

TOWARD DISCOVERING OPTIONS THAT ACHIEVE FASTER PLANNING

Anonymous authors

Paper under double-blind review

ABSTRACT

We propose a new objective for option discovery that emphasizes the computational advantage of using options in planning. In a sequential machine, the speed of planning is proportional to the number of elementary operations used to achieve a good policy. For episodic tasks, the number of elementary operations depends on the number of options composed by the policy in an episode and the number of options being considered at each decision point. To reduce the amount of computation in planning, for a given set of episodic tasks and a given number of options, our objective prefers options with which it is possible to achieve a high return by composing few options, and also prefers a smaller set of options to choose from at each decision point. We develop an algorithm that optimizes the proposed objective. In a variant of the classic four-room domain, we show that 1) a higher objective value is typically associated with fewer number of elementary planning operations used by the option-value iteration algorithm to obtain a near-optimal value function, 2) our algorithm achieves an objective value that matches it achieved by two human-designed options 3) the amount of computation used by option-value iteration with options discovered by our algorithm matches those with the human-designed options, 4) the options produced by our algorithm also make intuitive sense—they seem to move to and terminate at the entrance of each room.

1 INTRODUCTION

The options framework (Sutton, Precup, Singh 1999) is a way to achieve temporal abstraction, which is perceived as a cornerstone of artificial intelligence. The options are extended courses of actions, defining different ways of behaving. If the agent has a set of options, it could learn a model that predicts the outcomes of executing the options. One of the most important ways options can be useful is that planning with option models could be much faster than planning with action models because options specify jumpy moves.

A natural question to ask is where options come from, which is a challenging active research frontier (Section 17.2, Sutton & Barto 2018). The first and maybe the most important step towards the option discovery problem is to specify what options should be discovered, which involves defining an objective that can be used to rank options. Existing works provide various objectives. For example, some works argue that good options should lead to subgoal states in the environment and the subgoal states could be, for example, bottlenecks (McGovern & Barto 2001, Menache et al. 2002, van Dijk & Polani 2011, Bacon 2013) or salient events (Singh, Barto, Chentanez 2010) in the environment. Some works propose that good options are those such that, when choosing from them, the agent can achieve high performance in a given task(s) (Bacon, Harb, Precup 2017, Khetarpal et al. 2020, Veeriah et al. 2021). Others argue that a good set of options should, for example, be diversified and also part of solutions that achieve a high average return (e.g., Eysenbach et al. 2018), accelerate learning in future tasks (Brunskill & Li 2014), be represented efficiently (Solway et al. 2014, Harutyunyan et al. 2019), allow representing the value function easier (Konidaris & Barto 2009), improve exploration (Machado et al. 2017a;b, Jinnai et al. 2019), etc. Among all of the existing objectives, one (Jinnai et al. 2019) emphasizes the importance of the role of options in planning. Jinnai et al. proposes an algorithm that searches for the smallest set of "point options" so that planning converges in less than a given number of iterations in one task. A point option is a special kind of option – it can only be initialized in one state and terminate in a (possibly different) state. However, it is not appealing to restrict our focus on point options because more general options could achieve much faster planning.

We propose a simple objective that also emphasizes the importance of options in planning and the discovered options are general. We start by considering that all options can be initiated at any state. In this case, our objective prefers options with which fewer options are required to compose good solutions to multiple given tasks. The multi-task setting is critical to our objective. In fact, our objective reduces to the it proposed by Harb et al.(2018) if there is only one task. We argue in Section 3 that, with only one task, the best options degenerate to the entire solution to the given task and are typically not useful for other tasks that we may want to solve with the discovered options.

Some options may be interesting while others are not at a state. When considering an option at a state, it is appealing to only consider interesting options to save computation. For example, if the agent is full then it is not interesting to consider how to find food. In order to learn to reduce the number of options considered at each state, we generalize the definition of options by replacing the initiation set of an option with an *interest probability function*. The option’s initiation set can then be obtained by sampling according to the interest function and becomes stochastic. Our complete objective encourages options’ interests to be smaller. Therefore uninteresting options are less likely to be sampled and thus computation is saved.

We propose an algorithm to optimize the proposed objective. If there is only one task and all options can be initiated everywhere, our algorithm reduces to Harb et al.’s algorithm. We tested our algorithm in the four-room domain, which is a classic domain evaluating option algorithms. Empirical results show that the algorithm’s discovered options are comparable with two human-designed options that move to entrances of the four rooms in terms of the number of elementary operations used to achieve near-optimal value function used by option-value iteration.

2 THE OPTIONS FRAMEWORK WITH MULTIPLE TASKS

Our new option discovery objective involves solving a finite set of episodic tasks $\mathcal{N}$, all of which share the same state space $\mathcal{S}$ and action space $\mathcal{A}$. The state and action spaces are finite. For each task n , there is a corresponding set of terminal states $\perp^n$, which is a set containing one or multiple states in $\mathcal{S}$. The transition dynamics are shared across different tasks, except that starting from a terminal state of a task the agent stays at the state regardless of the action it takes. The reward settings for different tasks are typically different. Let $p^n : \mathcal{S} \times \mathcal{R} \times \mathcal{S} \times \mathcal{A} \rightarrow [0, 1]$ be the transition function of task n with $p^n(s', r | s, a) \doteq$ the probability of resulting in state s' and reward r given state-action pair (s, a) and task n .

