

# On How to Learn from a Stochastic Teacher or a Stochastic Compulsive Liar of Unknown Identity

B. John Oommen<sup>1</sup>, Govindachari Raghunath<sup>2</sup>, and Benjamin Kuipers<sup>3</sup>

<sup>1</sup> Fellow of the IEEE. School of Computer Science  
Carleton University, Ottawa, ON, K1S 5B6, Canada

oommen@scs.carleton.ca

<sup>2</sup> oommen@scs.carleton.ca

<sup>3</sup> Fellow of the AAAI and IEEE. Department of Computer Sciences  
University of Texas, Austin, Texas 78712, USA

kuipers@cs.utexas.edu

**Abstract.** We consider the problem of a learning mechanism (robot, or algorithm) that learns a parameter while interacting with either a *stochastic teacher* or a *stochastic compulsive liar*. The problem is modeled as follows: the learning mechanism is trying to locate an unknown point on a real interval by interacting with a stochastic environment through a series of guesses. For each guess the environment (teacher) essentially informs the mechanism, possibly erroneously, which way it should move to reach the point. Thus, there is a non-zero probability that the feedback from the environment is erroneous. When the probability of correct response is  $p > 0.5$ , the environment is said to be *Informative*, and we have the case of learning from a stochastic teacher. When this probability is  $p < 0.5$  the environment is deemed *Deceptive*, and is called a *stochastic compulsive liar*.

This paper describes a novel learning strategy by which the unknown parameter can be learned in both environments. To the best of our knowledge, our results are the first reported results which are applicable to the latter scenario. Another main contribution of this paper is that the proposed scheme is shown to operate equally well even when the learning mechanism is unaware whether the environment is Informative or Deceptive. The learning strategy proposed herein, called CPL-ATS, partitions the search interval into three equi-sized sub-intervals, evaluates the location of the unknown point with respect to these sub-intervals using fast-converging  $\epsilon$ -optimal  $L_{RI}$  learning automata, and prunes the search space in each iteration by eliminating at least one partition. The CPL-ATS algorithm is shown to be provably converging to the unknown point to an arbitrary degree of accuracy with probability as close to unity as desired. Comprehensive experimental results confirm the fast and accurate convergence of the search for a wide range of values for the environment's feedback accuracy parameter  $p$ . The above algorithm can be used to learn parameters for non-linear optimization techniques.

**Keywords:** Learning Automata, Statistical Learning, Parameter Estimation, Stochastic Optimization, Pattern Recognition.

## 1 Introduction

Consider the problem of a robot (algorithm, learning mechanism) moving along the real line attempting to locate a particular point  $\lambda^*$ . To assist the mechanism, we assume that it can communicate with an Environment (“Oracle”) which guides it with information regarding the direction in which it should go. If the Environment is deterministic the problem is the “Deterministic Point Location Problem” which has been studied rather thoroughly [1]. In its pioneering version [1] the problem was presented in the setting that the Environment could charge the robot a cost which was proportional to the distance it was from the point sought for. The question of having multiple communicating robots locate a point on the line has also been studied [1, 2]. In the stochastic version of this problem, we consider the scenario when the learning mechanism attempts to locate a point in an interval with stochastic (i.e., possibly erroneous) instead of deterministic responses from the environment. Thus when it should really be moving to the “right” it may be advised to move to the “left” and vice versa.

Apart from the problem being of importance in its own right, the stochastic point location problem also has potential applications in solving optimization problems. In many optimization solutions – for example in image processing, pattern recognition and neural computing [5, 9, 11, 12, 14, 16, 19], the algorithm works its way from its current solution to the optimal solution based on information that it currently has. A crucial question is one of determining the *parameter* which the optimization algorithm should use. In many cases the parameter of the scheme is related to the second derivative of the criterion function, which results in a technique analogous to a Newton’s root solving scheme. The disadvantages of the latter are well known – if the starting point of the algorithm is not well chosen, the scheme can diverge. Additionally, if the second derivative is small, the scheme is ill-defined. Finally, such a scheme requires the additional computation involved in evaluating the (matrix of) second derivatives [14, 16, 19]. To tackle this problem we suggest that our strategy to solve the stochastic point location problem can be invoked to learn the best parameter to be used in any algorithm.

The results that are claimed here are stated in the body of the paper. Sketches of their proofs are included in the Appendix. The more detailed proofs are found in the unabridged version of the paper [15].

## 2 The Stochastic Point Location Problem

The goal of the learning mechanism is to determine the optimal value of some parameter  $\lambda^* \in [0, 1)$ . Although the mechanism does not know the value of  $\lambda^*$ , we assume that it has responses from an intelligent environment  $E$  which is capable of informing it whether the current estimate  $\hat{\lambda}$  is too small or too big. To render the problem both meaningful and distinct from its deterministic version, we would like to emphasize that the response from this environment is assumed “faulty”. Thus,  $E$  may tell us to increase  $\hat{\lambda}$  when it should be decreased, and vice

versa with a non-zero probability  $1 - p$ . Note that the quantity  $p$  reflects on the “effectiveness” of the environment,  $E$ . Thus, whenever  $\hat{\lambda} < \lambda^*$ , the environment correctly suggests that we increase  $\hat{\lambda}$  with probability  $p$ . It could as well have incorrectly recommended that we decrease  $\hat{\lambda}$  with probability  $1 - p$ . Similarly, whenever  $\hat{\lambda} > \lambda^*$ , the environment tells us to decrease  $\hat{\lambda}$  with probability  $p$ , and to increase it with probability  $(1 - p)$ .

