

**Planning Under Time Constraints
in Stochastic Domains**

Thomas Dean Leslie Pack Kaelbling
Jak Kirman and Ann Nicholson

Department of Computer Science
Brown University
Providence, Rhode Island 02912

CS-93-55
December 1993

Planning Under Time Constraints in Stochastic Domains

Thomas Dean, Leslie Pack Kaelbling, Jak Kirman, Ann Nicholson

Department of Computer Science

Brown University, Providence, RI 02912

{tld, lpk, jak, aen}@cs.brown.edu

Abstract

We provide a method, based on the theory of Markov decision processes, for efficient planning in stochastic domains. Goals are encoded as reward functions, expressing the desirability of each world state; the planner must find a policy (mapping from states to actions) that maximizes future rewards. Standard goals of achievement, as well as goals of maintenance and prioritized combinations of goals, can be specified in this way. An optimal policy can be found using existing methods, but these methods require time at best polynomial in the number of states in the domain, where the number of states is exponential in the number of propositions (or state variables). By using information about the starting state, the reward function, and the transition probabilities of the domain, we restrict the planner's attention to a set of world states that are likely to be encountered in satisfying the goal. Using this restricted set of states, the planner can generate more or less complete plans depending on the time it has available.

Our approach employs several iterative refinement routines for solving different aspects of the decision making problem. We describe the meta-level control problem of *deliberation scheduling*, allocating computational resources to these routines. We provide different models corresponding to optimization problems that capture the different circumstances and computational strategies for decision making under time constraints. We consider *precursor* models in which all decision making is performed prior to execution and *recurrent* models in which decision making is performed in parallel with execution, accounting for the states observed during execution and anticipating future states. We describe experimental results for both the precursor and recurrent problems that demonstrate planning times that grow slowly as a function of domain size.

1 Introduction

In a completely deterministic world, it is possible for a planner simply to generate a sequence of actions, knowing that if they are executed in the proper order, the goal will necessarily result. In nondeterministic worlds, planners must address the question of what to do when things do not go as expected.

The method of triangle tables [Fikes *et al.*, 1972] made plans that could be executed robustly in any circumstance along the nominal trajectory of world states, allowing for certain classes of failures and serendipitous events. It is often the case, however, that an execution error will move the world to a situation that has not been previously considered by the planner. Many systems (SIPE, for example [Wilkins, 1988]) can monitor for plan “failures” and initiate replanning. Replanning is often too slow to be useful in time-critical domains, however. Schoppers, in his universal plans [Schoppers, 1987], gives a method for generating a reaction for every possible situation that could transpire during plan execution; these plans are robust and fast to execute, but can be very large and expensive to generate. There is an inherent contradiction in all of these approaches. The world is assumed to be deterministic for the purpose of planning, but its nondeterminism is accounted for by performing execution monitoring or by generating reactions for world states not on the nominal planned trajectory.

In this paper, we address the problem of planning in nondeterministic domains by taking nondeterminism into account from the very start. There is already a well-explored body of theory and algorithms addressing the question of finding optimal policies (universal plans) for nondeterministic domains. Unfortunately, these methods are impractical in large state spaces. However, if we know the start state, and have a model of the nature of the world’s nondeterminism, we can restrict the planner’s attention to a set of world states that are likely to be encountered on the way to the goal. Furthermore, the planner can generate more or less complete plans depending on the time it has available. In this way, we provide efficient methods, based on existing techniques of finding optimal strategies, for planning under time constraints in non-deterministic domains. Our approach addresses the uncertainty resulting from control error, but not sensor error; in most of the following, we assume certainty in observations, but discuss relaxing this assumption in Section 9.

We assume that the environment can be modeled as a stochastic automaton: a set of states, a set of actions, and a matrix of transition probabilities. In the simplest cases,

achieving a goal corresponds to performing a sequence of actions that results in a state satisfying some proposition. Since we cannot guarantee the length of a sequence needed to achieve a given goal in a stochastic domain, we are interested in building planning systems that minimize the expected number of actions needed to reach a given goal.

In our approach, constructing a plan to achieve a goal corresponds to finding a *policy* (a mapping from states to actions) that maximizes expected performance. Performance is based on the expected accumulated reward over sequences of state transitions determined by the underlying stochastic automaton. The rewards are determined by a *reward function* (a mapping from states to the real numbers) specially formulated for a given goal. A good policy in our framework corresponds to a universal plan for achieving goals quickly on average.

There are dynamic programming algorithms for computing the optimal policy given a stochastic model of the world. They are useful in small to medium-sized state-spaces, but become intractable on very large state-spaces. We address this difficulty by make some informal assumptions about the environments in which we are working that allow us to generate approximate solutions efficiently. In particular, we assume that the environment has the following properties:

- *high solution density*: it is relatively easy to find plausible (though perhaps not optimal) solutions
- *low dispersion rate*: from any given state, there are only a few states to which transitions can be made
- *continuity*: it is reasonable to estimate the values of states by considering the values of near-by states (where distance is measured as the expected number of steps between states)

Many large, realistic planning problems, such as those involving high-level navigation, have these properties.

In the following, we refer to the automaton modeling the environment as the *system* automaton. Instead of generating the optimal policy for the whole system automaton, we formulate a simpler or *restricted* stochastic automaton and then search for an optimal policy in this restricted automaton. The state space for the restricted automaton, called

the *envelope*, is a subset of the states of the system automaton, augmented with a special state OUT that represents being in any state outside of the envelope.

There are two basic types of operations on the restricted automaton. The first is called *envelope alteration* and serves to increase or decrease the number of states in the restricted automaton. The second is called *policy generation* and determines a policy for the system automaton using the restricted automaton. Note that, while the policy is constructed using the restricted automaton, it is a complete policy and applies to all of the states in the system automaton. For states outside of the envelope, the policy is defined by a set of *reflexes* that implement some default behavior for the agent.

The algorithm is implemented as an *anytime* algorithm [Dean and Boddy, 1988], one that can be interrupted at any point during execution to return an answer whose value, at least in certain classes of stochastic processes, improves in expectation as a function of the computation time. In this paper, *deliberation scheduling* refers to the problem of allocating processor time to envelope alteration and policy generation. We gather statistics on how envelope alteration and policy generation improve performance and use these statistics to compile expectations for allocating computational resources in time-critical situations.

We consider several decision models for deliberation scheduling. In the simpler models, called *precursor-deliberation* models, we assume that the agent has one opportunity to generate a policy and that, having generated a policy, the agent must use that policy thereafter. In more complicated models, called *recurrent-deliberation* models, we assume that the agent periodically replans and executes the resulting policy in parallel with planning the next policy.

Our approach is motivated by the intuitively appealing work of Drummond and Bresina on ‘anytime synthetic projection’ [Drummond and Bresina, 1990]. In this paper, we reformulate their basic framework in terms of Markov decision processes, cast the algorithmic issues in terms of approximations to specific optimization problems, provide a disciplined approach to allocating computational resources at run time, introduce techniques for specifying goals in stochastic domains, and describe how to extend the framework to deal with uncertainty in observation.

The structure of this paper is as follows. We begin with an introduction to stochastic decision making in Section 2, then define the structures necessary to cope with large state spaces in Section 3. In Section 4 we present basic algorithms for a variety of envelope

alterations and consider how iterative refinement versions of these algorithms may be used as anytime algorithms to provide expected improvements in value. Idealized decision models for precursor and recurrent deliberation are presented in Section 5; experimental results from the domain of robot path planning are given in Section 6. In Section 7 we show how the language of reward functions can be used to specify complex goals, including goals of achievement, maintenance of properties of the world, and prioritized combinations of primitive goals. In Section 8 we outline an extension of our approach to handle uncertainty in observation, based on the theory of partially observable Markov processes. The research presented in this paper is related to work in the literature on sequential decision making, stochastic control, and reinforcement learning; we review this work in Section 9.

2 Markov Decision Models

Following the work on Markov decision processes [Bellman, 1957, Bertsekas, 1987], we model the entire environment as a stochastic automaton. Let $\mathcal{S}$ be the finite set of world states; we assume that they can be reliably identified by the agent. Let $\mathcal{A}$ be the finite set of actions; every action can be taken in every state. The transition model of the environment is a function mapping elements of $\mathcal{S} \times \mathcal{A}$ into discrete probability distributions over $\mathcal{S}$. We write $\Pr(s_1, a, s_2)$ for the probability that the world will make a transition from state s_1 to state s_2 when action a is taken.

A *policy* π is a mapping from $\mathcal{S}$ to $\mathcal{A}$, specifying an action to be taken in each situation. An environment combined with a policy for choosing actions in that environment yields a Markov chain [Kemeny and Snell, 1960].

A *reward function* is a mapping from $\mathcal{S}$ to $\mathbb{R}$, specifying the instantaneous reward that the agent derives from being in each state. Given a policy π and a reward function R , the *value* of state $s \in \mathcal{S}$, $V_\pi(s)$, is the sum of the expected values of the rewards to be received at each future time step, discounted by how far into the future they occur. That is, $V_\pi(s) = \sum_{t=0}^{\infty} \gamma^t E(R_t)$, where R_t is the reward received on the t th step of executing policy π after starting in state s . The *discounting factor*, $0 \leq \gamma < 1$, controls the influence of rewards in the distant future. When $\gamma = 0$, the value of a state is determined entirely by rewards received on the next step; we are generally interested in problems with a longer horizon and set γ to be near 1. Due to properties of the exponential, the definition of V

can be rewritten as

$$V_{\pi}(s) = R(s) + \gamma \sum_{s' \in \mathcal{S}} \Pr(s, \pi(s), s') V_{\pi}(s') \quad . \quad (1)$$

We say that policy π *dominates* (is better than) π' if, for all $s \in \mathcal{S}$, $V_{\pi}(s) \geq V_{\pi'}(s)$, and for at least one $s \in \mathcal{S}$, $V_{\pi}(s) > V_{\pi'}(s)$. A policy is optimal if it is not dominated by any other policy.