The agent’s interaction with the environment produces a sequence of episodes. The first episode starts from state S_0 sampled from an initial state distribution d_0 . A task N_0 is chosen randomly from $\mathcal{N}$ and remains unchanged within the episode. The agent observes S_0 and N_0 and takes an action A_0 . The environment then emits task N_0 ’s reward R_1 and the next state S_1 according to the transition function p^{N_0} . Such an agent-environment interaction keeps going until the end of the episode when the agent reaches a terminal state. Let the time step at which the first episode terminates T . A new episode starts at $T + 1$, with a new initial state S_{T+1} sampled from d_0 and a new task N_{T+1} sampled from $\mathcal{N}$. The agent may choose its actions following a stationary policy $\pi : \mathcal{A} \times \mathcal{S} \rightarrow [0, 1]$ with $\pi : (a | s) \doteq$ the probability of choosing action a at state s . Denote the set of stationary policies Π .

Given a task n , assume that every stationary policy reaches a terminal state with positive probability in at most $|S \setminus \perp^n|$ steps, regardless of the initial state. Under this assumption, the value function of a policy $\pi \in \Pi$ in task $n \in \mathcal{N}$, $v_\pi^n(s)$ can be defined: $v_\pi^n(s) \doteq \mathbb{E}[R_1 + \dots + \gamma^{T-1}R_T | S_0 = s, N_0 = n, A_{0:T-1} \sim \pi]$, where T is a random variable denoting the time step at which the episode terminates and $\gamma \in [0, 1]$ is the discount factor. The optimal value function of task n is defined to be: $v_*^n(s) \doteq \max_{\pi \in \Pi} v_\pi^n(s)$. Here the max always exists. Standard reinforcement learning algorithms like Q-learning can be applied to obtain optimal values for these tasks (Section 5.6 Bertsekas & Tsitsiklis 1996).

The agent may maintain a set of options, which are behaviors that typically take more than one-step to finish (Sutton et al., 1999). An option o is a tuple $(\mathcal{I}_o, \pi_o, \beta_o)$, where $\mathcal{I}_o \subseteq \mathcal{S}$ is the option’s initiation set, $\pi_o \in \Pi$ is the option’ policy, and $\beta_o \in \Gamma$, where $\Gamma : \{f | f : \mathcal{S} \rightarrow [0, 1]\}$, is the option’s termination function. Denote $\mathbf{O}$ as the space of all possible options. That is, $\mathbf{O} \doteq \Gamma \times \Pi \times \Gamma$.

Suppose the agent has a set of k adjustable options, which can be learned from data, in addition to $|\mathcal{A}|$ primitive actions (non-adjustable options), making it a total of $|\mathcal{A}| + k$ options. Denote $\mathcal{O}$ as the

entire set of options the agent maintains. Let $\mathcal{H} \doteq \{1, 2, \dots, k + |\mathcal{A}|\}$ denote the set of indices of options¹ and $\mathcal{H}^{adj} \doteq \{1, 2, \dots, k\}$ denote the set of indices of adjustable options. Therefore options with indices $k + 1, \dots, k + |\mathcal{A}|$ correspond to the primitive actions. For each task n , we define a *meta-preference function* $f^n \in \mathcal{S} \times \mathcal{H} \rightarrow [0, \infty)$, representing the agent’s willing to choose each option at every state for task n . Let $\mathcal{F} \doteq \{f^n : n \in \mathcal{N}\}$ be the set of meta-preference functions. Note that a meta-preference function chooses from indices of options rather than options themselves because the agent’s options can change over time and while the indices don’t.

Respecting the fact that the agent chooses from option indices, we make the following definitions for precise presentation. For each s , define the set of indices whose associated options can be initiated at s : $\Omega(s) \doteq \{o \in \mathcal{O} : s \in \mathcal{I}_o\}$. Define a function $\pi_{\mathcal{O}} : \mathcal{A} \times \mathcal{S} \times \mathcal{H} \rightarrow [0, 1]$ with $\pi_{\mathcal{O}}(a | s, h) \doteq \pi_{o_h}(a | s) \forall a \in \mathcal{A}, s \in \mathcal{S}, h \in \mathcal{H}$, and a function $\beta_{\mathcal{O}} : \mathcal{S} \times \mathcal{H} \rightarrow [0, 1]$ with $\beta_{\mathcal{O}}(s, h) \doteq \beta_{o_h}(s), \forall s \in \mathcal{S}, h \in \mathcal{H}$. Note that the termination probabilities for non-adjustable options (primitive actions) are 1. That is, $\beta_{\mathcal{O}}(s, h) = 1, \forall s \in \mathcal{S}, \forall h \in \{k + 1, \dots, k + |\mathcal{A}|\}$. A primitive action’s policy chooses the action deterministically. That is, $\pi_{\mathcal{O}}(a | s, h) = \mathbf{I}(a = a_h), \forall s \in \mathcal{S}, h \in \{k + 1, \dots, k + |\mathcal{A}|\}, a \in \mathcal{A}$.

In order to choose actions, the agent could choose to execute its options. Such a way of behaving is called *call-and-return*. Specifically, at the first time step of each episode, the agent observes a task N_0 randomly chosen from $\mathcal{N}$ and a state S_0 randomly chosen from d_0 , it then chooses an option index H_0 with probability

$$\mu_{\Omega(S_0)}^{N_0}(H_0 | S_0) \doteq \frac{f^{N_0}(S_0, H_0)}{\sum_{h \in \Omega(S_0)} f^{N_0}(S_0, h)}, \quad (1)$$

and takes action A_0 with probability $\pi(A_0 | S_0, H_0)$. Having observed the next state S_1 and reward S_1 , the agent either keeps following H_0 with probability $1 - \beta(S_1, H_0)$ or terminates H_0 with probability $\beta(S_1, H_0)$. The task remains unchanged through the episode $N_0 = N_1 = \dots$ and the above process keeps going until the episode terminates. For simplicity, let Z_{t+1} denote the termination signal of the episode. That is, $Z_{t+1} = 0$ if the current episode continues ($S_{t+1} \notin \perp^{N_t}$) and $Z_{t+1} = 1$ if the the current episode is terminates ($S_{t+1} \in \perp^{N_t}$).