We further distinguish between two types of environments – *Informative* and *Deceptive*. An environment is said to be “Informative” if the probability  $p$  of its giving a correct feedback is greater than 0.5. If  $p < 0.5$ , the environment is said to be “Deceptive”. Thus a Deceptive environment is more likely to give erroneous feedback than an Informative environment. This together with the fact that the learning mechanism is not aware of the nature of the environment complicates the learning process and its convergence.

### 3 Related Work

Oommen [9] proposed and analyzed an algorithm that operates by discretizing the search space while interacting with an *Informative* environment. This algorithm takes advantage of the limited precision available in practical implementations to restrict the probability of choosing an action to only finitely many values from the interval  $[0, 1)$ . Its main drawback is that the steps are always very conservative. If the step size is increased the scheme converges faster, but the accuracy is correspondingly decreased. Bentley and Yao [3] solved the deterministic point location problem of searching in an unbounded space by examining points  $f(i)$  and  $f(i + 1)$  at two successive iterations between which the unknown point lies, and doing a binary search between these points. Although it may appear that similar binary search can be applied in the stochastic point location problem, the faulty nature of the feedback from the environment may affect the certainty of convergence of the search, and hence a more sophisticated search strategy is called for. Thus, whereas in Bentley’s and Yao’s algorithm we could confidently discard regions of the search space, we have to now resort to stochastic methods, and work so that we minimize the probability of rejecting an interval of interest. This is even more significant and sensitive when the learning mechanism is unaware whether the environment is Informative or Deceptive. A novel strategy combining learning automata and pruning was used in [10] to search for the parameter in the continuous space when interacting with an Informative environment. In this paper we extend the results of [10] further to operate also in the continuous space, but to work equally well for both *Informative* and *Deceptive* environments. To the best of our knowledge this is the first reported result for learning in a Deceptive environment.

A completely different approach to locating a point will be to partition the search space into intervals and choose the mid-point of each interval repeatedly as an estimate for the unknown point, and use a majority voting scheme on the feedbacks obtained to eliminate one interval. Applying Chernoff bounds would then allow us to precisely compute the number of steps sufficient for ensuring

correct pruning with a certain level of confidence [13]. The difference between such an “estimation” approach and a learning automaton approach such as ours is that in the former case, a fixed number of sampling steps has to be done *irrespective* of how strong the feedback is, before determining the pruning. In contrast, in a learning automata approach, when one of the actions is superior, the reinforcement mechanism ensures that this action is sampled more frequently than the other. Thus the better an action is with respect to another, the faster it’s optimality is determined. Another drawback with the simple majority voting scheme is that in each epoch it can prune at most one partition. On the other hand, our scheme allows more than two thirds of the search space to be pruned in a single epoch. A third and most compelling drawback of the former scheme in our problem setting is that it will not be feasible to detect a point in a deceptive environment where the probability of correct feedback  $p < 0.5$ .

## 4 Continuous Point Location with Adaptive Tertiary Search

The solution presented in this paper is based on the Continuous Point Location with Adaptive Tertiary Search (CPL–ATS) strategy introduced in [10]. The basic idea behind both the solutions is to systematically explore a current interval for the parameter. This exploration is a series of guesses, each one more accurate than the previous one. Initially, we guess the mid-point of the given interval to be our estimate. We then partition the given interval into disjoint sub-intervals, eliminating at least one of the subintervals from further search and recursively searching the remaining interval until the search interval is at least as small as the required resolution of estimation. Crucial to the CPL–ATS scheme is the construction of partitions and the elimination process.

The given search interval is divided into three partitions. Each of these three partitions is independently explored using an  $\epsilon$ -optimal fast converging two-action learning automaton where the two actions are those of selecting a point from the left and right half of the partition under consideration. The elimination process then utilizes the  $\epsilon$ -optimality property of the underlying automata and the monotonicity of the intervals to systematically eliminate at least one of the partitions. The resultant is a new pruned interval consisting of at most two contiguous partitions from the original three. At each stage the mid-point of the remaining interval is taken to be the estimate of  $\lambda^*$ . We shall assume that the individual learning automaton used is the well-known  $L_{RI}$  scheme with parameter  $\theta$ , although any other  $\epsilon$ -optimal scheme can be used just as effectively.

## 5 Notations and Definitions

Let  $\Delta = [\sigma, \gamma]$  s.t.  $\sigma \leq \lambda^* < \gamma$  be the current search interval containing  $\lambda^*$  whose left and right (smaller and greater) boundaries on the real line are  $\sigma$  and  $\gamma$  respectively. We partition  $\Delta$  into three equi-sized disjoint partitions  $\Delta^j$ ,  $j \in \{1, 2, 3\}$ ,

such that,  $\Delta^j = [\sigma^j, \gamma^j]$ . To formally describe the relative locations of intervals we define an interval relational operator  $\prec$  such that,  $\Delta^j \prec \Delta^k$  iff  $\gamma^j < \sigma^k$ . Since points on the real interval are monotonically increasing, it is easy to see that,  $\Delta^1 \prec \Delta^2 \prec \Delta^3$ . For every partition  $\Delta^j$ , we define  $L^j$  and  $R^j$  as its *left* half and *right* half respectively as:

$$L^j = \{x \mid \sigma^j \leq x < \text{mid}(\Delta^j)\}, \text{ and } R^j = \{x \mid \text{mid}(\Delta^j) \leq x < \gamma^j\},$$

where  $\text{mid}(\Delta^j)$  is the mid point of the partition  $\Delta^j$ . A point  $x \in L^j$  will be denoted by  $x_L^j$ , and a point  $x \in R^j$  by  $x_R^j$ .