One of the most common goals is to achieve a certain condition *p as soon as possible*. If we define the reward function as $R(s) = 0$ if p holds in state s and $R(s) = -1$ otherwise, and represent all goal states as being absorbing, then the optimal policy will result in the agent reaching a state satisfying p as soon as possible. A state is *absorbing* if all actions result in that same state with probability 1; that is, $\forall a \in \mathcal{A}$, $\Pr(s, a, s) = 1$. Making the goal states absorbing ensures that we go to the “nearest” state in which p holds, independent of the states that will follow. The language of reward functions is quite rich, allowing us to specify much more complex goals, including the maintenance of properties of the world and prioritized combinations of primitive goals; this is explored in Section 7.

Given a state-transition model, a reward function, and a value for γ , it is possible to compute the optimal policy using either the policy iteration algorithm [Howard, 1960] or the value iteration algorithm [Bellman, 1957]. We use the policy iteration algorithm because it is guaranteed to converge in a finite number of steps—generally a small number of steps in the domains that we have experimented with—and thus simplifies debugging our computational experiments. The policy iteration algorithm works as follows:

1. Let π' be any policy on $\mathcal{S}$
2. While $\pi \neq \pi'$ do loop
 - a. $\pi := \pi'$
 - b. For all $s \in \mathcal{S}$, calculate $V_{\pi}(s)$ by solving the set of $|\mathcal{S}|$ linear equations in $|\mathcal{S}|$ unknowns given by equation 1
 - c. For all $s \in \mathcal{S}$, if there is some action $a \in \mathcal{A}$ s.t.
 $[R(s) + \gamma \sum_{s' \in \mathcal{S}} \Pr(s, s', a) V_{\pi}(s')] > V_{\pi}(s)$, then $\pi'(s) := a$; otherwise $\pi'(s) := \pi(s)$

3. Return π

The algorithm iterates, generating at every step a policy that strictly dominates the previous policy, and terminates when a policy can no longer be improved, yielding an optimal policy. In every iteration, the values of the states under the current policy are computed. This is done by solving a system of equations, which requires time on the order of $|\mathcal{S}|^3$. The algorithm then improves the policy by looking for states s in which doing some action a other than $\pi(s)$ for one step, then continuing with π , would result in higher expected reward than simply executing π . When such a state is found, the policy is changed so that it always chooses action a in that state. The algorithm is guaranteed to converge in a number of iterations polynomial in $|\mathcal{S}|$.

3 Coping with Large State Spaces

As the size of our state spaces grows, even a polynomial-time algorithm such as policy iteration becomes too inefficient. We will assume that our environment is such that, for any given reward function and initial starting state, it is sufficient to consider a highly-restricted subset of the entire state space in our planning.

A *partial policy* is a mapping from a subset of $\mathcal{S}$ into actions; the domain of a partial policy π is called its *envelope*, $\mathcal{E}_\pi$. The *fringe* of a partial policy, $\mathcal{F}_\pi$, is the set of states that are not in the envelope of the policy, but that may be reached in one step of policy execution from some state in the envelope. That is, $\mathcal{F}_\pi = \{s \in \mathcal{S} - \mathcal{E}_\pi \mid \exists s' \in \mathcal{E}_\pi \text{ s.t. } \Pr(s', \pi(s'), s) > 0\}$.

To construct a restricted automaton, we take an envelope $\mathcal{E}$ of states and add the distinguished state `OUT`. For any states s and s' in $\mathcal{E}$ and action a in $\mathcal{A}$, the transition probabilities remain the same. Further, for every $s \in \mathcal{E}$ and $a \in \mathcal{A}$, we define the probability of going out of the envelope as

$$\Pr(s, a, \text{OUT}) = 1 - \sum_{s' \in \mathcal{E}} \Pr(s, a, s') \ .$$

The `OUT` state is absorbing.

The cost of falling out of the envelope is a parameter that depends on the domain. If it is possible to re-invoke the planner when the agent falls out of the envelope, then one

approach is to assign $V(\text{OUT})$ to be the estimated value of the state into which the agent fell minus some function of the time required to construct a new partial policy. Under the reward function described earlier, the value of a state is negative, and its magnitude is the expected number of steps to the goal; if time spent planning is to be penalized, it can simply be added to the magnitude of the value of the OUT state with a suitable weighting function.

4 Basic Algorithms

As a concession to complexity, in generating a policy, our algorithms consider only a subset of the state space of the stochastic process. The algorithm starts with an initial policy and a restricted state space (or *envelope*), extends that envelope, and then computes a new policy. We would like it to be the case that the new policy π' is an improvement over (or at the very least no worse than) the old policy π in the sense that $V_{\pi'}(s_0) - V_{\pi}(s_0) \geq 0$.

In general, however, we cannot guarantee that the policy will improve without extending the state space to be the entire space of the system automaton, which results in computational problems. The best that we can hope for is that the algorithm improves in *expectation*. Suppose that the initial envelope is just the initial state and the initial policy is determined entirely by the reflexes. The difference $V_{\pi'}(s_0) - V_{\pi}(s_0)$ is a random variable, where π is the reflex policy and π' is the computed policy. We would like it to be the case that $E[V_{\pi'}(s_0) - V_{\pi}(s_0)] > 0$, where the expectation is taken over start states and goals drawn from some fixed distribution. Although it is possible to construct system automata for which even this improvement in expectation is impossible, we believe many moderately benign environments are well-behaved in this respect. In particular, *navigation* environments (excluding mazes) in which transitions are restricted by spatio-temporal constraints generally satisfy our requirements.

Our basic algorithm consists of two stages: envelope alteration followed by policy generation. The algorithm takes an envelope and a policy as input and generates as output a new envelope and policy. We assume that the algorithm has access to the state transition matrix for the stochastic process. In general, we assume that the algorithm is applied in the manner of iterative refinement, with more than one invocation of the algorithm. We also treat envelope alteration and policy generation as separate, so we cast the overall process of policy formation in terms of some number of rounds of envelope alteration followed

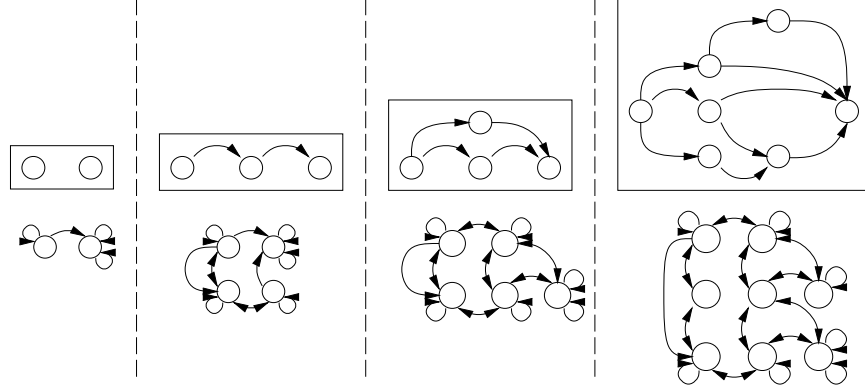


Figure 1: Sequence of restricted automata and associated paths through state space

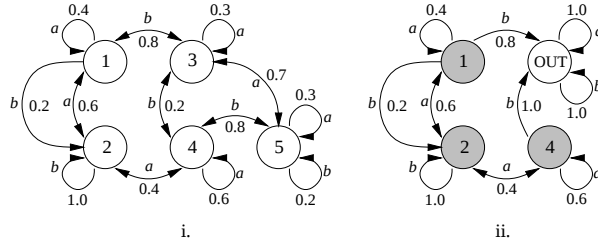


Figure 2: Stochastic process and a restricted version

by policy generation, resulting in a sequence of policies. Figure 1 depicts a sequence of automata generated by iterative refinement along with corresponding paths through state space from the initial state to a goal state.

Policy generation is itself an iterative algorithm that improves an initial policy by estimating the value of policies with respect to the restricted-state-space stochastic process mentioned earlier. When run to completion, policy generation continues to iterate until it finds a policy that it cannot improve with respect to its estimate of value; this policy is guaranteed to be optimal with respect to the restricted-state-space stochastic process. We say that policy generation is *inflexible* if the i th round of policy generation is always *run to completion* on the restricted automaton available at the i th round.

Figure 2.i shows an example system automaton consisting of five states. Suppose that the initial state is 1, and state 4 satisfies the goal. The path $1 \xrightarrow{a} 2 \xrightarrow{a} 4$ goes from the initial

state to a state satisfying the goal and the corresponding envelope is $\{1, 2, 4\}$. Figure 2.ii shows the restricted automaton for that envelope. Let $\pi(s)$ be the action specified by the policy π to be taken in state s ; the optimal policy for the restricted automaton shown in Figure 2.ii is defined by $\pi(1) = \pi(2) = \pi(4) = a$ on the states of the envelope and the reflexes by $\pi(\text{OUT}) = b$ (*i.e.*, $\forall s \notin \{1, 2, 4\}, \pi(s) = b$) (the reflex actions need not be the same in all states).

First we consider the high level algorithms for a precursor and a recurrent model of planning and execution. Execution of an explicit policy is trivial, so we describe only the algorithm for generating policies. We then look at the various component phases of the planning algorithms, which are shared between both models. A wider range of precursor and recurrent models are described in Section 5, where we consider how to schedule deliberation in both models.

4.1 High Level Algorithms

4.1.1 Precursor Deliberation Model

In the precursor deliberation model, there are two separate phases of operation: planning and execution. The planner constructs a policy that is followed by the agent until a new goal must be pursued or until the agent falls out of the current envelope. In the simplest precursor models, a deadline is specified indicating when planning stops and execution begins.

The high level planning algorithm, given a description of the environment and start state s_0 or a distribution over start states, is as follows:

1. Generate an initial envelope $\mathcal{E}$
2. While $(\mathcal{E} \neq \mathcal{S})$ and (not deadline) do
 - a. Extend the envelope $\mathcal{E}$
 - b. Generate an optimal policy π for restricted automaton with state set $\mathcal{E} \cup \{\text{OUT}\}$
3. Return π

The algorithm first finds a small subset of world states and calculates an optimal policy over those states. Then it gradually adds new states in order to make the policy robust by decreasing the chance of falling out of the envelope. After new states are added, the optimal policy over the new envelope is calculated. Note the interdependence of these steps: the choice of which states to add during envelope extension may depend on the current policy, and the policy generated as a result of optimization may be quite different depending on which states were added to the envelope. The algorithm terminates when a deadline has been reached or when the envelope has been expanded to include the entire state space.