Define $\bar{q}_{\mathcal{O}, \mathcal{F}}^n(s, h, a)$ to be the expected cumulative reward with a set of options $\mathcal{O}$ and a set of preferences $\mathcal{F}$ if the agent starts from state s , chooses option o_h , and action a , and behaves in the call-and-return fashion. Define option-value function $\bar{q}_{\mathcal{O}, \mathcal{F}}^n(s, h) \doteq \sum_a \pi_{\mathcal{O}}(a | s, h) \bar{q}_{\mathcal{O}, \mathcal{F}}^n(s, h, a)$, and value function $\bar{v}_{\mathcal{O}, \mathcal{F}}^n(s) \doteq \sum_h \mu^n(h | s) \bar{q}_{\mathcal{O}, \mathcal{F}}^n(s, h)$ where $\mu^n(h | s)$ is the marginal probability of choosing option with index h at state s for task n , $\mu^n(h | s) \doteq \sum_{\Omega \in \mathcal{P}(\mathcal{H})} \Pr(\Omega | s) \mu_{\Omega}^n(h | s)$, where μ_{Ω}^n is defined in (1), and $\mathcal{P}(\mathcal{H})$ is the power set of $\mathcal{H}$. It can be seen that the best achievable value function when choosing and following options is the optimal value function. That is, $\max_{\mathcal{O} \in \mathbf{O}^k, \mathcal{F} \in \mathbf{F}^{|\mathcal{N}|}} \bar{v}_{\mathcal{O}, \mathcal{F}}^n(s) = \bar{v}_*^n(s), \forall n \in \mathcal{N}, s \in \mathcal{S}$. This equation holds because any pair $(\mathcal{O}, \mathcal{F})$ defines a non-stationary flat policy (policy over primitive actions). We know for any non-stationary flat policy π , $\bar{v}_{\pi}^n(s) \leq \bar{v}_*^n(s), \forall n, s$ (Puterman 1994). Also, by setting $\mathcal{O}$ and $\mathcal{F}$ appropriately, all policies that only choose from primitive actions can be obtained and therefore an optimal policy, thus $\max_{\mathcal{O}, \mathcal{F}} \bar{v}_{\mathcal{O}, \mathcal{F}}^n(s) = \bar{v}_*^n(s), \forall n \in \mathcal{N}, s \in \mathcal{S}$.

3 AN OPTION DISCOVERY OBJECTIVE FOR FASTER PLANNING

In this section, we formally define our new objective and explain why it enables faster planning.

The speed of planning is proportional to the number of elementary operations in a sequential machine. To understand how the number of elementary operations used in planning is affected by a set of options, consider applying value iteration, which is a classic iterative planning algorithm, to a set of options. The complexity of the value iteration algorithm for computing a near-optimal policy is the number of operations per-iteration times the number of iterations required to achieve a near-optimal value function, from which a near-optimal policy can be derived. The algorithm requires a model as input. For each iteration, the algorithm queries the model, for each state-option pair, the probability distribution over next states. Therefore the per-iteration complexity is $\sum_{s \in \mathcal{S}} \Omega(s) \times |\mathcal{S}|$. This shows that the per-iteration complexity is proportional to the average number of options that can be initiated at each state.

¹Throughout the paper, all sets are indexed and given a set $\mathcal{X}$, we use the notation x_i to denote the i -th element of $\mathcal{X}$.

A direct way to estimate the number of required iterations to achieve a near-optimal value function is to apply value iteration directly. This would be just fine if the set of options are fixed and our only goal is to evaluate these options. However, if the set of options are continually changing, which is unavoidable in the option discovery problem, estimating the iterations in this way would have several drawbacks. First, this approach is computationally expansive because whenever an option changes, its model needs to be re-estimated and value iteration needs to be re-applied. Second, it is unclear how to improve the options to reduce the number of iterations.

We consider estimating the other quantity, the expected number of options being executed by a good policy to reach the terminal state. This quantity can be estimated by a model-free learning algorithm so that neither model learning nor planning is involved. Furthermore, as we will show shortly, it is possible to derive an algorithm that adjusts options to reduce this quantity.

This quantity is also closely related to the number planning iterations. To understand their relation, consider the MRP shown in Figure 1. Value iteration would need 3 iterations to find the policy shown in the figure (with all estimated values pessimistically initialized). The expected number of options executed by the policy to reach the terminal state, is 2.5. The difference comes from the stochasticity in option transitions – without stochasticity, it is clear that the number of iterations used by value iteration to find a policy is exactly the number of options executed by the policy to reach the terminal state.

We are now ready to present our new objective, which we call the *fast planning option discovery objective*. Intuitively, the new objective prefers a set of options such that 1) there exists a policy choosing from them achieving a high expected return and 2) the total number of options considered at each decision point when following the policy is small. Note that the second point combines the need of having a small average number of options considered at each state and the need of having a small number of options chosen by a policy to reach the terminal state.

Define $\tilde{q}_{\mathcal{O}, \mathcal{F}}^n(s, h, a) \doteq -\mathbb{E}[|\Omega(S_0)| + |\Omega(S_{t_1})| + \dots + |\Omega(S_{t_{N-1}})| | S_0 = s, H_0 = h, A_0 = a]$ as the negative expected cumulative number of options that the agent needs to consider before reaching the terminal state, if the agent starts from state s , chooses option o_h , and action a , and follows policy μ^n . Here t_i denotes the time step at which the $i - 1$ th option terminates and t_N is the time step at which the terminal state is reached. Let $\tilde{q}_{\mathcal{O}, \mathcal{F}}^n(s, h) \doteq \sum_a \pi_{\mathcal{O}}(a | s, h) \tilde{q}_{\mathcal{O}, \mathcal{F}}^n(s, h, a)$, and $\tilde{v}_{\mathcal{O}, \mathcal{F}}^n(s) \doteq \sum_h \mu^n(h | s) \tilde{q}_{\mathcal{O}, \mathcal{F}}^n(s, h) - |\Omega(s)|$.

