discovery for Autonomous Generation of Symbolic Knowledge. In: *International Conference of the Italian Association for Artificial Intelligence*. Springer; p. 153–167 (2022)
40. Barto, A.G., Mahadevan, S.: Recent advances in hierarchical reinforcement learning. *Discrete event dynamic systems*. **13**(1), 41–77 (2003)
41. Konidaris, G., Barto, A. G.: Skill discovery in continuous reinforcement learning domains using skill chaining. In: *Advances in neural information processing systems*; p. 1015–1023 (2009)
42. Niel, R., Wiering, M. A.: Hierarchical Reinforcement Learning for Playing a Dynamic Dungeon Crawler Game. In: *2018 IEEE Symposium Series on Computational Intelligence (SSCI)*. IEEE; p. 1159–1166 (2018)
43. Rafati, J., Noelle, D. C.: Unsupervised Methods For Subgoal Discovery During Intrinsic Motivation in Model-Free Hierarchical Reinforcement Learning. In: *KEG@ AAAI*; p. 17–25 (2019)
44. Parisi, S., Dean, V., Pathak, D., Gupta, A.: Interesting object, curious agent: Learning task-agnostic exploration. *Adv. Neural. Inf. Process. Syst.* **34**, 20516–20530 (2021)
45. Romero, A., Baldassarre, G., Duro, R. J., Santucci, V. G.: Analysing autonomous open-ended learning of skills with different interdependent subgoals in robots. In: *2021 20th International Conference on Advanced Robotics (ICAR)*. IEEE; p. 646–651 (2021)
46. Bagaria, A., Senthil, J. K., Konidaris, G.: Skill Discovery for Exploration and Planning using Deep Skill Graphs. In: Meila M, Zhang T, editors. *Proceedings of the 38th International Conference on Machine Learning*. vol. 139 of *Proceedings of Machine Learning Research*. PMLR; p. 521–531 (2021) Available from: <https://proceedings.mlr.press/v139/bagaria21a.html>
47. Veeriah, V., Zahavy, T., Hessel, M., Xu, Z., Oh, J., Kemaev, I., et al.: Discovery of options via meta-learned subgoals. *Adv. Neural. Inf. Process. Syst.* **34**, 29861–29873 (2021)
48. Nau, D., Ghallab, M., Traverso, P.: *Automated Planning: Theory & Practice*. Morgan Kaufmann Publishers Inc., San Francisco, CA, USA (2004)
49. Russell, S.J.: *Artificial intelligence a modern approach*. Pearson Education, Inc. (2010)
50. Chitnis, R., Silver, T., Tenenbaum, J. B., Kaelbling, L. P., Lozano-Pérez, T.: Glib: Efficient exploration for relational model-based reinforcement learning via goal-literal babbling. In: *Proceedings of the AAAI Conference on Artificial Intelligence*. **35**, 11782–11791 (2021)
51. Mikita, H., Azuma, H., Kakiuchi, Y., Okada, K., Inaba, M.: Interactive symbol generation of task planning for daily assistive robot. In: *2012 12th IEEE-RAS International Conference on Humanoid Robots (Humanoids 2012)*. IEEE; p. 698–703 (2012)
52. Rodríguez-Lera, F. J., Martín-Rico, F., Matelán-Olivera, V.: Generating symbolic representation from sensor data: inferring knowledge in robotics competitions. In: *2018 IEEE International Conference on Autonomous Robot Systems and Competitions (ICARSC)*. IEEE; p. 261–266 (2018)

53. Ganapini, M. B., Campbell, M., Fabiano, F., Horesh, L., Lenchner, J., Loreggia, A., et al. Combining Fast and Slow Thinking for Human-like and Efficient Navigation in Constrained Environments. arXiv preprint [arXiv:2201.07050](https://arxiv.org/abs/2201.07050). (2022)
54. Garnelo, M., Arulkumaran, K., Shanahan, M.: Towards deep symbolic reinforcement learning. arXiv preprint [arXiv:1609.05518](https://arxiv.org/abs/1609.05518). (2016)
55. Garnelo, M., Shanahan, M.: Reconciling deep learning with symbolic artificial intelligence: representing objects and relations. *Curr. Opin. Behav. Sci.* **29**, 17–23 (2019). <https://doi.org/10.1016/j.cobeha.2018.12.010>. (SI: **29: Artificial Intelligence (2019)**)
56. Amado, L., Pereira, R. F., Aires, J., Magnaguagno, M., Granada, R., Meneguzzi, F.: Goal recognition in latent space. In: 2018 International Joint Conference on Neural Networks (IJCNN). IEEE; p. 1–8 (2018)
57. Asai, M., Fukunaga, A.: Classical Planning in Deep Latent Space: Bridging the Subsymbolic-Symbolic Boundary. *Proceedings of the AAAI Conference on Artificial Intelligence*. 32(1), (2018). <https://doi.org/10.1609/aaai.v32i1.12077>
58. Asai, M.: Unsupervised grounding of plannable first-order logic representation from images. In: *Proceedings of the International Conference on Automated Planning and Scheduling*. vol. 29; . p. 583–591 (2019)
59. Konidaris, G., Kaelbling, L.P., Lozano-Perez, T.: From skills to symbols: Learning symbolic representations for abstract high-level planning. *Journal of Artificial Intelligence Research*. **61**, 215–289 (2018)
60. Ghallab, M., Howe, A., Knoblock, C., Mcdermott, D., Ram, A., Veloso, M., et al.: PDDL—The Planning Domain Definition Language. Available from: <http://citeseerx.ist.psu.edu/viewdoc/summary?doi=10.1.1.37.212>
61. Oddi, A., Rasconi, R., Santucci, V. G., Sartor, G., Cartoni, E., Mannella, F., Baldassarre, G. (2020). An intrinsically motivated planning architecture for curiosity-driven robots. 6th Italian Workshop on Artificial Intelligence and Robotics, AIRO 2019 (Vol. 2594, pp. 19–24). CEUR-WS.