4.1.2 Recurrent Deliberation Model

A more sophisticated model of interaction between planning and execution is one in which the planner runs concurrently with the execution, sending new or expanded strategies to the executor as they are developed.

In recurrent-deliberation models, the agent has to repeatedly decide how to allocate time to deliberation, taking into account new information obtained during execution. The details of such models are discussed in Section 5; here we provide just a rough sketch. We assume two separate modules: one for planning and a second for execution. In the simplest model, the planner and executor operate in a rigid cycle with a period determined by fixed length of time. At the beginning of each cycle, the planner is given the current state by the execution module; the planner spends the fixed length of time working on a new policy; at the end of the fixed time, the planner gives the new policy to the execution module. For the time being we assume that the execution module can identify the current state with certainty; in Section 8 we consider the case in which there is uncertainty in observation.

In the recurrent models, it is often necessary to remove states from the envelope in order to lower the computational costs of generating policies from the restricted automata. For instance, in the mobile-robot domain, it may be appropriate to remove states corresponding to portions of a path the robot has already traversed if there is little chance of returning to those states. Figure 3 shows a typical sequence of changes to the envelope corresponding to the state space for the restricted automaton. The current state is indicated by $\blacklozenge$ and the goal state is indicated by $\square$.

The recurrent planning algorithm, given a description of the environment, the policy π_c that is currently being followed by the agent, and the state of the agent at the beginning

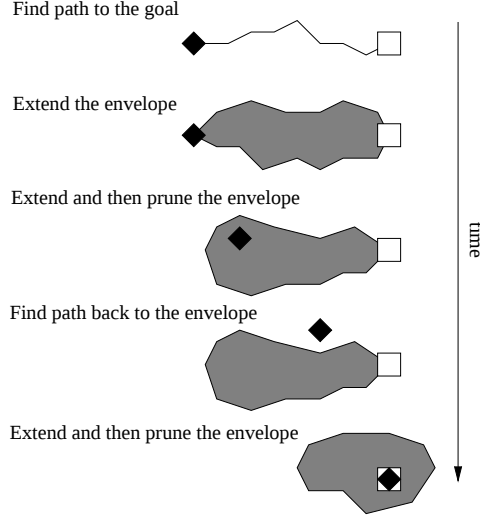


Figure 3: Typical sequence of changes to the envelope

of the planning interval, s_c , is as follows:

While (not goal) do

1. Set s_c to be the current state for planning purposes
2. While (not end of current planning interval) do
 - a. Extend the envelope $\mathcal{E}$
 - b. Prune the envelope $\mathcal{E}$
 - c. Generate an optimal policy π' for restricted automaton with state set $\mathcal{E} \cup \{\text{OUT}\}$
3. Set π_c to be the new policy π'

The details of the extension and pruning of the envelope will depend on the agent's expected state at the end of the planning interval. This can be determined from the state transitions, s_c and π_c by forward simulation.

4.2 Algorithm Components

In the following sections, we consider each subcomponent of the precursor and recurrent algorithms in more detail. Each of the subcomponents can be implemented as an anytime

algorithm; in Section 5, we cast the problem of allocating computational resources to the subcomponents as an optimization problem and then describe decision-theoretic techniques to compute approximations.

4.2.1 Policy Generation

Given a restricted automaton with envelope $\mathcal{E}$, we use the policy iteration algorithm to generate the optimal policy. Although this is potentially an $O(|\mathcal{E}|^3)$ operation, most realistic environments cannot transition from every state to every other, so the transition matrix is sparse, allowing much more efficient solution of the equations. This algorithm requires a number of iterations at most polynomial in the number of states; in practice, in our robot path planning domains with 660 to 16,000 world states, it has never taken more than 15 to 50 iterations respectively. When we use this as a subroutine in our planning algorithm, we generate a plausible policy for the first step, and then for all subsequent steps we use the old policy as the starting point for policy iteration. We are considering, but have not tried, Kushmerick *et al.*'s method for generating plausible initial policies [Kushmerick *et al.*, 1993]. Because, in general, the policy does not change radically when the envelope is extended, it requires very few iterations of the policy iteration algorithm to generate the optimal policy for the extended envelope (typically 2 or 3 iterations for the smaller domains, up to 10 for the larger domains). Occasionally, when a very dire consequence or an exceptional new path is discovered, the whole policy must be changed.

4.2.2 Initial Trajectory Planning

The high-level precursor and recurrent algorithms work no matter how the initial envelope is chosen, but if it is done with some intelligence, the early policies are much more useful, and the time taken to reach the goal is shorter. In our examples, we consider the goal of being in a state satisfying p as soon as possible. For such simple goals of achievement, a good initial envelope is one containing a chain of states from the initial state, s_0 , to some state satisfying p such that, for each state, there is some action with a non-zero probability of moving to the next state in the chain.

In the implemented system, we generate a path from start to goal by doing a depth-first search from s_0 considering the most probable outcome for each action in decreasing order of

probability. This yields a set of states that can be traversed to a goal state. We then check for any shortcuts within this path, deleting intermediate states if this does not decrease the probability of reaching the goal for the resultant path. We then attempt to improve the robustness of the path by adding a successor to a path state if it in turn has the next state in the path as its successor and the combined transition probabilities are higher than the original single transition probability.

We use this method to generate a small number of paths from the start to the goal, say 10, and choose the shortest path (which is usually the path with the highest probability) to form the initial envelope to the policy. We can then use the nominal path from the start to goal as the initial policy, which makes the optimization of the initial envelope much faster than if we began with a completely random policy for the envelope. More sophisticated techniques or more complicated heuristics could be used to generate a good initial envelope; our strategy is to spend as little time as possible doing this, so that a plausible policy is available as soon as possible.

4.2.3 Envelope Alteration

Envelope alteration can be classified in terms of three basic operations on the envelope: *trajectory planning*, *envelope extension*, and *envelope pruning*. Trajectory planning consists of searching for a path from some initial state to a state satisfying the goal; this method need not make use of the current restricted automaton. Envelope extension adds states to the envelope. Envelope pruning removes states from the envelope and is generally used only in recurrent-deliberation models. Both envelope extension and envelope pruning will typically make use of the current restricted automaton; for example, envelope extension may add those the states outside of the envelope that the agent is most likely to reach given the current policy, and pruning may delete states that the agent is unlikely to end up in.

Extending the envelope There are a number of possible strategies for extending the envelope; the most appropriate depends on the domain. The aim of the envelope extension is to judiciously broaden the subset of the world states, by including states that are outside the envelope of the current policy but that may be reached upon executing the policy. One simple strategy is to add the entire fringe of the current policy, $\mathcal{F}_\pi$; this would result in

adding states uniformly around the current envelope. It will often be the case, however, that some of the states in the fringe are very unlikely to be reached given the current policy.

A more reasonable strategy, similar to one advocated by Drummond and Bresina [Drummond and Bresina, 1990], is to look for the N most likely fringe states. We do this by simulating the restricted automaton and accumulating the probabilities of falling out into each fringe state. We then have a choice of strategies. We can add each of the N most likely fringe states. Alternatively, for goals of achievement, we can take each element of this subset of the fringe states and find a path from the state back to some state in the envelope. Following Drummond and Bresina, we call this class of envelope extension methods *robustification*. Robustification may also be combined with trajectory planning by taking a fringe state and adding a path to a state that satisfies the goal.

Trajectory planning Trajectory planning is performed in much the same way as initial trajectory planning except that the notions of initial and goal states may vary. For instance, we often wish to find a path (trajectory) from some fringe state back to some state inside the envelope or back to a state satisfying the goal. Apart from this difference, the actual search techniques are exactly the same as in initial trajectory planning.

Pruning In the recurrent models, it is often necessary to remove states from the envelope in order to lower the computational costs of generating policies from the restricted automaton. One obvious method is to prune states from the current envelope on the grounds that the agent is unlikely to end up in those states and therefore need not consider them in formulating a policy. However we need to be careful about how this is done. It may be that a state has a low instantaneous reward, or is some kind of sink (*e.g.*, all non-self transitions have low probability). In these situations the current policy will direct the agent away from that state, resulting in a low probability of that state being reached. We do not always want to remove this kind of state; its presence in the envelope directs the agent away from an area it should avoid. We want to distinguish between states that have a low probability of being reached because they are in an area to be avoided but are still somehow between the agent and the goal, and those which the agent has gone past. In order to prune the latter category, for the results given in Section 6, we prune the N least likely states which

also have a lower value than the value of the current state.¹

4.3 Example: Robustification in the Precursor Model

In our approach, unlike that of Drummond and Bresina, extending the current policy is coupled tightly and naturally to *changing* the policy as required to keep it optimal with respect to the restricted view of the world. The following example illustrates how such changes are made using the precursor algorithm described above.

The example domain is high-level mobile-robot path planning. It was chosen so that it would be easy to understand the policies generated by our algorithms. The floor plan is divided into a grid of 166 locations, $\mathcal{L}$, with four directional states associated with each location, $\mathcal{D} = \{N, S, E, W\}$, corresponding to the direction the robot is facing, resulting in a total of 664 world states, representing the layout of the fourth floor of the Brown University Computer Science department (see Figure 4). The robot is given a task to navigate from some starting location to some target location. The actions available to the robot are $\{\text{STAY}, \text{GO}, \text{TURN-RIGHT}, \text{TURN-LEFT}, \text{TURN-ABOUT}\}$. The transition probabilities for the outcome of each action may be obtained empirically. In our experimental simulation, the STAY action is guaranteed to succeed. The probability of success for GO and turning actions in most locations is 0.8, with the remainder of the probability mass divided between undesired results such as overshooting, over-rotating, slipping sideways, etc. The world also contains *sinks*, locations that are difficult or impossible to leave. In the mobile-robot domain, a sink might correspond to a stairwell that the robot could fall into. The reward function for the sequential decision problem associated with a given initial and target location assigns 0 to the four states corresponding to the target location and -1 to all other states. On average each state has 15.6 possible successors. This is a dispersion rate of about 2.3% for the basic fourth floor domain, however it is reduced 0.1% when the domain is expanded to 16,000 states; this is the type of low dispersion rate domain we identified earlier as suitable for our approach.