We propose to maximize the following objective w.r.t. $\mathcal{O}, \mathcal{F}$:

$$J(\mathcal{O}, \mathcal{F}, c) \doteq \sum_n \sum_s d_0(s) \tilde{v}_{\mathcal{O}, \mathcal{F}}^n(s) + c \sum_n \sum_s d_0(s) \tilde{v}_{\mathcal{O}, \mathcal{F}}^n(s), \quad (2)$$

where $c > 0$ is a problem parameter. The objective is a weighted sum of two terms. The first term is the expected values weighted over initial state distribution summed over all tasks. The second term is the total number of options that the agent will have to consider choosing from within an episode when following meta-policies μ^n , again starting from states sampled from d_0 , summed over all tasks. The problem parameter c specifies the agent’s relative interests over the two terms.

Remark: If there is only one task and all options can be initiated at all states, our objective J reduces to the objective proposed by Harb et al. (2018). However, with only one task, the best strategy is to learn an option that solves the entire task and to always choose that option so that there would be only one option executed by the policy and the second term in J is minimized. On the contrary, if there are multiple tasks and the number of tasks is more than the number of options, the agent can not assign for each task an option, and learned options would better be an *overlapped part* of the solutions to different tasks. The next example illustrates this point.

Example. Consider a two-room domain with four episodic tasks shown in Figure 2a. Each green cell marks a terminal state of a task. All rewards are -1. Each episode starts from a state randomly picked from the left room. The agent has four primitive actions (left, right, up, down) and

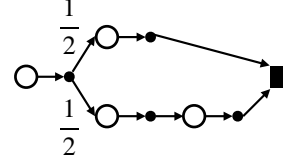


Figure 1: Example illustrating of the relation of the number of planning iterations and the number of options being executed by a good policy. The empty circles are normal states. The solid square is a terminal state. The solid circles are options. For each state there are multiple options can be chosen (not shown) and the one shown in the figure is the best one. From the bottom state, taking the option shown in the figure moves to the three states in the middle with equal probabilities.

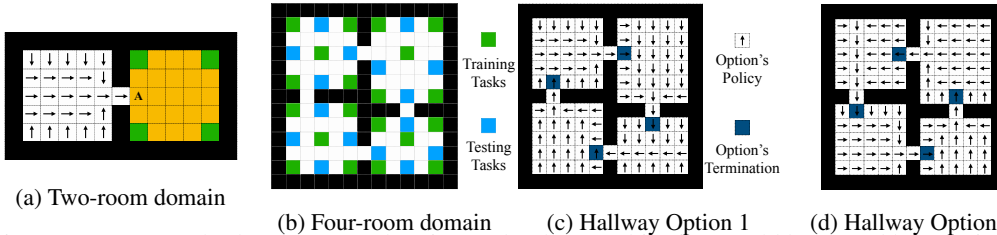


Figure 2: (a) Example showing a good option under the our objective would be shared by solutions to multiple tasks. (b) Illustration of the four-room domain. In both two- and four-room domains, the agent is randomly initialized over all empty (white + green + blue) cells. All rewards are -1 and the discount factor is 1. (c) The policy and the termination probability of the first hallway option. This option’s policy travels through the four rooms in the clock-wise direction. This option is terminated whenever the agent moves to a new room (marked by deep blue colored cells). (d) The policy and the termination probability of the second hallway option, which travels in the counter-clock-wise direction.

one adjustable option. Consider an option that can be initiated everywhere. For cells in the left room and the hallway cell, actions chosen by the option’s policy are marked by arrows. For cells in the right room, the option’s policy is uniformly random. The option terminates deterministically at yellow cells. Note that the option is shared by solutions of all four tasks. Also, note that if the option terminates at the hallway cell instead of cell *A*, solutions to all four tasks need to take a primitive action `right`, introducing an additional decision to make. Therefore this option is worse than the one plotted according to our objective. On the other hand, if the option does not terminate at cell *A*, no matter what action is assigned to this cell, the action is not optimal for some tasks. As long as c is relatively smaller than the magnitude of the per-step reward. The hypothesized option induces a worse objective value compared with the one plotted in the figure.

4 DISCOVERING OPTIONS THAT CAN BE INITIATED EVERYWHERE

In this section, we consider the case where all options can be initiated at all states. Therefore the options being considered at each decision point is the whole set of options (primitive actions plus adjustable options). In this case, the second term in the objective is proportional to the expected number of options being executed in each episode. Recall that our objective reduces to the objective by Harb et al. (2018) if there is a single task and all options are considered at every state, therefore a multi-task extension of the single-agent tabular version of the Asynchronous Advantage Option-Critic (A2OC) Algorithm (Harb et al. 2018) can be applied to optimize the objective². We call this multi-task extension the *Multi-task advantage Option-Critic* (MAOC) algorithm. The description of the algorithm is provided in Appendix A.