To relate the various intervals to  $\lambda^*$  we introduce the following relational operators.

$$\begin{array}{ll} \lambda^* \ominus \Delta^j \text{ iff } \lambda^* < \sigma^j. & \text{i.e., } \lambda^* \text{ is to the left of the interval } \Delta^j. \\ \lambda^* \oplus \Delta^j \text{ iff } \lambda^* > \gamma^j. & \text{i.e., } \lambda^* \text{ is to the right of the interval } \Delta^j. \\ \lambda^* \ominus \Delta^j \text{ iff } \sigma^j \leq \lambda^* < \gamma^j. & \text{i.e., } \lambda^* \text{ is contained in the interval } \Delta^j. \\ \lambda^* \otimes \Delta^j \text{ iff } \lambda^* \ominus \Delta^j \text{ or } \lambda^* \oplus \Delta^j & \text{i.e., } \lambda^* \text{ is either to the left of or inside } \Delta^j. \\ \lambda^* \oslash \Delta^j \text{ iff } \lambda^* \oplus \Delta^j \text{ or } \lambda^* \ominus \Delta^j & \text{i.e., } \lambda^* \text{ is either to the right of or inside } \Delta^j. \end{array}$$

These operators can be shown to satisfy the usual laws of transitivity.

## 6 Construction of the Learning Automata

In the CPL-ATS strategy, with each partition  $\Delta^j$  we associate a 2-action  $L_{RI}$  automaton  $\mathcal{A}^j$  ( $\Sigma^j, \Pi^j, \Gamma^j, \Upsilon^j, \Omega^j$ ) where,  $\Sigma^j$  is the set of actions,  $\Pi^j$  is the set of action probabilities,  $\Gamma^j$  is the set of feedback inputs from the environment,  $\Upsilon^j$  is the set of action probability updating rules and  $\Omega^j$  is the set of possible decision outputs of the automaton at the end of each epoch. The environment E is characterized by the probability of correct response  $p$  which is later mapped to the penalty probabilities  $c_k^j$  for the two actions of the automaton  $\mathcal{A}^j$ . The overall search strategy CPL-ATS, in addition uses a decision table  $\Lambda$  to prune the search interval by comparing the output decisions  $\Omega^j$  for the three partitions. Thus  $\mathcal{A}^j$ ,  $j \in \{1, 3\}$ , together with E and  $\Lambda$  completely define the CPL-ATS strategy.

### 1. The set of actions of the automaton( $\Sigma^j$ )

The two actions of the automaton are  $\alpha_{k=0,1}^j$ , where,  $\alpha_0^j$  corresponds to selecting the *left* half  $L^j$  of the partition  $\Delta^j$ , and  $\alpha_1^j$  corresponds to selecting the *right* half  $R^j$ .

### 2. The action probabilities( $\Pi^j$ )

$P_k^j(n)$  represent the probabilities of selecting action  $\alpha_{k=0,1}^j$  at step  $n$ . Initially,  $P_k^j(0) = 0.5$ , for  $k = 0, 1$ .

### 3. The feedback inputs from the environment to each automaton ( $\Gamma^j$ )

It is important to recognize a subtle, but crucial point in the construction of the learning automata in CPL-ATS. From the automaton's point of view, the two actions are those of selecting either the left or the right half from

its partition. However, from the environment's point of view, the automaton presents a current guess for the true value of  $\lambda^*$ , and it gives a feedback based on the relative position (or direction) of the guessed value with respect to  $\lambda^*$ . Thus there is a need to map the intervals to a point value and the feedback on the point value to the feedback on the choice of the intervals. Let the automaton select either the left or right half of the partition and then pick a point randomly (using a continuous uniform probability distribution) from this sub-interval which is presented as the current estimate for  $\lambda^*$ . Then the feedbacks  $\beta(n)$  at step  $n$  are defined by the conditional probabilities,

$$\begin{aligned} Pr[\beta(n) = 0 \mid x_L^j \in L^j \text{ and } x_L^j \geq \lambda^*] &= p \\ Pr[\beta(n) = 1 \mid x_L^j \in L^j \text{ and } x_L^j < \lambda^*] &= p \\ Pr[\beta(n) = 0 \mid x_R^j \in R^j \text{ and } x_R^j < \lambda^*] &= p \\ Pr[\beta(n) = 1 \mid x_R^j \in R^j \text{ and } x_R^j \geq \lambda^*] &= p \end{aligned} \quad (1)$$

Note that, the condition  $x_L^j \in L^j$  indicates that the action  $\alpha_0^j$  was selected, and the condition  $x_R^j \in R^j$  indicates the other action  $\alpha_1^j$  was selected.

4. *The action probability updating rules ( $\Upsilon^j$ )*

First of all, since we are using the  $L_{RI}$  scheme, we ignore all the penalty responses. Upon reward, we obey the following updating rule :

If  $\alpha_{k=0,1}^j$  was rewarded then,

$$\begin{aligned} P_{1-k}^j[n+1] &\leftarrow \theta \cdot P_{1-k}^j[n] \\ P_k^j[n+1] &\leftarrow 1 - P_{1-k}^j[n+1] \end{aligned}$$

where  $0 \ll \theta < 1$  is the  $L_{RI}$  reward parameter.