Figure 5 shows a subset of the domain corresponding to the locations surrounding a stairwell. The stairwell states taken as a whole correspond to what we call a *complete sink*; there are no nonzero transitions out of a complete sink. The stairwell states are only

¹Note that sinks that are between the current state and the goal will have low value and still be deleted, so this heuristic is still not ideal.

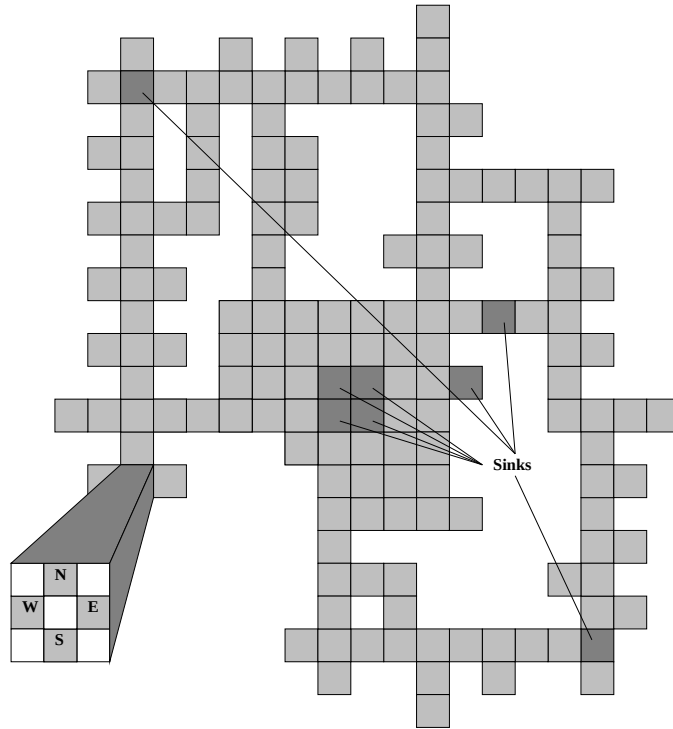


Figure 4: Plan of the Brown CS department fourth floor used for the mobile-robot domain. Each square is a location corresponding to four states, one for each heading of the robot, as shown in the expanded square in the lower left. Sinks are indicated by darker shading. The larger domains are constructed by replicating this map in a grid layout.

accessible from one direction, the north. In this figure there are four small squares associated with each location, one for each possible heading; thus each small square corresponds to a state, the direction of the arrow shows the policy for the robot in that location and with that heading. Figure 5 (a) shows the optimal policy for a small early envelope; Figures 5(b) and (c) show two subsequent envelopes where the policy changes to direct the robot to circumvent the stairwell, reflecting aversion to the risk involved in taking the shortest path.

5 Deliberation Scheduling

Deliberation scheduling is the problem of allocating processor time to envelope alteration and policy generation. It is natural to think of deliberation scheduling in terms of optimization even if the combinatorics dictate that an optimal solution is not computationally feasible. Having said this, it still remains to determine what optimization problem we are trying to solve. We have to specify exactly what options are allowed and what information is available; such a characterization is generally referred to as a *decision model*.

In the following, we present a number of decision models. It should be pointed out that for each instance of the problems that we consider there are a large number of possible decision models. By specifying different decision models, we can make deliberation scheduling easy or hard. Our selection of which decision models to investigate is guided by our interest in providing insight into the problems of time-critical decision making and our anticipation of the combinatorial problems involved in deliberation scheduling. In this section, we ignore the time spent in deliberation scheduling; for practical reasons, however, we are interested in decision models for which the on-line time spent in deliberation scheduling is negligible.

In the simpler *precursor-deliberation* models, we assume that the agent has one opportunity to generate a policy and that having generated a policy the agent has to stick to that policy thereafter. Precursor-deliberation models include those in which

1. a deadline is given in advance specifying when to stop deliberating and start acting according to the generated policy (the algorithm for this was given in Section 4.1.1)
2. the agent is given an unlimited amount of time to respond with a time cost of delay specified as a fixed function

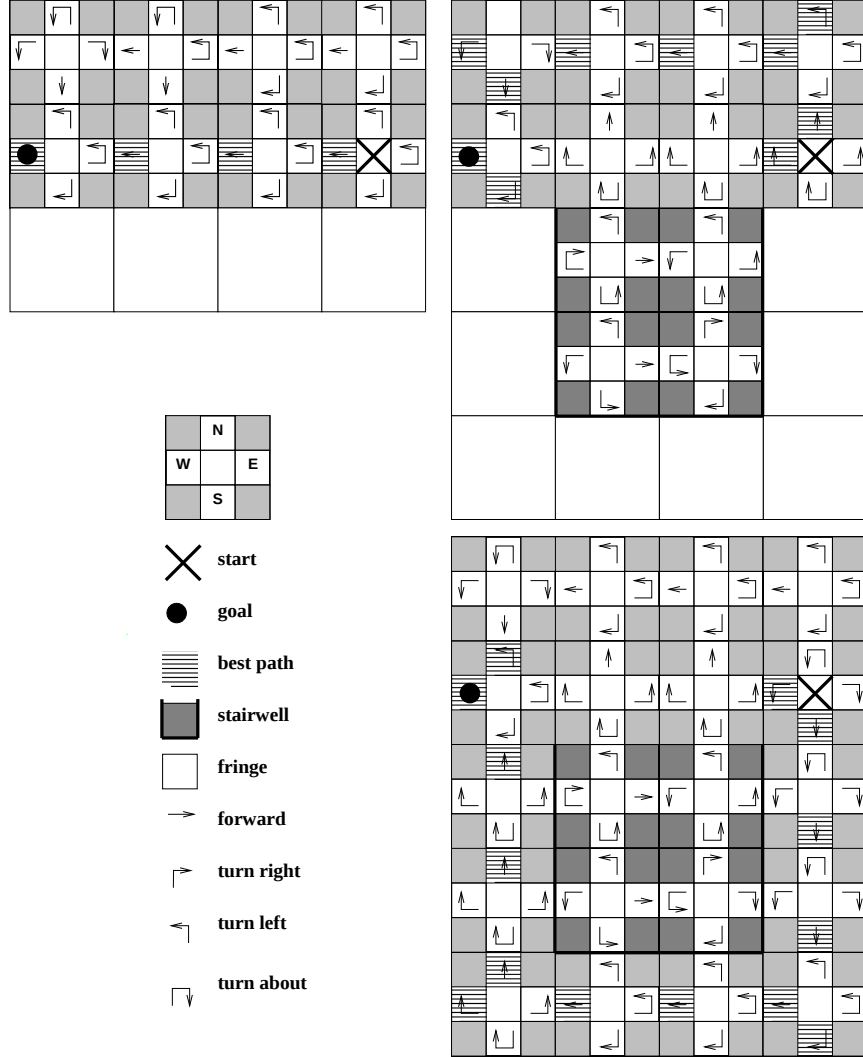


Figure 5: Example of policy change for different envelopes near a complete sink. The direction of the arrow indicates the current policy for that state. (a) Sink not in the envelope: the policy chooses the straightforward shortest path. (b) Sink included: the policy skirts north around it. (c) All states surrounding the stairwell included: the barriers on the south, east and west sides allow the policy take a longer but safer path. For this run $\gamma = 0.999999$ and $V(\text{OUT}) = -4000$.

In the following two models, there is a trigger event that occurs indicating that the agent must begin following its policy immediately with no further refinement.

3. the trigger event can occur at any time in a fixed interval with a uniform distribution
4. the trigger event is governed by a more complicated distribution, *e.g.*, a normal distribution centered on an expected time

Models (3) and (4) are not considered in this paper; models (1) and (2) are treated in Section 5.1.

In the more complicated *recurrent-deliberation* models, we assume that the agent periodically replans. Recurrent-deliberation models include those in which

5. the agent performs further envelope alteration and policy generation if and only if it ‘falls out’ of the envelope defined by the current restricted automaton
6. the agent performs further envelope alteration and policy generation in parallel with execution, tailoring the restricted automaton and its corresponding policy to states anticipated in the near future (the algorithm for this was given in Section 4.1.2)

Model (6) is treated in Section 5.2.

5.1 Precursor-Deliberation Models

In the following we consider four cases of precursor deliberation with known deadlines and one case of precursor deliberation with unlimited time to respond and a cost for delay. Let t_{TOT} be the total amount of time from the current time until the deadline. If there are k rounds of envelope alteration and policy generation, then we have

$$t_{EA_1} + t_{PG_1} + \cdots + t_{EA_k} + t_{PG_k} = t_{TOT},$$

where t_{EA_i} (t_{PG_i}) is the time spent in the i th round of envelope alteration (policy generation). Let π_i represent the policy after the i th round of envelope alteration followed by policy generation.

Single round; inflexible policy generation; deadlines In the simplest case, policy generation does not inform envelope alteration and so we might as well do all of the envelope alteration before policy generation, and

$$t_{EA_1} + t_{PG_1} = t_{TOT}.$$

In order to schedule time for EA_1 and PG_1 , we need:

1. the expected improvement of the value of a random initial state between the reflex policy and the policy resulting from policy generation given a fixed amount of time allocated to envelope alteration, $E[V_{\pi_1}(s_0) - V_{\pi_0}(s_0)|t_{EA_1}]$;
2. the expected size of the envelope given the time allocated to the first round of envelope alteration, $E[|\mathcal{E}_1| | t_{EA_1}]$; and
3. the expected time required for policy generation given the size of the envelope after the first round of envelope alteration, $E[t_{PG_1} | |\mathcal{E}_1|]$.

Each of (1), (2) and (3) can be determined empirically, and, at least in principle, the optimal allocations to envelope alteration and policy generation can be determined. If we assume no variance in run times and envelope sizes, then optimal deliberation scheduling corresponds to finding that t_{EA_1} maximizing $E[V_{\pi_1}(s_0) - V_{\pi_0}(s_0)|t_{EA_1}]$ subject to the constraint that

$$t_{EA_1} + E[t_{PG_1} | E[|\mathcal{E}_1| | t_{EA_1}]] \leq t_{TOT}.$$

Note that, because policy generation is itself an iterative refinement algorithm, we can interrupt it at any point to obtain a policy. Although the particular model considered here assumes inflexible policy generation for the purpose of deliberation scheduling, we might use a more flexible approach to handling of deadlines at run time.