We now present an experiment to evaluate 1) whether the MAOC algorithm achieves a high objective value across a set of training tasks and 2) whether, with the set discovered options, planning to obtain near-optimal value functions for a set of testing tasks requires fewer planning iterations. The experiment consists of a training phase and a testing phase. In the training phase, the agent observes a set of training tasks and learns options to solve these tasks. In the testing phase, the agent fixes the learned options, learns a model of these options, and plans with the learned model to solve a set of testing tasks, which are similar to the training tasks. The planning algorithm is the classic value iteration algorithm applied to both primitive actions and learned options. If the algorithm uses fewer iterations to achieve near-optimal performance than the same algorithm using only primitive actions, we say that the learned options enable faster planning. In this experiment, the environment is a four-room grid world with 20 training tasks corresponding to 20 green cells, and 12 testing tasks corresponding to 16 blue cells (Figure 2b). At the beginning of each episode, the agent observes a task that is randomly sampled from the 20 training tasks in the training phase (or the 16 testing tasks in the testing phase). The initial state of the episode is uniformly randomly sampled from all empty (white + green + blue) cells. The green/blue cell corresponding to the current task is

²The A2OC algorithm is a multi-agent algorithm because there are multiple agents sharing parameters simultaneously interacting with their own copies of the environment.

the only cell that terminates the episode when the agent reaches it. All other cells, including other green/blue cells, are just normal states. The agent has four actions, up, down, left, right, which deterministically take the agent to the intended neighbor cell if the cell is empty. Otherwise, the agent stays at its current cell. The reward is -1 for each step, including the step transitioning to the terminal state, regardless of the action the agent takes. The experiment consists of 10 runs, each of which consists of 10^8 training steps. Every 10^6 training steps, the agent is evaluated for 500 episodes. No parameter updates are performed in evaluation episodes. At the beginning of each evaluation episode, the agent chooses an option that it finds to be the best, using a greedy policy w.r.t. its option-value estimate Q . The agent then executes the chosen option until it terminates. The agent then repeats the above process until the episode terminates. For each evaluation episode, we compute the compound return, which is the return minus the $c \times |\mathcal{H}| \times$ number of decision points, where c is chosen to be 0.2. We use the *average compound return* to denote the compound return averaged over 500 evaluation episodes. The agent uses two adjustable options.

We report a learning curve of the algorithm in Figure 3a. The parameter setting used is the one that results in the highest compound return averaged over the last 5 times evaluation and 10 runs. In the same figure, we also show in the dotted line the best compound return given only primitive actions, which terminate every time step. In addition, we show the best compound return given two hallway options (Figure 2c, Figure 2d) that move to and terminate at entrances of rooms and the primitive actions. Details of the experiment are presented in Appendix B.3.

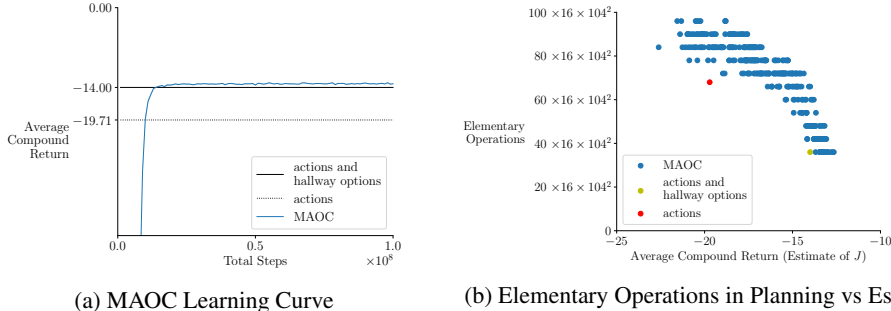
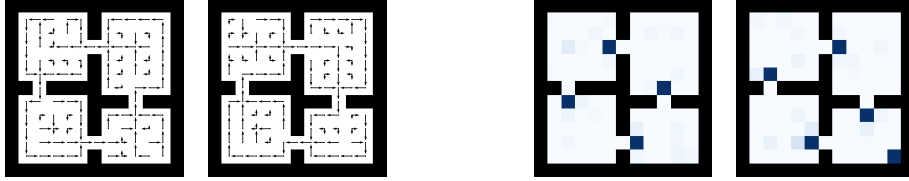


Figure 3: (a) Learning curve of the MAOC algorithm (with initiation sets containing all states) shows that the algorithm reliably achieves a compound return that is even slightly better than the return achieved with the two built-in hallway options. The x -axis is the number of training steps, the y -axis is the average compound return. The dotted line represents the best compound return given only four primitive actions, in which case options (primitive actions) terminates every time step. The solid line represents the best possible compound return given two built-in hallway options and four primitive actions. The shading area represents the standard error over 10 runs. (b) Relation between the average compound return (estimate of the objective with $c = 0.2$) in a set of training tasks and the number of elementary operations used by option-value iteration to achieve near-optimal values in a set of testing tasks. The best discovered options match the human-designed options in terms of number of elementary operations in planning.

Figure 3b helps understand the relation between the objective J , defined over the set of *training* tasks (blue colored cells in Figure 2b), and the number of planning operations used to achieve near-optimal value functions for the set of *testing* tasks (green colored cells in Figure 2b). Each blue dot is obtained by performing an experiment with a different parameter setting or a different seed. The x -axis is the estimate average compound return. The y -axis is the total number of operations used to obtain near-optimal value functions for all testing tasks. As a baseline, we plot a yellow dot, corresponding to the number of operations and the best average compound return when the set of options consists of two built-in hallway options (Figure 5) and primitive actions. We also plot in red the number of planning operations given only primitive actions. It can be seen that the average compound return and the number of iterations are negatively correlated – higher estimate of the objective value for training tasks is typically associated with fewer planning operations for testing tasks. In addition, it can be seen that the discovered options in the best runs achieve the same number of planning operations compared with the two built-in options.

We also plot the learned options in Figure 4. The two discovered options are pretty similar to the human-designed options. First, in both cases, there is one option that moves in the clockwise direction

and the other one moves in the opposite direction. Second, in both cases, the two options terminate at the entrances of rooms. There is one difference between the two sets of options. At the entrance cell of a room, our discovered option sometimes chooses to not go directly to the next room, but visits the edges and corners of the room (e.g., see the first option’s policy in the upper right room). This is actually because the goals of training tasks are distributed in the corners of the room and this option can be used by these training tasks to reduce the cost term in J . On the other hand, for the two built-in options, the agent has to take primitive actions to reach goals in corners of a room once it enters that room. This is also the reason why the discovered options achieve a slightly higher objective value compared with the built-in options.