5. *The decision outputs at each epoch ( $\Omega^j$ )*

From the action probabilities we infer the decision  $\Omega^j$  of the  $L_{RI}$  automaton  $\mathcal{A}^j$ , after a fixed number  $N_\infty$ , of iterations. Typically,  $N_\infty$  is chosen so as to be reasonably sure that the automaton has converged.  $\Omega^j$  indicates that the automaton has inferred whether  $\lambda^*$  is to the *left*, *right* or *inside* the partition. The set of values that  $\Omega^j$  can take and the preconditions are given by:

$$\Omega^j = \begin{cases} \textit{Left} & \text{if } P_0^j(N_\infty) \geq 1 - \epsilon. \\ \textit{Right} & \text{if } P_1^j(N_\infty) \geq 1 - \epsilon. \\ \textit{Inside} & \text{Otherwise.} \end{cases}$$

6. *The decision table for pruning the search space ( $\Lambda$ )*

Once the individual automata for the three partitions have made a decision regarding where they reckon  $\lambda^*$  to be, the CPL-ATS reduces the size of the search interval by eliminating at least one of the three partitions. The new pruned search interval  $\Delta^{new}$  for the subsequent learning phase (epoch) is generated according to the pruning decision table  $\Lambda$  shown in Table 1.

**Table 1.** Decision table( $\Delta$ ) to prune the search space based on the automata outputs  $\Omega^j$ 

| $\Omega^1$ | $\Omega^2$ | $\Omega^3$ | New Sub-interval $\Delta^{new}$ |
|------------|------------|------------|---------------------------------|
| Left       | Left       | Left       | $\Delta^1$                      |
| Inside     | Left       | Left       | $\Delta^1$                      |
| Right      | Left       | Left       | $\Delta^1 \cup \Delta^2$        |
| Right      | Inside     | Left       | $\Delta^2$                      |
| Right      | Right      | Left       | $\Delta^2 \cup \Delta^3$        |
| Right      | Right      | Inside     | $\Delta^3$                      |
| Right      | Right      | Right      | $\Delta^3$                      |

**Theorem 1.** If the automaton  $\mathcal{A}^j = (\Sigma^j, \Pi^j, \Gamma^j, \Upsilon^j, \Omega^j)$  interacts with the environment  $E$  and gets feedbacks obeying Equation (1), then the effective penalty probabilities  $c_{k=0,1}^j$  for the two actions  $\alpha_{k=0,1}^j$  are given by:

$$c_0^j = (1 - p) + (2p - 1) \cdot \Pr(x_L^j < \lambda^* \mid \alpha_0^j \text{ was chosen}) \quad (2)$$

$$c_1^j = p - (2p - 1) \cdot \Pr(x_R^j < \lambda^* \mid \alpha_1^j \text{ was chosen}) \quad (3)$$

From the theory of learning automata [6, 8], we know that for any 2-action  $L_{RI}$  scheme, if  $\exists k \in \{0, 1\}$  such that  $c_k^j < c_{1-k}^j$ , then the action  $\alpha_k^j$  is optimal and for this action  $P_k^j(N_\infty) \rightarrow 1$ .

By the construction of the automaton, once the left or right sub-interval is chosen, a point is picked from this interval using a uniform probability distribution. Therefore, the cumulative probability distributions for picking a guess value in a selected interval are given by:

$$\Pr(x_L^j < x \mid x_L^j \in L^j) = \begin{cases} 0 & \text{if } x < \sigma^j \\ \frac{x - \sigma^j}{\text{mid}(\Delta^j) - \sigma^j} & \text{if } \sigma^j \leq x < \text{mid}(\Delta^j) \\ 1 & \text{if } x \geq \text{mid}(\Delta^j). \end{cases} \quad (4)$$

$$\Pr(x_R^j < x \mid x_R^j \in R^j) = \begin{cases} 0 & \text{if } x < \text{mid}(\Delta^j) \\ \frac{x - \text{mid}(\Delta^j)}{\gamma^j - \text{mid}(\Delta^j)} & \text{if } \text{mid}(\Delta^j) \leq x < \gamma^j \\ 1 & \text{if } x \geq \gamma^j \end{cases} \quad (5)$$

By substituting  $\lambda^*$  for  $x$ , in the above equations, we get  $\Pr(x_L^j < \lambda^* \mid x_L^j \in L^j)$  and  $\Pr(x_R^j < \lambda^* \mid x_R^j \in R^j)$  in Equations (2) and (3) respectively.

## 7 Properties of CPL-ATS in an Informative Environment

We shall now state some results about how CPL-ATS behaves in an Informative Environment.

Lemmas 1 and 2 essentially use the  $\epsilon$ -optimality property of  $L_{RI}$  automata to prove that they produce the correct decision output for each partition under an

Informative environment. Lemma 3 uses the monotonicity of the real interval to establish some restrictions on the combination of decision outputs for adjacent partitions. These are then used in Theorem 2 to prove that the decision table in Table 1 is complete for the Informative environment. Theorem 3 establishes that after elimination of one or more partitions, the remaining interval still contains  $\lambda^*$ , thereby assuring convergence. It should be borne in mind that these are still probabilistic results, although the probability is shown to be potentially as close to unity as we want, provided that we choose the parameters for the  $L_{RI}$  automata appropriately.