In the case of nondegenerate distributions over run times and envelope sizes, optimal deliberation scheduling would require consideration of cases in which the actual run times violate the specified deadline. This is relatively straightforward to model, but considerably more difficult to implement.

Multiple rounds; inflexible policy generation; deadlines Assume that policy generation can profitably inform envelope alteration, *i.e.*, that the policy after round i provides guidance in extending the environment during round $i + 1$. In this case, we have k rounds and

$$t_{EA_1} + t_{PG_1} + \cdots + t_{EA_k} + t_{PG_k} = t_{TOT}.$$

Recall that the fringe states for a given envelope and policy correspond to those states outside the envelope that can be reached with a non-zero probability in a single step by following the policy starting from some state within the envelope. Let the most likely *falling-out* state with respect to a given envelope and policy correspond to that fringe state that is most likely to be the first fringe state reached by following the policy starting in the initial state. We might consider a very simple method of envelope alteration in which we just add the most likely falling-out state and then the next most likely and so on. Suppose that adding each additional state takes a fixed amount of time. Let

$$E[V_{\pi_i}(s_0) - V_{\pi_{i-1}}(s_0) | |\mathcal{E}_{i-1}| = m, |\mathcal{E}_i| = m + n]$$

denote the expected improvement in the value of the initial state after the i th round of envelope alteration and policy generation given that there are n states added to the m states that were already in the envelope after the $i - 1$ round.

Again, the expectations described above can be obtained empirically. Coupled with the sort of expectations described for the previous single-round case (*e.g.*, $E[t_{PG_i} | |\mathcal{E}_i|]$), one could, at least in principle, determine the optimal number of rounds k and the allocations to t_{EA_i} and t_{PG_i} for $1 \leq j \leq k$. In practice, we use slightly different statistics and heuristic methods for deliberation scheduling to avoid the combinatorics.

Single round; flexible policy generation; deadlines Actually, this case is simpler in concept than the case with inflexible policy generation assuming that we can compile the following statistics.

$$E[V_{\pi_1}(s_0) - V_{\pi_0}(s_0) | t_{EA_1}, t_{PG_1}]$$

Multiple rounds; flexible policy generation; deadlines Again, with additional statistics, *e.g.*,

$$E[V_{\pi_i}(s_0) - V_{\pi_{i-1}}(s_0) | |\mathcal{E}_{i-1}| = m, |\mathcal{E}_i| = m + n, t_{PG_{i-1}}],$$

this case is not much more difficult than the earlier cases.

Single round; inflexible policy generation; cost of delay Deliberation models that assume no fixed deadline but specify a time cost of delay as a fixed function can be handled similarly to the cases considered above. For instance, in the case of single round, inflexible policy generation, if we assume no variance in run times and envelope sizes, optimal deliberation scheduling corresponds to finding that t_{EA_1} maximizing the sum of $E[V_{\pi_1}(s_0) - V_{\pi_0}(s_0)|t_{EA_1}]$ and $\text{Cost}(t_{EA_1} + E[t_{PG_1} | E[|\mathcal{E}_1| | t_{EA_1}]])$.

5.2 Recurrent-Deliberation Models

In recurrent deliberation models, the agent has to decide repeatedly how to allocate time to deliberation, taking into account new information obtained during execution. In this section, we consider a particular model for recurrent deliberation in which the agent allocates time to deliberation only at prescribed intervals. We assume that the agent has separate planning and execution modules that run in parallel and communicate by message passing; the planning module sends policies to the execution module and the execution module sends observed states to the planning module.

We call the models considered in this section the *discrete, weakly-coupled, recurrent deliberation* models. *Discrete* because each tick of the clock corresponds to exactly one state transition; *recurrent* because the execution module gets a new policy from the planning module periodically; *weakly coupled* in that the two modules communicate by having the execution module send the planning module the current state and the planning module send the execution module the latest policy.

As mentioned earlier, in the recurrent models, it is often necessary to remove states from the envelope in order to lower the computational costs of generating policies from the restricted automata. In general, there are many more possible strategies for deploying envelope alteration and policy generation in recurrent models than in the case of precursor models. To cope with the attendant combinatorics, we raise the level of abstraction slightly and assume that we are given a small set of strategies that have been determined empirically to improve policies significantly in various circumstances. Each strategy corresponds to some fixed schedule for allocating processor time to envelope alteration and policy generation routines.

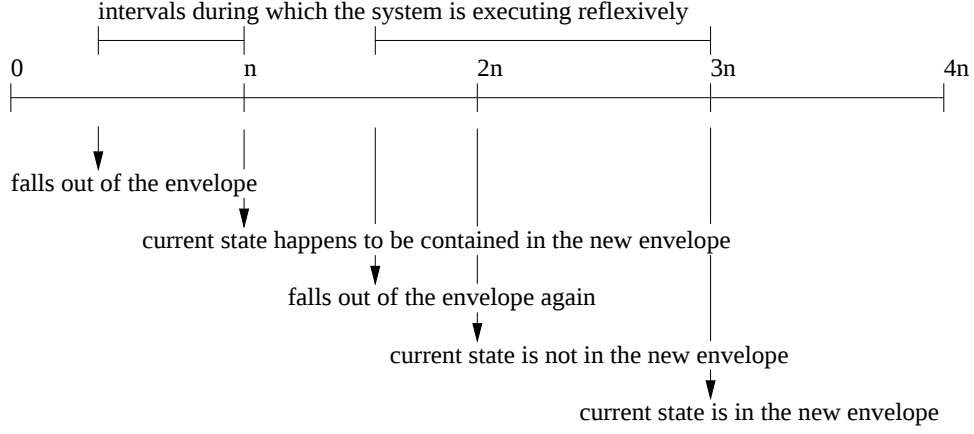


Figure 6: Recurrent deliberation

Discrete; weakly-coupled; fixed intervals We first consider the case in which communication between the two modules occurs exactly once every n execution steps or *ticks*; at times $n, 2n, 3n, \dots$, the planning module sends off the policy generated in the last n ticks, receives the current state from the execution module, and begins deliberating on the next policy. Strategies would be tuned to a particular n -tick planning cycle. One strategy might be to use a particular pruning algorithm to remove a specified number of states and then use whatever remains of the n ticks to generate a new policy. In this regime, deliberation scheduling consists of choosing which strategy to use at the beginning of each n -tick interval.

Before we get into the details of the decision model, consider some complications that arise in recurrent deliberation problems. At any given moment, the agent is executing a policy, π , defined on the current envelope and augmented with a set of reflexes for states falling outside the envelope. The agent begins executing π in state s . At the end of the current n -tick interval, the execution module is given a new policy π' , and the planning module is given the current state s' . It is possible that s' is not included in the envelope for π' ; if the reflexes do not drive the robot inside the envelope then the agent's behavior throughout the next n -tick interval will be determined entirely by the reflexes. Figure 6 shows a possible run depicting intervals in which the system is executing reflexively and intervals in which it is using the current policy; for this example, we assume reflexes that enable an agent to remain in the same state indefinitely.

Let $\delta_n(s, \pi, s')$ be the probability of ending up in s' starting from s and following π for n steps. Suppose that we are given a set of deliberation strategies $\{F_1, F_2, \dots\}$. As is usual in such combinatorial problems with indefinite horizons, we adopt a myopic decision model with a limited horizon. In particular, we assume that, at the beginning of each n -tick interval, we are planning to follow the current policy π for n steps, follow the policy $F(\pi)$ generated by some strategy F attempting to improve on π for the next n steps, and thereafter follow the optimal policy π^* . If we assume that it is impossible to get to a goal state in the next $2n$ steps, the expected value of using strategy F is given by

$$\left[- \sum_{i=0}^{2n-1} \gamma^i + \gamma^{2n} \sum_{s' \in S} \delta_n(s, \pi, s') \left[\sum_{s'' \in S} \delta_n(s', F(\pi), s'') V_{\pi^*}(s'') \right] \right] - V_{\pi}(s),$$

where $0 \leq \gamma < 1$ is a discounting factor, controlling the degree of influence of future results on the current decision.

Extending the above model to account for the possibility of getting to the goal state in the next $2n$ steps is straightforward; computing a good estimate of V_{π^*} is not, however. We might use the value of some policy other than π^* , but then we run the risk of choosing strategies that are optimized to support a particular suboptimal policy when in fact the agent may be able to do much better. In general, it is difficult to estimate the long term prospects for sequential decision problems of indefinite duration. In the next model, we consider an alternative decision model that avoids computing or even estimating the value of the optimal policy, but has related problems in practice.

Discrete; weakly-coupled; variable intervals One practical problem with the fixed-interval model is that it is difficult to design strategies for a fixed n -tick interval. In this case, we allow variable planning intervals and assume that we can predict reasonably accurately the time required for a given deliberation strategy to run. Also, in anticipation of combinatorial issues that arise in our experimental studies, we adopt a simpler myopic decision model. In this case, we assume that the agent will apply exactly one deliberation strategy and commit to the resulting policy thereafter. The expected value of using strategy F on π assuming that F will take k steps is just

$$\left[- \sum_{i=0}^{k-1} \gamma^i + \gamma^k \sum_{s' \in S} \delta_k(s, \pi, s') V_{F(\pi)}(s') \right] - V_{\pi}(s),$$

where the first term corresponds to the value of using π for the first k steps and $F(\pi)$ thereafter and the second term corresponds to the case in which we do no deliberation whatsoever and use π forever. As in the model described in the previous section, we assume that the goal cannot be reached in the next k steps; again it is simple to extend the analysis to the case in which the goal may be reached in fewer than k steps.

The above decision model does not require that we compute the value of the optimal policy. The model does, however, require that we compute the long-term performance of policies. In practice, of course, we will only compute an estimate, but this estimation will turn out to be rather difficult.

6 Experimental Results

This section reports on experiments conducted in the simulated robot-navigation environment described in Section 4.3.

6.1 Greedy Precursor Deliberation

In general, computing the optimal deliberation schedule for the multiple-round precursor-deliberation models described above is computationally complex. We have experimented with a number of simple, greedy and myopic scheduling strategies; we report on one such strategy here.