62. Sutton, R.S., Precup, D., Singh, S.: Between MDPs and semi-MDPs: A framework for temporal abstraction in reinforcement learning. *Artif. Intell.* **112**(1), 181–211 (1999)
63. Younes, H. L., Littman, M. L.: PPDDL 1.0: An extension to PDDL for expressing planning domains with probabilistic effects. *Techn Rep CMU-CS-04-162*. 2:99 (2004)
64. Schmidhuber, J.: Formal theory of creativity, fun, and intrinsic motivation (1990–2010). *IEEE Trans. Auton. Ment. Dev.* **2**(3), 230–247 (2010)
65. Santucci, V. G., Baldassarre, G., Mirolli, M.: Biological Cumulative Learning through Intrinsic Motivations: A Simulated Robotic Study on the Development of Visually-Guided Reaching. In: *EpiRob*. Citeseer; (2010)
66. Barto, A., Mirolli, M., Baldassarre, G.: Novelty or surprise? *Front. Psychol.* **4**, 907 (2013)
67. Mirolli, M., Baldassarre, G.: Functions and mechanisms of intrinsic motivations: The knowledge versus competence distinction. *Intrinsically motivated learning in natural and artificial systems*. p. 49–72 (2013)
68. Oudeyer, P. Y., Kaplan, F.: What is intrinsic motivation? A typology of computational approaches. *Frontiers in neurorobotics*. p. 6 (2009)
69. Ester, M., Kriegel, H. P., Sander, J., Xu, X.: A Density-Based Algorithm for Discovering Clusters in Large Spatial Databases with Noise. In: *Proceedings of the Second International Conference on Knowledge Discovery and Data Mining*. KDD'96. AAAI Press; p. 226–231 (1996)
70. Cortes, C., Vapnik, V.: Support-Vector Networks. *Mach. Learn.* **20**(3), 273–297 (1995)
71. Rosenblatt, M.: Remarks on Some Nonparametric Estimates of a Density Function. *Ann. Math. Stat.* **27**(3), 832–837 (1956)
72. Parzen, E.: On Estimation of a Probability Density Function and Mode. *Ann. Math. Stat.* **33**(3), 1065–1076 (1962)
73. Rolf, M., Steil, J.J., Gienger, M.: Goal babbling permits direct learning of inverse kinematics. *IEEE Trans. Auton. Ment. Dev.* **2**(3), 216–229 (2010)
74. Ecoffet, A., Huizinga, J., Lehman, J., Stanley, K. O., Clune, J.: Go-explore: a new approach for hard-exploration problems. arXiv preprint [arXiv:1901.10995](https://arxiv.org/abs/1901.10995). (2019)
75. Bharadhwaj, H., Garg, A., Shkurti, F.: Leaf: Latent exploration along the frontier. In: 2021 IEEE International Conference on Robotics and Automation (ICRA). IEEE; p. 677–684 (2021)
76. Jin, B., Xu, X.: Wholesale price forecasts of green grams using the neural network. *Asian Journal of Economics and Banking*. (2024);(ahead-of-print)
77. Jin, B., Xu, X.: Forecasting wholesale prices of yellow corn through the Gaussian process regression. *Neural Comput. Appl.* **36**(15), 8693–8710 (2024)
78. James, S., Rosman, B., Konidaris, G.: Autonomous learning of object-centric abstractions for high-level planning. In: *International Conference on Learning Representations*; (2021)

Publisher's Note Springer Nature remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.