**Lemma 1.** *For an Informative environment  $E$ , given the  $L_{RI}$  scheme with a parameter  $\theta$  which is arbitrarily close to unity, the following is true:*

$$\begin{aligned} \text{If } (\lambda^* \ominus \Delta^j), \text{ then } Pr(\Omega^j = \text{Left}) &\rightarrow 1. \\ \text{If } (\lambda^* \ominus \Delta^j), \text{ then } Pr(\Omega^j = \text{Right}) &\rightarrow 1. \\ \text{If } (\lambda^* \ominus \Delta^j), \text{ then } Pr(\Omega^j = \{\text{Left, Inside or Right}\}) &\rightarrow 1. \end{aligned}$$

**Lemma 2.** *For an Informative environment  $E$ , given the  $L_{RI}$  scheme with a parameter  $\theta$  which is arbitrarily close to unity, the following is true:*

$$\begin{aligned} \text{If } (\Omega^j = \text{Left}) \text{ then } Pr(\lambda^* \otimes \Delta^j) &\rightarrow 1 \\ \text{If } (\Omega^j = \text{Right}) \text{ then } Pr(\lambda^* \otimes \Delta^j) &\rightarrow 1 \\ \text{If } (\Omega^j = \text{Inside}) \text{ then } Pr(\lambda^* \otimes \Delta^j) &\rightarrow 1 \end{aligned}$$

**Lemma 3.** *In an Informative environment  $E$ , if the CPL-ATS learning mechanism uses the same  $L_{RI}$  scheme at all levels of the recursion, and the parameter  $\theta$  is arbitrarily close to unity, the following is true:*

$$\begin{aligned} \text{If } (\Omega^j = \text{Left}) \text{ then } Pr(\Omega^{j+1} = \text{Left}) &\rightarrow 1 \\ \text{If } (\Omega^j = \text{Inside}) \text{ then } Pr(\Omega^{j+1} = \text{Left}) &\rightarrow 1 \\ \text{If } (\Omega^j = \text{Right}) \text{ then } Pr(\Omega^{j+1} = \{\text{Left, Right or Inside}\}) &\rightarrow 1. \end{aligned}$$

**Theorem 2.** *If the environment is Informative and if the partitions use the same  $L_{RI}$  scheme with parameters  $\theta$  as close to unity as needed, then the decision table given in Table 1 is complete.*

A consequence of this theorem is that any entry not shown in the decision table is said to be *inconsistent* in the sense that for an Informative environment and appropriate  $\theta$ , the probability of occurrence of this entry is arbitrarily close to zero.

**Theorem 3.** *If the algorithm uses the same  $L_{RI}$  scheme at all levels of the recursion with a parameter  $\theta$  arbitrarily close to unity and  $N_\infty$  sufficiently large, then for an Informative environment, the unknown  $\lambda^*$  is always contained in the new search-interval  $\Delta^{\text{new}}$  resulting from the application of the decision rules of Table 1.*

## 8 Properties of CPL-ATS in a Deceptive Environment

Let  $E$  be an environment of a 2-action learning automaton. Let  $c_{k \in \{0,1\}}$  be the penalty probabilities of the two actions  $\alpha_{k \in \{0,1\}}$  of any automaton  $\mathcal{A}^j$  that operates in this environment. Then another environment  $E^*$  with penalty probabilities  $c'_{k \in \{0,1\}}$ , is said to be the *dual* of  $E$  if and only if under  $E^*$ ,  $c'_k = 1 - c_k$ .

The following lemma is a natural corollary of the above definition.

**Lemma 4.** *If  $\alpha_k^j$  is the  $\epsilon$ -optimal action for an  $L_{RI}$  automaton  $\mathcal{A}^j$  under a given environment  $E$ , then  $\alpha_{1-k}^j$  is the  $\epsilon$ -optimal action under its dual environment  $E^*$ , and vice versa.*

**Lemma 5.** *Let  $E$  be an environment with the probability of correct feedback  $p$  and corresponding penalty probabilities  $c_{k \in \{0,1\}}^j$  respectively. If  $E^*$  is a new environment constructed such that its probability of correct feedback  $p' = 1 - p$ , then the penalty probabilities  $c_{k \in \{0,1\}}^{j'}$  for its two actions  $\alpha_{k \in \{0,1\}}^j$  are  $c_k^{j'} = 1 - c_k^j$ .*

The above lemma follows easily by substituting  $1 - p$  in place of  $p$  in Equations (2) and (3) respectively. Thus we arrive at a dual of a given environment merely by complementing its parameter  $p$ .

Let  $E$  be the given Deceptive environment. By definition then, we have its probability of correct response  $p < 0.5$ . We now construct a dual  $E^*$  of this environment with a corresponding probability of correct response  $p' = 1 - p$ . Then this dual environment is Informative since,  $p' > 0.5$ . Thus if the learning automaton can somehow determine whether a given environment is Deceptive, then Lemma 4 and 5 assure us that by interchanging the actions (or equivalently the penalties and rewards), the automaton will still be able to converge to the optimal action with as high a probability as we want.

The following results are quite straightforward.

**Theorem 4.** *Given a Deceptive environment  $E$ ,*

$$\begin{aligned} & \text{If } (\lambda^* \odot \Delta^j), \text{ then } Pr(\Omega^j = \text{Right}) \rightarrow 1 \\ & \text{If } (\lambda^* \otimes \Delta^j), \text{ then } Pr(\Omega^j = \text{Left}) \rightarrow 1 \\ & \text{If } (\lambda^* \ominus \Delta^j), \text{ then } Pr(\Omega^j = \{\text{Left}, \text{Inside or Right}\}) \rightarrow 1. \end{aligned}$$

**Theorem 5.** *Suppose it is given that  $\lambda^* \odot \Delta^1$ ,  $\lambda^* \ominus \Delta^2$  and  $\lambda^* \otimes \Delta^3$ . Then under a Deceptive environment, the decision output vector  $\Omega$  for the three automata  $\mathcal{A}^j, j \in \{1, 2, 3\}$ , will be inconsistent with the decision table of Table 1. Conversely, if for the given environment and  $\lambda^*$  as above, the decision output vector  $\Omega$  of the automata is inconsistent with the decision table for large  $N_\infty$  and  $\theta \rightarrow 1$ , then the environment is Deceptive.*

Theorem 5 suggests a simple mechanism for determining whether or not an environment is Deceptive.