We gathered a variety of statistics on how extending the envelope increases value. The statistics that proved most useful corresponded to the expected improvement in value for different numbers of states added to the envelope. Instead of conditioning just on the size of the envelope prior to alteration we found it necessary to condition on both the size of the envelope and the estimated value of the current policy (*i.e.*, the value of the optimal policy computed by policy iteration on the restricted automaton). At run time, we use the size of the automaton and the estimated value of the current policy to index into a table of *performance profiles* giving expected improvement as a function of number of states added to the envelope.

Using the mobile-robot domain (the single fourth floor, 664 states), we generated 1,600,000 data points to compute statistics of the sort described above plus estimates of the

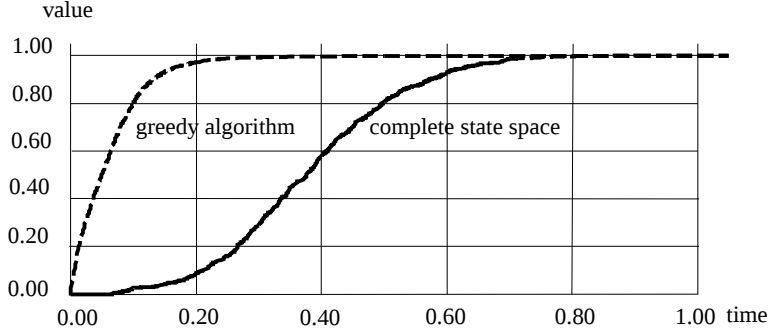


Figure 7: Comparison of planning algorithm using greedy deliberation strategy (dashed line) with the policy iteration optimization method (solid line): Average over 630 runs

time required for one round of envelope alteration followed by policy generation given the size of the envelope, the number of states added, and value of the current policy. We use the following simple greedy strategy for choosing the number of states to add to the envelope on each round. For each round of envelope alteration followed by policy generation, we use the statistics to determine the number of states which, added to the envelope, maximizes the ratio of performance improvement to the time required for computation.

We compared the performance of (1) our planning algorithm using the greedy deliberation strategy with (2) policy iteration optimizing the policy for the whole domain. Our results show that the planning algorithm using the greedy deliberation strategy supplies a good policy early, and typically converges to a policy that is close to optimal before the whole domain policy iteration method does. Figure 7 shows average results from 620 runs, where a single run involves a particular start state and goal state. The graph shows the average value of the start state under the policy available at time t , $V_{\hat{\pi}}(s_0)$, as a function of time. In order to compare results from different start/goal runs, we show the average ratio of the value of the current policy to the value of the optimal policy for the whole domain, plotted against the ratio of actual time to the time, T_{opt} , that the policy iteration takes to reach that optimal value.

The greedy deliberation strategy performs significantly better than the standard optimization method. We also considered very simple strategies such as adding a small fixed number of fringe states each iteration, and adding the whole fringe each iteration, which performed fairly well for this domain, but not as well as the greedy policy. Further experi-

mentation is required to draw definitive conclusions about the comparative performance of these deliberation strategies for particular domains.

6.2 Recurrent Deliberation

In this section, we present results for recurrent-deliberation problems of indefinite duration using statistical estimates of the value of a variety of deliberation strategies. We do this for the discrete, weakly-coupled decision model which allows variable-length intervals for deliberation. Although fixed-length intervals facilitate exposition, it is much easier to collect useful statistical estimates of the utility of deliberation strategies if the deliberation interval is allowed to vary. For the remainder of this section, a deliberation strategy is just a particular sequence of invocations of envelope alteration and policy generation routines.

6.2.1 Gathering Statistics

The utility of a deliberation strategy is characterized as a function of attributes of the policy to which it will be applied, such as the estimated value of the policy and the size of the envelope. For example, the function $EIV(F, \hat{V}_\pi, |\mathcal{E}_\pi|)$ provides an estimate of the expected improvement in value from using the strategy F assuming that the estimated value of the current policy and the size of the corresponding envelope fall within specified ranges. This function is implemented as a table in which each entry is indexed by a strategy F and a set of ranges over the attributes. We determine the EIV function off line by gathering statistics for F running on a wide variety of policies. At run time, the deliberation scheduler computes an estimate of the value of the current policy $\hat{V}_\pi$, determines the relevant attributes of current policy, for example the size $|\mathcal{E}_\pi|$ of the corresponding envelope, and chooses the strategy F maximizing EIV for those attributes. Note that actual results given in Section 6.2.2 use EIV contingent on other attributes also.

To build a table of estimates of function EIV off line, we begin by gathering data on the performance of strategies ranging over possible initial states, goals, and policies. For a particular strategy F , initial state x , and policy π , we run F on π , determine the elapsed number of steps k , and compute the estimated improvement in value as defined in the section describing the discrete, weakly-coupled, variable interval deliberation model. Given

data of the sort described above, we build the table for $EIV(F, \hat{V}_\pi, |\mathcal{E}_\pi|)$ by appropriately dividing the data into buckets with equal numbers of elements.

One unresolved problem with this approach is exactly how to compute $\hat{V}_\pi(x)$. Recall that π is only a partial policy defined on a subset of $\mathcal{S}$ augmented with a set of reflexes to handle states outside the current envelope. In estimating the value of a policy, we are really interested in estimating the value of the augmented partial policy. If the reflexes kept the agent in the same place indefinitely, then as long as there was some nonzero probability of falling out of the envelope with a given policy starting in a given state, the actual value of the policy in that state would be $-1/(1 - \gamma)$. Of course, this is an extremely pessimistic estimate for the long term value of a particular policy since in the recurrent model the agent will periodically compute a new policy based on where it is in the state space. The problem is that we cannot directly account for these subsequent policies without extending the horizon of the myopic decision model and absorbing the associated computational costs in off-line data gathering and online deliberation scheduling.

To avoid complicating the online decision making, we have adopted the following expedient, which allows us to keep our one-step-lookahead model. We modify the transition probabilities for the restricted automaton so that there is always a non-zero probability of getting back into the envelope after having fallen out of it. Exactly what this probability should be is difficult to determine. The particular value chosen will determine just how concerned the agent will be with the prospect of falling out of the envelope. In fact, the value is dependent on the actual strategies chosen by deliberation scheduling which, in our particular case, depends on EIV and this value of falling back in. We might possibly resolve the circularity by solving a large and very complicated set of simultaneous equations; in practice, we have found that it is not difficult to find a value that works reasonably well.

6.2.2 Experimental Results

Domain The experimental results for the recurrent model were obtained on the mobile-robot domain in a range of sizes. Table 1 shows the numbers of locations and states for the different sized domains in the columns **nLocs** and **nStates** respectively. The larger domains are obtained by combining multiples of the floor plan shown in Figure 4 into two-dimensional grids.

The actions available to the agent were the same as those described in Section 4.3 and

World	nLocs	nStates	ExpCost	CostOut
1x1	158	632	-500	-1
2x2	632	2528	-1000	-4
3x3	1422	5688	-1500	-9
4x4	2528	10112	-2000	-16
5x5	3950	15800	-2500	-25

Table 1: Information about the domains used to obtain experimental results for the recurrent deliberation model

used to obtain the precursor-model results. The transition probabilities were also the same, except that the domain was modified slightly to no longer contain any complete sinks; this means that even if the agent falls into a *semi-sink* state (i.e. one with low but non-zero probabilities of making transitions into another state), it can eventually reach the goal. In each case, the discount factor, γ , was 0.9999.

In addition to the numbers of locations and states, Table 1 also shows the values used for **CostOut** and **ExpCost** when generating the statistics; **CostOut** is the estimate of how long the agent must wait once it has fallen out of the envelope using only its reactive policy and **ExpCost** is the estimate of the average value of the states in the world. These values are conservative estimates given the relatively benign nature of the domain — about 3% of the states are semi-sinks.²

Deliberation strategies Our implementation provided a number of phases, including envelope **optimization** (0) and the following types of envelope alteration:

1. **findfirstpath** (FFP): find 10 paths from the agent’s current state, x_{cur} , to a goal state, and chose the shortest path to be the initial envelope
2. **findpath** (FP): if the agent’s current state x_{cur} is not in the envelope, find a path from x_{cur} back to the envelope, and add this path to the envelope

²Given $\gamma = 0.9999$, the value of a complete sink is $-1/(1 - \gamma) = -10,000$.

3. **robustify** ($R[N]$): we used the following heuristic to extend the envelope: find the N most likely fringe states and add them to the envelope
4. **prune** ($P[N]$): of the states that have a worse value than the current state, remove the N least likely to be reached using the current policy.

The **findFirstPath** phase is executed only once, to obtain the initial envelope, then the deliberation module chooses between a set of 24 hand-crafted strategies. Each of these 24 strategies begins with a **findpath** phase and ends with the **optimization** phase. Between these first and last phases, robustification, pruning and optimization are used in different combinations with different numbers of states to be added or deleted. In order to be able to compare the same strategy for the different sized domains, we formulated the number of states to be added and deleted in terms of the dimensions of the world; for world size $n \times n$, where $n = 1, \dots, 5$, the number of states to be added or deleted was one of $\{5n^2, 10n^2, 20n^2\}$. Examples strategies are:

```
{FP R[10n2] 0}
{FP P[20n2] 0}
{FP P[5n2] R[10n2] 0}
{FP R[20n2] P[10n2] 0}
{FP R[10n2] 0 P[10n2] 0}
```

Statistics We collected statistics over a large number of runs (where a run is a particular start/goal pair) for each size domain, generating data points for strategy execution as shown in Table 2.

The start/goal pairs were chosen uniformly at random from all states excluding the semi-sinks and we ran the simulated robot in parallel with the planner until the goal was reached. The planner executed **findFirstpath** (FFP) to obtain the initial envelope, then executed the following loop: choose one of the 24 strategies uniformly at random, execute that strategy, and then pass the new policy to the simulated robot.

Conditioning attributes We found the following conditioning variables to be significant: the envelope size, $|E|$, the estimated value of the current state $\hat{V}_\pi$, the “fatness” of the envelope (the ratio of envelope size to fringe size), and the Manhattan distance, M , between

World	# Stats Runs	# Data Points
1x1	5858	155696
2x2	6423	168215
3x3	5758	128611
4x4	5363	110830
5x5	4963	94021
total	28365	657373

Table 2: Statistics for different sized domains

the start and goal locations. We then built the lookup tables of the expected improvement in value as a function of $|\mathcal{E}|$, $\hat{V}_\pi$, the fatness, M and the strategy s . The lookup table granularity used was 3 buckets per attribute dimension.