(a) Learned Options’ Policy (b) Learned Options’ Termination Probabilities
Figure 4: The learned option by the MAOC algorithm using the same parameter setting used to produce the learning curve Figure 3a. For termination probabilities, darker blue means terminating with higher probability.

5 FROM INITIATION SETS TO INTEREST FUNCTIONS

In many of the follow-up papers (Machado et al. 2017a;b, Bacon et al. 2017, Harb et al., 2018, Harutyunyan., 2019) of the original options framework paper by Sutton et al. (1999), the initiation sets degenerate as they are assumed to be fixed and contain all states. A natural question to ask is why would we ever want to discover initiation sets.

One possible justification is that an option’s initiation set represent ‘affordance’ of the option. That is, whether or not an option *can be* initiated at a state. However, this justification appears questionable to us because it does not justify why options that can not be executed at certain states is better than options that can be initiated everywhere. The computational justification, introduced in Section 3, appears to be more sound and natural to us. That is, initiation sets would better to be more selective so that fewer options are considered at each state to save computation in planning. We believe that this is the first time that the advantage of using initiation sets is justified from the computational perspective.

The question then becomes how to learn these sets from data. A possible solution is the generate-and-test approach – one just randomly generates these sets and evaluates the corresponding objective value of J , and chooses the one that results in the highest value. However, it is clear that this approach is computationally expensive. We look for a more efficient algorithm searching for initiation sets. To this end, we generalize the definition of options by allowing initiation sets to be stochastic. With this modification, it is possible to apply stochastic gradient descent to reduce number of options that can be initiated at each state.

We now provide the more general definition of options. An option o is a tuple (i_o, π_o, β_o) , where $i_o \in \Gamma$ is the option’s interest function, $\Gamma : \{f \mid f : \mathcal{S} \rightarrow [0, 1]\}$, $\pi_o \in \Pi$ is the option’s policy, and $\beta_o \in \Gamma$ is the option’s termination function. At each state s , the options that can be initiated are sampled according to their interest functions: $\Pr(o \in \Omega(s)) = i_o(s)$. In addition, assume that primitive actions can be initiated at all state $i_a(s) = 1 \forall s \in \mathcal{S}, \forall a \in \mathcal{A}$. Finally, just as what we defined in Section 2, given a set of options $\mathcal{O}$, define a function $i_{\mathcal{O}} : \mathcal{S} \times \mathcal{H} \rightarrow [0, 1]$ with $i_{\mathcal{O}}(s, h) \doteq i_{o_h}(s), \forall s \in \mathcal{S}, h \in \mathcal{H}$.

The call-and-return agent-environment interaction is also modified given the new definition of options. Suppose that the agent is solving task N by executing a set of options, with the set $\mathcal{H}$ denoting the indices of options, and suppose that at state S , the previous option terminates and a new option needs to be chosen, the agent then samples a set $\Omega \in \mathcal{P}(\mathcal{H})$, with $\Pr(\Omega \mid s) \doteq \prod_{h \in \Omega} i(s, h) \prod_{\bar{h} \notin \Omega} i(s, \bar{h})$

probability, and chooses an option indexed by H with probability $\mu_{\Omega}^N(H | S)$. Note that Ω can not be empty because primitive actions' interests are 1 and are thus in Ω .

The idea of using interest functions in options has also been explored by Khetarpal et al. (2020). Note that their interest functions are not necessarily probabilities and are different as ours. The most important difference is that the interest functions in our work are used to generate initiation sets while the initiation sets are no longer considered in their work. While they have observed that smaller initiation sets reduce the computational cost, they replaced initiation sets with interest functions, and still allowed all options to be initiated at every state.

6 DISCOVERING INTEREST FUNCTIONS

In order to maximize the objective J by updating the options and the meta-preference functions, it is common to apply stochastic gradient ascent. Suppose that adjustable options' interests, policies, termination probabilities, as well as the meta-policy preferences are jointly parameterized by θ . In this case, both $\mathcal{O}$ and $\mathcal{F}$ are parameterized by θ . With this in mind, we replace $\mathcal{O}$ and $\mathcal{F}$ by θ and present the gradient w.r.t. J in the following theorem. The proof of the theorem is presented in Appendix B.1.

Theorem 6.1. *Let $v_{\theta}^n(s) \doteq \bar{v}_{\theta}^n(s) + c\tilde{v}_{\theta}^n(s)$, $q_{\theta}^n(s, h) \doteq \bar{q}_{\theta}^n(s, h) + c\tilde{q}_{\theta}^n(s, h)$ and $q_{\theta}^n(s, h, a) \doteq \bar{q}_{\theta}^n(s, h, a) + c\tilde{q}_{\theta}^n(s, h, a)$. Suppose that $J(\theta)$ is differentiable w.r.t. θ ,*

$$\begin{aligned} \nabla J(\theta) = & \sum_n \sum_s d_0(s) m_{\theta}^n(s) + \sum_{n,s,h} d_{\gamma,\theta}^n(s, h) \left(\sum_a \nabla \pi_{\theta}(a | s, h) q_{\theta}^n(s, h, a) \right. \\ & \left. + \gamma \sum_{a,s',r} \pi_{\theta}(a | s, h) p(s', r | s, a) \left(-\nabla \beta_{\theta}(s', h) (q_{\theta}^n(s', h) - v_{\theta}^n(s')) + \beta_{\theta}(s', h) m_{\theta}^n(s') \right) \right), \end{aligned}$$