**Theorem 6.** *Let  $\mathcal{I} = [0, 1)$  be the original search interval in which  $\lambda^*$  is to be found. Let  $\mathcal{I}' = [-1, 2)$  be the initial search interval used by CPL-ATS. Then, CPL-ATS always determines whether or not an environment is Deceptive after a single epoch.*

Theorem 6 gives us a simple and practical mechanism for detecting Deceptive environments. Thus, we start with an initial search interval  $\mathcal{I}' = [-1, 2)$ , equipartition it and run the three  $L_{RI}$  automata  $\mathcal{A}^j, j = 1, 2, 3$ . for one epoch. At the end of the epoch, we use Table 1 to check if the decision vector  $\Omega$  has an entry. By appealing to Theorem 6, we conclude that the environment is Informative or Deceptive accordingly. If the environment was found to be Deceptive, we simply flip the probability update rules. i.e., we essentially treat every reward as a penalty and vice versa. Lemma 4 guarantees that we will then converge to the optimal action. If instead, the environment was found to be Informative we simply proceed with the search.

Note that the expanded interval  $\mathcal{I}'$  is needed only for the first epoch, to detect the environment's nature. Once this is detected, we use the original interval  $\mathcal{I}$  to search for  $\lambda^*$ . It is assumed that the fact that we use the expanded intervals  $[-1, 0)$  and  $[1, 2)$  does not affect the responses given by the environment.

## 9 Experimental Results

The parameter learning mechanism, CPL-ATS, described in this paper was experimentally evaluated to verify the validity of our analytic results and to examine its rate of convergence. To verify the power of the scheme and to study its effectiveness for various conditions, simulation experiments were conducted for various values of  $\theta$ , the reward factor of the  $L_{RI}$  automata, and for various values of  $p$ , the probability of the environment correctly providing the feed back. In all the experiments it was assumed that  $\lambda^* \in [0, 1)$ . In each case, a single screening epoch was run using the expanded interval  $[-1, 2)$  to detect whether or not the environment is Informative. After that, the given original search interval  $[0, 1)$  was used as the starting point. Each epoch consisted of 250 iterations ( $N_\infty$ ) of the three  $L_{RI}$  automata. At the end of each epoch the decision table Table 1 was consulted to prune the current search interval and the algorithm was recursively evoked. The recursion was terminated when the width of the interval was less than twice the desired accuracy.

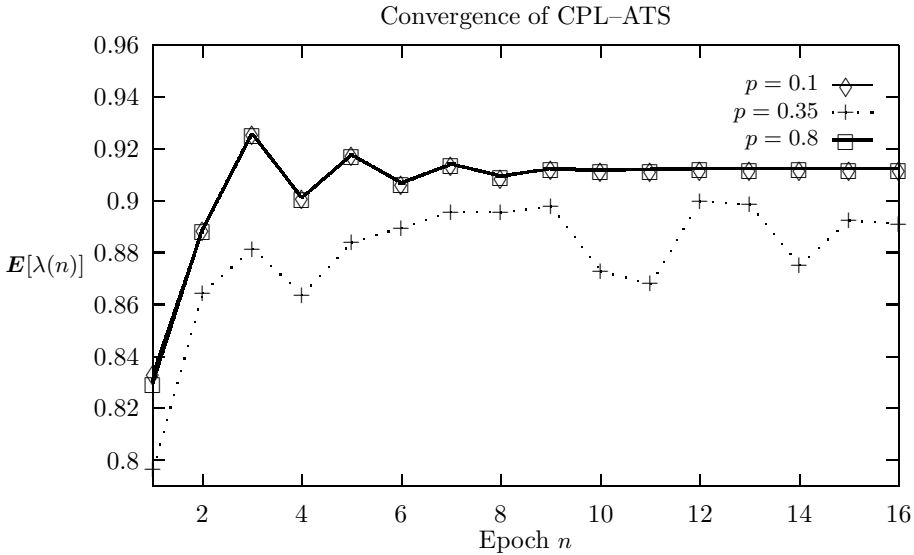
The results of our experiments are truly conclusive and confirm the power of the CPL-ATS scheme. Although several experiments were conducted using various  $\lambda^*$  and parameter values, we report for brevity sake, only one set of results, namely, those for  $\lambda^* = 0.9123$ . For this value, several independent replications with different random number streams were performed to minimize the variance of the reported results. The reported results are averaged over the replications.

The mean asymptotic value of the estimate for  $\lambda^*$  are shown in Table 2 for various values of  $p$  and  $\theta$ . The final estimates agree with the true value 0.9123 of  $\lambda^*$  to the first three decimal places for *all* values of  $p$  shown in the table. Figure 1 shows the convergence of the algorithm for  $p = 0.1, 0.35$  and  $0.8$ . This figure plots the running estimate at the end of each epoch. The values shown are actually the averaged over 50 replications for each set of parameters.