Simulation results To test our algorithm, we took 50 pairs of start and goal states from each world, chosen uniformly at random from all states excluding the semi-sink states. For each pair we ran the simulated robot in parallel with the following deliberation mechanisms:

- recurrent-deliberation with strategies chosen using statistical estimates of EIV (LOOKUP)
- dynamic programming policy iteration over the entire domain, with a new policy given to the robot
 - after each iteration (ITER)
 - only after it has been completely optimized (WHOLE)

We found that the statistics were not particularly sensitive to the size of the domain; statistics gathered for the smaller-sized worlds transferred fairly well to the larger-sized worlds. The LOOKUP results use the statistics lookup table compiled from the approximately 660,000 data points.

Figure 8 shows the average number of steps taken by the agent to reach the goal for the various algorithms. For the smaller domains, the recurrent-deliberation algorithm does

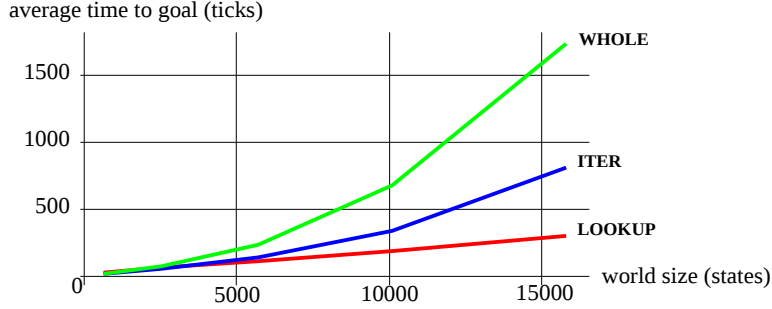


Figure 8: Comparison of the recurrent algorithm to policy iteration over varying domain size

not perform better than either of the policy iteration algorithms. However, as we move to larger domains, the improvement is marked. As we might expect, **WHOLE** is exponential and becomes computationally infeasible as the size of the domain increases. **ITER** also shows a non-linear degradation in the time to goal. **LOOKUP** shows linear behavior, clearly performing better as the domain size increases.

The implementation used to obtain these experimental results did not include a separate trajectory planning phase that looks for new paths to goal states. This lack of exploration meant that the planner did not look for shortcuts either to states in the envelope but significantly closer to the goal, or to the goal itself. Therefore, the performance depended more on the quality of the first path used as the initial envelope than it might have if trajectory planning had been implemented. Without the exploratory trajectory planning, all the significant path planning is done when the initial envelope is found; in this case, if that initial envelope contains a path of length close to the Manhattan distance, the recurrent algorithm performs quite well. However if the first path is not a good one, the recurrent algorithm ends up exploring most of the state space and loses its performance gains over the dynamic programming algorithm. The addition of trajectory planning would be likely to improve the performance of the system as a whole.

7 Representing Goals with Reward Functions

In early AI work on planning, it was traditional to have a goal of achievement specified by a logical expression over properties of world states. This translates into having a set of desirable world states and the implicit goal to reach one of these states in the least possible amount of time. At the same time, work in temporal and dynamic logics gave us the notions of a proposition being *always* true or *eventually* true and of one predicate being true *until* another became true. These ideas have been attractive to AI researchers because they give us a more complex, compositional language in which to express goals. Unfortunately, these expressions are not suitable for direct use as goals in AI applications. An agent with the goal of *eventually*(p) has the option of postponing p indefinitely; there is no requirement to achieve p sooner rather than later. Similarly, to have the goal *always*(p) is to require that p be maintained true into the infinite future, which is impossible in any sort of real world.

7.1 Basic Goal Types

One way to retain these ideas but make them more useful is to replace *eventually* by *asap*, meaning, intuitively, *as soon as possible* and to replace *always* by *alap*, meaning *as long as possible*. In stochastic domains, we can convert goals of this kind into reward functions and apply the same algorithms for finding good policies with respect to the reward functions.

Given a goal of *asap*(p), we can generate a policy that takes actions in such a way as to minimize the expected number of steps taken before a state in which p holds is entered. First, we generate the reward function

$$r_p(s) = \begin{cases} 1 & \text{if } p(s) \\ -1 & \text{otherwise} \end{cases}$$

We must also modify the environment so that all states s such that $p(s)$ are absorbing; this ensures that we go to the “nearest” state in which p holds, independent of the states that will follow.

Similarly, given a goal of *alap*(p), we can generate a policy that takes actions in such a way as to maximize the number of steps taken before a state in which $\neg p$ holds is entered. We can use the reward function given above, but this time, we modify the environment so

that all states s such that $\neg p(s)$ are absorbing; this ensures that no good results can ensue after encountering a state in which p does not hold.

If we have a longer-term goal of staying in states in which p holds *as much as possible*, that is $amap(p)$, then we need only adopt the reward function above and make no changes to the environment.

7.2 Goal Combination

It is often useful to think of an agent as having multiple goals simultaneously. Goals based on reward functions can be combined to achieve this effect, although the kinds of combination that are appropriate are different than for logical goals.

Goals of achievement can be disjoined in two ways: either $asap(p) \vee asap(q)$ or $asap(p \vee q)$. The first method requires that either p be achieved as soon as possible or that q be achieved as soon as possible, but is indifferent between them. This kind of combination will rarely be useful, because it would allow p to be pursued, even though it takes much longer than achieving q . We therefore prefer the second form, which can be performed by taking the maximum of the reward functions at each state and making any state with zero reward absorbing. Prioritized disjunction (in which, for instance, states in which p holds are to be preferred to states in which q holds, but only if the length of time to achieve p is not too much greater) can be achieved by taking the maximum of scaled versions of the reward functions:

$$r(s) = \max(\alpha r_p(s), \beta r_q(s))$$

where α and β encode the relative desirability of achieving p and achieving q . The same technique can be applied to create reward functions for the goals $alap(p \vee q)$ and $amap(p \vee q)$.

External conjunction of *asap* goals is problematic: what does it mean to achieve p as soon as possible *and* to achieve q as soon as possible? In general, these goals will be conflicting and the conjunction is meaningless. However, we can accommodate internal conjunction by taking the minimum of the reward function at each state. Again, we can apply the same technique to create reward functions for the goals $alap(p \wedge q)$ and $amap(p \wedge q)$.

Another useful goal combination is *asap-maint*(p, q), in which the goal is to achieve p as soon as possible, and to maintain q until p has been achieved. This can be accomplished using the same reward function as before, but making all states in which $\neg q$ holds absorbing

as well. Even though there is no difference in instantaneous value between states in which q does and does not hold, the $\neg q$ states are both bad and absorbing, which will give them a very high negative value. This is essentially what was done in the experimental domain described earlier, with the stairwells being absorbing states.

Many languages for the combination of goals allow sequencing, in which it is specified that p is to be achieved, then q is to be achieved. If it is not possible for the agent to perceive or remember that p has been achieved, then sequenced goals cannot be specified using reward functions. If the agent can perceive that p has been achieved (notated $prev(p)$), then the goal $asap(prev(p) \wedge q)$ will have the desired effect.

7.3 Modifying the Planning Algorithm

The planning algorithm described earlier can be applied directly to the whole range of possible reward functions, but some aspects can be tuned to improve the early behavior of the algorithm.

For *asap* goals, it makes sense for the initial envelope to include some path, however tenuous, from the current state to some goal state. If this is not the case, there is no basis for assigning value to the states in the envelope and the policy will essentially be random. For *alap* goals, it is sufficient for the initial envelope to simply be the current state. A good envelope for an *alap* goal need only contain a cycle or set of states that the agent can stay in with high probability.

In addition, the appropriateness of various deliberation strategies will also depend on the type of the goal. This dependence can be handled directly by the statistics-based deliberation-scheduling mechanisms described above.

8 Uncertainty in Observation

Howard's and Bellman's formulations of dynamic programming address sequential decision problems of indefinite duration in which the system dynamics can be described as a Markov chain. They assume that the decision-making agent is able to observe its actual state with absolute certainty when executing a policy. In the previous sections, we have adopted this same basic model for decision making. In this section, we are concerned with decision

problems of indefinite duration in which the assumption regarding observation does not hold. For instance, if an agent is in state x then some of the time (perhaps most of the time) it observes x and the rest of the time it observes some other state. In the case in which it observes a state other than its actual state, we say that the agent has made an error in observation. We can specify the uncertainty as a conditional probability distribution relating observed and actual states. The state transitions are still governed by a Markov chain, but the ability of the agent to execute a given policy is dependent on its ability to correctly identify states.

Simple framework In this section, we cast the problem of decision making with uncertainty in observation within the framework described in previous sections. In particular, we demonstrate how observation errors can be modeled using Markov decision models, and sketch algorithms for heuristic trajectory planning and envelope alteration that account for uncertainty in observation. We describe an approach in which the only information used by the agent in estimating its current state is the observation made on entering that state. In addition, we describe a more general approach in which state estimation accounts for the entire past history of observations and actions. We begin with some terminology.

We decompose the problem into designing an *observer* to identify the current state and a *regulator* to determine what actions to take assuming that the agent can identify the current state with absolute certainty. In the control literature this approach would be said to assume the *separability* of observation and regulation. The approach is described as follows. First, determine the optimal regulator assuming that observation is perfect; this can be done using Howard’s policy iteration. Second, design the optimal observer using a model for the observation errors; there is a substantial literature on using Bayesian decision theory to design optimal observers [Bar-Shalom and Fortmann, 1988]. Finally, couple the optimal observer to the optimal regulator.

The problem with the above strategy is that the regulator, by assuming a perfect observer, may drive the agent into regions of the state space in which it is very hard to correctly identify the current state and so having the optimal regulator is not much help. We can avoid this problem by designing a regulator that takes into account the abilities of the observer. In the following, we describe approaches to doing this that fit within the framework presented in previous sections.

Here is the standard formulation for sequential decision problems (which differs slightly in notation from the one we used earlier):

- $p_{ij}(u)$ — the probability of making a transition from x_i to x_j having executed action u in state x_i
- $R(x_i)$ — the reward for being in x_i

We abbreviate the state-transition probabilities for a given policy π as $p_{ij} = p_{ij}(\pi(x_i))$.