where $d_{\gamma,\theta}^n(s, h) \doteq \mathbb{E} \left[\sum_{t=0}^T \gamma^t \mathbf{I}(S_t = s, H_t = h) | S_0 \sim d_0, \theta, n \right]$ is the discounted number of time steps on average option o_h is used in state s in a single episode, and

$$\begin{aligned} m_{\theta}^n(s) \doteq & \sum_{\Omega \in \mathcal{P}(\mathcal{H})} \Pr(\Omega | s) \sum_{h \in \Omega} \mu_{\Omega,\theta}^n(h | s) \left(\sum_{\bar{h} \in \Omega} \nabla \log i_{\theta}(s, \bar{h}) + \sum_{\bar{h} \notin \Omega} \nabla \log(1 - i_{\theta}(s, \bar{h})) \right) \\ & + \nabla \log \mu_{\Omega,\theta}^n(h | s) q_{\theta}^n(s, h) - c \sum_{h \in \mathcal{H}} \nabla i_{\theta}(s, h), \end{aligned}$$

and $\mu_{\Omega,\theta}^n(h | s) \doteq \frac{f_{\theta}^n(s, h)}{\sum_{\bar{h} \in \Omega} f_{\theta}^n(s, \bar{h})}$.

The policy update term $\sum_a \nabla \pi_{\theta}(a | s, h) q_{\theta}^n(s, h, a)$, termination probability update term $-\nabla \beta_{\theta}(s', h) (q_{\theta}^n(s', h) - v_{\theta}^n(s'))$, and meta-policy update term $\sum_{h \in \Omega} \mu_{\Omega,\theta}^n(h | s) \nabla \log \mu_{\Omega,\theta}^n(h | s) q_{\theta}^n(s, h)$ also appear in the gradient of the option-critic objective (Bacon et al. 2017) and the objective by Harb et al. (2017).

The update to the interest function

$$\begin{aligned} & \sum_{\Omega \in \mathcal{P}(\mathcal{H})} \Pr(\Omega | s) \sum_{h \in \Omega} \mu_{\Omega,\theta}^n(h | s) \left(\sum_{\bar{h} \in \Omega} \nabla \log i_{\theta}(s, \bar{h}) + \sum_{\bar{h} \notin \Omega} \nabla \log(1 - i_{\theta}(s, \bar{h})) \right) q_{\theta}^n(s, h) \\ & - c \sum_{h \in \mathcal{H}} \nabla i_{\theta}(s, h), \end{aligned}$$

is new. It can be understood in the following way. The second term $-c \sum_{h \in \mathcal{H}} \nabla i_{\theta}(s, h)$ term encourages all interests to be smaller. Now for the first term, if $\sum_{h \in \Omega} \mu_{\Omega,\theta}^n(h | s) q_{\theta}^n(s, h)$ is high, it means that the set Ω contains good options, then when one update the interest function using the gradient of J , the interests of all the options in Ω will be increased by a large amount and the interests of all the options not in Ω will be decreased by a large amount

Our full algorithm, which we call Fast Planning Option Critic (FPOC), is a tabular algorithm that searches for a pair $(\mathcal{O}, \mathcal{F})$ that maximizes J using a sample of the gradient given in Theorem 6.1. Due to the space limit, we present the full algorithm in Appendix A.

Applying FPOC to the four-room domain the same way as introduced in Section 4, we obtain Figure 5b. This time in order to show the effect of learning interest function, we increase the adjustable options to 8 because with only 2 adjustable options, it seems that both of them are quite useful for many of the states in order to achieve a high return and their interests would better to high. Readers are welcomed to see Figure 8 for the same figure with 2 and 4 adjustable options. As a comparison, we show the same figure but for MAOC in Figure 5a. Note that for MAOC, many of the runs result in options that require more than $100 \times 16 \times 10^4$ operations and are not shown in the figure. It is clear that learning the interest function significantly reduces the elementary operations. Visualization of the discovered options as well as other empirical results can be found in Appendix B.3.

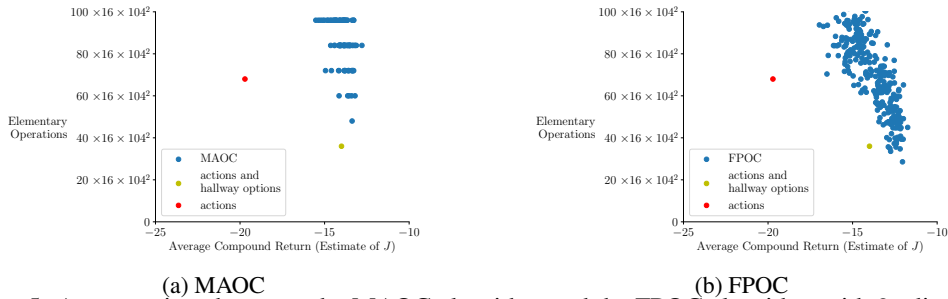


Figure 5: A comparison between the MAOC algorithm and the FPOC algorithm with 8 adjustable options. It can be seen that learning to reduce number of options being considered at each state significantly reduces the number of operations used by option-value iteration.

7 DISCUSSION AND FUTURE WORK

We have introduced a new objective for option discovery that highlights the computational advantage of planning with options when there are a set of given tasks and a given number of options. To the best of our knowledge, this is the first objective that explicitly advocates discovering general options that achieve fast planning. Optimizing our objective maximizes the returns of solutions to the given tasks while minimizing the number of options used to compose the solutions as well as the number of options considered at each decision point. Empirically, we showed that higher objective values are associated with fewer elementary operations to achieve near-optimal value functions. To optimize the proposed objective, a key step is to generalize the classic definition of options so that initiation sets are no longer fixed, but are sampled following the options’ interest functions. We proposed a new algorithm that optimizes the objective by following the sample gradient of the objective. In the four-room domain, we show that the proposed algorithm, with or without learning to adapt the interest function, achieves a high objective value, but also discovers options that are comparable with two human-designed options in terms of the number of planning operations used when applying option-value iteration to solve a set of testing tasks with the options. We believe that our objective and algorithm provide a new perspective for option discovery and the more general temporal abstraction problem.