We first note that the convergence of the algorithm is almost identical for  $p = 0.1$  and  $p = 0.8$  even though the former represents a highly unreliable

**Table 2.** Asymptotic value of  $\hat{E}[\lambda(N_\infty)]$  as it varies with  $p$  and  $\theta$ . In all the cases,  $\lambda^* = 0.9123$ ,  $N_\infty = 250$  and  $\epsilon = 0.005$ . The values shown are averaged over 50 independent experiments

| $p$  | $\theta = 0.8$ | $\theta = 0.85$ | $\theta = 0.9$ |
|------|----------------|-----------------|----------------|
| 0.10 | 0.912298       | 0.912273        | 0.912194       |
| 0.15 | 0.912312       | 0.912298        | 0.912222       |
| 0.20 | 0.912193       | 0.912299        | 0.912236       |
| 0.80 | 0.912317       | 0.912284        | 0.912234       |
| 0.85 | 0.912299       | 0.912275        | 0.912202       |
| 0.90 | 0.912302       | 0.912275        | 0.912202       |



**Fig. 1.** Convergence of estimate  $\hat{E}[\lambda(n)]$  for  $\theta = 0.8$ . The value shown are averaged over 50 replications. The unknown parameter  $\lambda^* = 0.9123$

environment whereas the latter is a highly reliable environment. This is not surprising because, after having detected in epoch 0 that the environment is not Informative, the CPL-ATS strategy switches the reward and penalty feedbacks, effectively behaving as  $p' = 1 - p$ . Thus, even in the very first epoch a value of 0.83 is achieved which is within 9 percent error of the true value of  $\lambda^* = 0.9123$ . In two more epochs, for all the four  $p$  values shown, the estimate is 0.925926 which is within 1.5 percent error. Note however, that for  $p = 0.35$  the convergence is relatively sluggish. It took 25 epochs for it to reach the terminal value of 0.906453 which is within 0.64 percent of the true  $\lambda^*$ . Thus, as  $p$  is increasingly closer to 0.5,  $\theta$  must be set close to unity and  $N_\infty$  (the number of iterations

per epoch) has to be increased to achieve convergence.<sup>1</sup> Thus, a  $p$  value very close to 0.5 represents a highly inconsistent feedback wherein half the time the environment directs the learning mechanism in one direction and the rest of the time in another.

## 10 Conclusions

In this paper we have considered the problem of a learning mechanism locating a point on a line when it is interacting with a random stationary environment which essentially informs it, possibly erroneously, which way it should move to reach the point being sought. The first reported paper to solve this problem [9] presented a solution which operated in a discretized space. Later in [10], we presented a new scheme by which the point can be learned using a combination of various learning principles. The heart of this strategy involved performing a controlled random walk on the underlying space using the  $L_{RI}$  updating rule and then intelligently pruning the space using an adaptive tertiary search. The overall learning scheme is shown to be  $\epsilon$ -optimal. We have also shown that the results of [10] can also be generalized to the cases when the environment is a stochastic compulsive liar (i.e., is Deceptive). Although the problem is solved in its generality, its application in non-linear optimization has also been suggested. We are currently investigating how our scheme can be utilized to catalyze the convergence of various optimization problems including those involving neural networks as described in [9, 11, 19].

## Acknowledgements

The first author is grateful to Professor Riccardo Baeza-Yates for introducing him to this very interesting problem in Santiago, Chile. The first author was partially supported by the Natural Sciences and Engineering Research Council of Canada. The work of the second author was supported by the Canadian Commonwealth Scholarships and Fellowships Program.

## References

- [1] Baeza-Yates, R. A., Culberson, J. C. and Rawlins, G. J. E., "Searching with uncertainty", *Proc. Scandinavian Workshop on Algorithms and Theory*, SWAT 88, LNCS 318, 1988, pp. 176-189. 25
- [2] Baeza-Yates, R. A. and Schott, R., "Parallel Searching in the Plane", *Proc. 1992 International Conference of the Chilean Computer Society, IC-SCCC 92*, 1992, pp. 269-279. 25
- [3] Bentley, J. L. and Yao, A. C-C., "An Almost Optimal Algorithm for Unbounded Searching", *Information Processing Letters*, Aug. 1976, pp.82-87. 26

---

<sup>1</sup> For the parameter values that we experimented, we found that the algorithm consistently converges for  $p$  in the ranges (0,0.4) and (0.6,1.0).

- [4] Karlin, S., and Taylor, H.M., *A First Course on Stochastic Processes*, 2nd. ed., Academic Press, 1974.
- [5] Kashyap, R. L., and Oommen, B. J., "Scale Preserving Smoothing of Polygons", *IEEE Trans. on Pat. Anal. and Mach. Intel.*, Nov. 1983, pp. 667-671. [25](#)
- [6] Lakshmivarahan, S., *Learning Algorithms Theory and Applications*, Springer-Verlag, 1981. [30](#)
- [7] Lancôt, J. K. and Oommen, B. J., "Discretized Estimator Learning Automata", *IEEE Trans. on Syst. Man and Cybern.*, Vol. SMC-22, November/December 1992, pp. 1473- 1483.
- [8] Narendra, K. S., and Thathachar, M. A. L., *Learning Automata*, Prentice-Hall, 1989. [30](#)
- [9] Oommen, B. J., "Stochastic Searching on the Line and its Applications to Parameter Learning in Nonlinear Optimization", *IEEE Transactions on Systems, Man and Cybernetics*, Vol.SMC-27, 1997, pp.733-739. [25](#), [26](#), [35](#)
- [10] Oommen, B. J. and Raghunath, G., "Automata Learning and Intelligent Tertiary Searching for Stochastic Point Location", *IEEE Transactions on Systems, Man and Cybernetics*, Vol. SMC-28B, 1998, pp. 947-954. [26](#), [27](#), [35](#)
- [11] Pao, Y.-H., *Adaptive Pattern Recognition and Neural Networks*, Addison Wesley, Reading, Mass., 1989. [25](#), [35](#)
- [12] Pavlidis, T., *Structural Pattern Recognition*, Springer-Verlag, New York, 1977. [25](#)
- [13] Pelc, A., "Searching with Known Error Probability", *Theoretical Computer Science*, 63: pp. 185-202. [27](#)
- [14] Press, W. H., Flannery, B. P., Teukolsky, S. A., and Vetterling, W. T., *Numerical Recipes : The Art of Scientific Computing*, Cambridge University Press, Cambridge, 1986. [25](#)
- [15] Oommen, B. J., Raghunath, G., and Kuipers, B., "Parameter Learning from Stochastic Teachers and Stochastic Compulsive Liars", Technical Report in preparation, School of Computer Science, Carleton University, Ottawa. The report can be obtained by contacting the first author. [25](#), [36](#)
- [16] Rao, S. S., *Optimization :Theory and Applications*, John Wiley, New Delhi 2nd. Ed., 1984. [25](#)
- [17] Rich, E. and Knight, K. *Artificial Intelligence*. McGraw Hill Inc., 1991.
- [18] Tsetlin, M. L., *Automaton Theory and the Modelling of Biological Systems*, New York and London, Academic, 1973.
- [19] Wasserman, P. D., *Neural Computing : Theory and Practice*, Van Nostrand Reinhold, New York, 1989. [25](#), [35](#)