We assume that the observer takes in observations and generates as output the most likely current state based on its observations. The observer can take into account only the most recent observation or observations arbitrarily far into the future. In the simplest case, we can model the uncertainty of the observer's output by:

- g_{ij} — the probability the observer outputs x_i when the state is really x_j

The standard approach uses Howard's policy iteration algorithm to find the optimal policy, π^* , which satisfies the following system of equations:

$$V_{\pi^*}(x_i) = R(x_i) + \max_a \gamma \sum_{x_j \in X} p_{ij}(a) V_{\pi^*}(x_j) .$$

When we have uncertainty in observation, it is possible to define the value function over world states given a particular policy π as

$$V_{\pi}(x_i) = R(x_i) + \gamma \sum_{x_j \in X} \sum_{x_k \in X} g_{ki} p_{ij}(\pi(x_k)) V_{\pi}(x_j) .$$

The value of a state x_i is the instantaneous reward plus the expected value of the next state. The probability of making a transition from state x_i to some next state x_j is more complicated when there is uncertainty in observation. It is the probability of believing that you are really in some state x_k given that you are actually in state x_i (that is, g_{ki}), times the probability of making a transition from x_i to x_j given that you take the action appropriate to being in x_k (that is, $p_{ij}(x_k)$).

Now, it is possible to define the optimal uncertain-observation policy as

$$\pi^*(x_i) = \arg \max_a \sum_{x_j \in X} p_{ij}(a) V_{\pi^*}(x_j) ,$$

which simply requires you to take the action that would be the best for the state you believe yourself to be in.

Unfortunately, the techniques of policy iteration cannot be applied to compute this policy. Any change of action in one state affects the values of all the other states in a complex way. It may be the case that there is no alternative but to enumerate the policies in order to find the optimal. We are currently investigating approximation methods for this problem.

It is important to point out that this approach does not result in an optimal observer; it results in an optimal regulator assuming a particular observer. The regulator does what it can to work within the limitations of the chosen observer and requires that the performance of the observer be modeled simply by the g_{ij} values. In the following, we consider an observer that accounts for all of the past observations and actions.

Partially observable Markov decision processes In operations research, Markov decision processes that involve uncertainty in observation are called *partially observable*. There is a large literature on partially observable Markov decision processes (see [Monahan, 1982] for a survey). A standard technique is to convert the partially observable process into a perfectly observable process with states that correspond to distributions over the original states, summarizing the agent's knowledge of its current state given all of its past actions and observations. The resulting, perfectly-observable process is Markov.

The problem with the above approach is that the state space of the perfectly-observable process is no longer finite and, in some cases, a sequence of policies rather than a single policy is required for optimality. A number of techniques have been developed that serve to partition the state space into a finite set of regions or consider only a finite, reachable subspace of the infinite state space. Smallwood and Sondik describe state-space-partitioning algorithms for the finite-horizon [Smallwood and Sondik, 1973] and indefinite-duration (infinite-horizon) [Sondik, 1978] cases. In the worst case, the number of regions in the partition can grow exponentially in the time horizon and so much of the current research involves approximation algorithms [Lovejoy, 1991, White and Scherer, 1989].

Integration into our approach The above discussion was meant to provide some evidence that our basic approach can be extended to take into account uncertainty in obser-

vation. The simplest approach to extension would be to make use of existing methods for optimization and approximation and provide analogs of the envelope alteration routines introduced earlier to handle the case of uncertainty in observation. In the following, we sketch what these analogs would look like.

In earlier sections, we describe algorithms for trajectory planning and envelope extension for the case in which there are no errors in observation. For trajectory planning, we use a heuristic function that takes into account the probability of traversing a trajectory with a fixed policy. During trajectory planning, each trajectory is associated with a particular partial policy and a traversal probability. For a trajectory ending in x_i to be extended to x_j by action u , either the associated policy already maps x_i to u or the policy must be extended so that $\pi(x_i) = u$. The traversal probability is updated to be the traversal probability for the unextended trajectory times the probability of making the most recent transition, $p_{ij}(u)$.

In the case of observation errors, the basic idea is quite similar. In the case of the simple decision model described at the beginning of this section, the traversal probability for the unextended trajectory is multiplied by the probability that the agent will be able to follow the most recent transition given both errors in observation and errors in movement as summarized by $f_{ij} = \sum_{x_k \in X} g_{ki} p_{ij}(\pi(x_k))$. In the case in which the states of the Markov process correspond to distributions, the approach is more complicated to describe without introducing more notation than seems warranted but the idea is similar. The heuristic methods for adding fringe states and pruning are likewise similar in their implementation and intuitive basis.

9 Related Work

Our primary interest is in applying the sequential decision making techniques of Bellman [Bellman, 1957] and Howard [Howard, 1960] in time-critical applications. Our initial motivation for the methods discussed here came from the ‘anytime synthetic projection’ work of Drummond and Bresina. [Drummond and Bresina, 1990]. We improve on the Drummond and Bresina work by providing (i) coherent semantics for goals in stochastic domains, (ii) theoretically sound probabilistic foundations, (iii) and decision-theoretic methods for controlling inference.

The approach described in this paper represents a particular instance of time-dependent planning [Dean and Boddy, 1988] and borrows from, among others, Horvitz' [Horvitz, 1988] approach to flexible computation. Boddy [Boddy, 1991] describes solutions to related problems involving dynamic programming. Hansson and Mayer's BPS (Bayesian Problem Solver) [Hansson and Mayer, 1989] supports general state-space search with decision-theoretic control of inference; it may be that BPS could be used as the basis for envelope extension thus providing more fine-grained decision-theoretic control. Christiansen and Goldberg [Christiansen and Goldberg, 1990] and Kushmerick, Hanks, and Weld [Kushmerick *et al.*, 1993] also address the problem of planning in stochastic domains. Boddy [Boddy, 1991] describes solutions to related problems involving dynamic programming. For an overview of resource-bounded decision making methods, see chapter 8 of the text by Dean and Wellman [Dean and Wellman, 1991].

Acknowledgements

Thomas Dean's work was supported in part by a National Science Foundation Presidential Young Investigator Award IRI-8957601, in part by the Advanced Research Projects Agency of the Department of Defense monitored by the Air Force under Contract No. F30602-91-C-0041, and in part by the National Science foundation in conjunction with the Advanced Research Projects Agency of the Department of Defense under Contract No. IRI-8905436. Leslie Kaelbling's work was supported in part by a National Science Foundation National Young Investigator Award IRI-9257592 and in part by ONR Contract N00014-91-4052, ARPA Order 8225.

References

- [Bar-Shalom and Fortmann, 1988] Bar-Shalom, Yaakov and Fortmann, Thomas E. 1988. *Tracking and Data Association*. Academic Press, New York.
- [Bellman, 1957] Bellman, Richard 1957. *Dynamic Programming*. Princeton University Press.
- [Bertsekas, 1987] Bertsekas, Dimitri P. 1987. *Dynamic Programming*. Prentice-Hall, Englewood Cliffs, N.J.

- [Boddy, 1991] Boddy, Mark 1991. Anytime problem solving using dynamic programming. In *Proceedings AAAI-91*. AAAI. 738–743.
- [Christiansen and Goldberg, 1990] Christiansen, Alan and Goldberg, Ken 1990. Robotic manipulation planning with stochastic actions. In *DARPA Workshop on Innovative Approaches to Planning, Scheduling and Control*. San Diego, California.
- [Dean and Boddy, 1988] Dean, Thomas and Boddy, Mark 1988. An analysis of time-dependent planning. In *Proceedings AAAI-88*. AAAI. 49–54.
- [Dean and Wellman, 1991] Dean, Thomas and Wellman, Michael 1991. *Planning and Control*. Morgan Kaufmann, San Mateo, California.
- [Drummond and Bresina, 1990] Drummond, Mark and Bresina, John 1990. Anytime synthetic projection: Maximizing the probability of goal satisfaction. In *Proceedings AAAI-90*. AAAI. 138–144.
- [Fikes *et al.*, 1972] Fikes, Richard E.; Hart, Peter E.; and Nilsson, Nils J. 1972. Learning and executing generalized robot plans. *Artificial Intelligence* 3:251–288.
- [Hansson and Mayer, 1989] Hansson, Othar and Mayer, Andrew 1989. Heuristic search as evidential reasoning. In *Proceedings of the Fifth Workshop on Uncertainty in AI*. 152–161.
- [Horvitz, 1988] Horvitz, Eric J. 1988. Reasoning under varying and uncertain resource constraints. In *Proceedings AAAI-88*. AAAI. 111–116.
- [Howard, 1960] Howard, Ronald A. 1960. *Dynamic Programming and Markov Processes*. MIT Press, Cambridge, Massachusetts.
- [Kemeny and Snell, 1960] Kemeny, J. G. and Snell, J. L. 1960. *Finite Markov Chains*. D. Van Nostrand, New York.
- [Kushmerick *et al.*, 1993] Kushmerick, Nicholas; Hanks, Steve; and Weld, Daniel 1993. An algorithm for probabilistic planning. Unpublished Manuscript.
- [Lovejoy, 1991] Lovejoy, William S. 1991. Computationally feasible bounds for partially observed markov decision processes. *Operations Research* 39(1):162–175.

- [Monahan, 1982] Monahan, George E. 1982. A survey of partially observable markov decision processes: Theory, models, and algorithms. *Management Science* 28(1):1–16.
- [Schoppers, 1987] Schoppers, Marcel J. 1987. Universal plans for reactive robots in unpredictable environments. In *Proceedings IJCAI 10*. IJCAI. 1039–1046.
- [Smallwood and Sondik, 1973] Smallwood, Richard D. and Sondik, Edward J. 1973. The optimal control of partially observable markov processes over a finite horizon. *Operations Research* 21:1071–1088.
- [Sondik, 1978] Sondik, Edward J. 1978. The optimal control of partially observable markov processes over the infinite horizon: Discounted cost. *Operations Research* 26(2):282–304.
- [White and Scherer, 1989] White, Chelsea C. III and Scherer, William T. 1989. Solution procedures for partially observed markov decision processes. *Operations Research* 37(5):791–797.
- [Wilkins, 1988] Wilkins, David E. 1988. *Practical Planning: Extending the Classical AI Planning Paradigm*. Morgan-Kaufmann, Los Altos, California.