There are several ways in which our work is limited. The most important way is that our FPOC algorithm only treats the tabular setting. However, we do not see technical difficulty of extending this algorithm to the function approximation (FA) setting and would like to leave the analysis of the FA algorithm as a future work. The second way in which our work is limited is that we assume a given set of tasks. We believe that these tasks are sub-tasks that the agent finds interesting/useful to solve in order to achieve better performance in the agent’s main task and should be discovered by the agent itself. Discovering sub-tasks or sub-problems is a fundamental open research problem. Several typical recent approaches to the problem include Sutton et al. (2022), Veeriah et al. (2021), Nachum et al. (2018), and Vezhnevets et al. (2017). The third way in which our work is limited is that the number of options is assumed to be fixed and specified by a human.

REFERENCES

- Bacon, P. L. (2013). *On the bottleneck concept for options discovery*. McGill University.
- Bacon, P. L. (2018). *Temporal Representation Learning*. McGill University.
- Bacon, P. L., Harb, J., Precup, D. (2017). The Option-Critic Architecture. *AAAI Conference on Artificial Intelligence*.
- Bertsekas, D. P., Tsitsiklis, J. N. (1996). *Neuro-dynamic programming*. Athena Scientific.
- Brunskill, E., Li, L. (2014). PAC-inspired Option Discovery in Lifelong Reinforcement Learning. *International Conference on Machine Learning*.
- Eysenbach, B., Gupta, A., Ibarz, J., Levine, S. (2018). Diversity is All You Need: Learning Skills without a Reward Function. *International Conference on Learning Representations*.
- Harb, J., Bacon, P. L., Klissarov, M., Precup, D. (2018). When waiting is not an option: Learning options with a deliberation cost. *AAAI Conference on Artificial Intelligence*.
- Harutyunyan, A., Dabney, W., Borsa, D., Heess, N., Munos, R., Precup, D. (2019). The termination critic. *International Conference on Artificial Intelligence and Statistics*.
- Jinnai, Y., Abel, D., Hershkowitz, D., Littman, M., Konidaris, G. (2019). Finding options that minimize planning time. *International Conference on Machine Learning*.
- Jinnai, Y., Park, J. W., Abel, D., Konidaris, G. (2019). Discovering options for exploration by minimizing cover time. *International Conference on Machine Learning*.
- Khetarpal, K., Klissarov, M., Chevalier-Boisvert, M., Bacon, P. L., Precup, D. (2020). Options of interest: Temporal abstraction with interest functions. *AAAI Conference on Artificial Intelligence*.
- Konidaris, G., Barto, A. (2009). Skill discovery in continuous reinforcement learning domains using skill chaining. *Neural Information Processing Systems*.
- Machado, M. C., Bellemare, M. G., Bowling, M. (2017). A Laplacian Framework for Option Discovery in Reinforcement Learning. *International Conference on Machine Learning*.
- Machado, M. C., Rosenbaum, C., Guo, X., Liu, M., Tesauro, G., Campbell, M. (2017). Eigenoption discovery through the deep successor representation. *International Conference on Learning Representations*.
- McGovern, A., Barto, A. G. (2001). Automatic Discovery of Subgoals in Reinforcement Learning using Diverse Density. *International Conference on Machine Learning*.
- Menache, I., Mannor, S., & Shimkin, N. (2002). Q-Cut—Dynamic Discovery of Sub-goals in Reinforcement Learning. *European Conference on Machine Learning*.
- Nachum, O., Gu, S. S., Lee, H., Levine, S. (2018). Data-efficient hierarchical reinforcement learning. *Neural Information Processing Systems*.
- Puterman, M. L. (1994). *Markov Decision Processes: Discrete Stochastic Dynamic Programming*. John Wiley & Sons.
- Şimşek, Ö., & Barto, A. G. (2004). Using Relative Novelty to Identify Useful Temporal Abstractions in Reinforcement Learning. *International Conference on Machine Learning*.
- Singh, S., Barto, A. G., Chentanez, N. (2004). Intrinsically Motivated Reinforcement Learning. *Neural Information Processing Systems*.
- Solway, A., Diuk, C., Córdova, N., Yee, D., Barto, A. G., Niv, Y., Botvinick, M. M. (2014). Optimal behavioral hierarchy. *PLoS Computational Biology*.
- Sutton, R. S., Precup, D., Singh, S. (1999). Between MDPs and Semi-MDPs: A Framework for Temporal Abstraction in Reinforcement Learning. *Artificial Intelligence*.
- Sutton, R. S., Barto, A. G. (2018). *Reinforcement Learning: An Introduction*. MIT Press.

van Dijk, S. G., Polani, D. (2011). Grounding Subgoals in Information Transitions. *IEEE Symposium on Adaptive Dynamic Programming and Reinforcement Learning*.

Veeriah, V., Zahavy, T., Hessel, M., Xu, Z., Oh, J., Kemaev, I., van Hasselt, H., Silver, D., Singh S. (2021). Discovery of Options via Meta-Learned Subgoals. *Neural Information Processing Systems*.

Vezhnevets, A. S., Osindero, S., Schaul, T., Heess, N., Jaderberg, M., Silver, D., Kavukcuoglu, K. (2017). Feudal networks for hierarchical reinforcement learning. *International Conference on Machine Learning*.