## Appendix

In the following we provide only sketches of proof for major lemmas and theorems due to limitations of space. The complete proofs for all the lemmas and theorems stated in this paper can be found in the unabridged version which is also available as a technical report [15].

**Theorem 1** *If the automaton  $\mathcal{A}^j = (\Sigma^j, \Pi^j, \Gamma^j, \Upsilon^j, \Omega^j)$  interacts with the environment  $E$  gets feedbacks obeying equation 1, then the effective penalty probabilities  $c_{k=0,1}^j$  for the two actions  $\alpha_{k=0,1}^j$  are given by:*

$$\begin{aligned}
 c_0^j &= (1 - p) + (2p - 1) \cdot \Pr(x_L^j < \lambda^* \mid \alpha_0^j \text{ was chosen}) \\
 c_1^j &= p - (2p - 1) \cdot \Pr(x_R^j < \lambda^* \mid \alpha_1^j \text{ was chosen})
 \end{aligned}$$

**Proof:** By definition of the penalty probability we have,

$$\begin{aligned}
c_0^j &= Pr(\beta(n) = 1 \mid \text{sub-interval } L^j \text{ is chosen at step } n) \\
&= Pr(\beta(n) = 1 \mid x_L^j \in L^j, x_L^j < \lambda^*) \cdot Pr(x_L^j < \lambda^* \mid x_L^j \in L^j) + \\
&\quad Pr(\beta(n) = 1 \mid x_L^j \in L^j, x_L^j \geq \lambda^*) \cdot Pr(x_L^j \geq \lambda^* \mid x_L^j \in L^j) \\
&= p \cdot Pr(x_L^j < \lambda^* \mid x_L^j \in L^j) + (1-p) \cdot (1 - Pr(x_L^j < \lambda^* \mid x_L^j \in L^j)) \quad (\text{by eq. 1}) \\
&= (1-p) + (2p-1) \cdot Pr(x_L^j < \lambda^* \mid x_L^j \in L^j) \\
&= (1-p) + (2p-1) \cdot Pr(x_L^j < \lambda^* \mid \alpha_0^j \text{ was chosen}).
\end{aligned}$$

In the same line of reasoning for the action  $\alpha_1^j$  we can derive  $c_1^j$ .  $\square$

**Lemma 1** *For an Informative environment  $E$ , given the  $L_{RI}$  scheme with a parameter  $\theta$  which is arbitrarily close to unity, the following is true:*

$$\begin{aligned}
&\text{If } (\lambda^* \ominus \Delta^j), \text{ then } Pr(\Omega^j = \text{Left}) \rightarrow 1. \\
&\text{If } (\lambda^* \otimes \Delta^j), \text{ then } Pr(\Omega^j = \text{Right}) \rightarrow 1. \\
&\text{If } (\lambda^* \ominus \Delta^j), \text{ then } Pr(\Omega^j = \{\text{Left}, \text{Inside or Right}\}) \rightarrow 1.
\end{aligned}$$

**Proof:** Consider first the case  $\lambda^* \otimes \Delta^j$ . From Equation (4) and (5), we get  $Pr(x_L^j < \lambda^*) = Pr(x_R^j < \lambda^*) = 0$ . Substituting these in equations 2 and 3 we get the values  $c_0^j = 1-p$  and  $c_1^j = p$ . Since for an informative environment  $p > 0.5$ , we immediately have  $c_0^j < c_1^j$ . Learning automata theory then assures that for any  $\epsilon$ -optimal scheme such as the  $L_{RI}$  scheme used here for each  $\mathcal{A}^j$ ,  $\alpha_0^j$  is the optimal action and hence  $P_0^j[N_\infty] \rightarrow 1$ . Similar arguments for  $\lambda^* \otimes \Delta^j$  lead to the conclusion  $P_1^j(N_\infty) \rightarrow 1$ .

Now consider the third case when  $\lambda^* \ominus \Delta^j$ . By the definition of  $\ominus$  we have,  $\sigma^j \leq \lambda^* < \gamma^j$ . Consider first the possibility,  $\sigma^j \leq \lambda^* < \text{mid}(\Delta^j)$ . In this case, by Equations (4), (5) we get

$$\begin{aligned}
c_0^j &= 1-p + (2p-1) \cdot \frac{\lambda^* - \sigma^j}{\text{mid}(\Delta^j) - \sigma^j} \\
c_1^j &= p
\end{aligned}$$

Since  $0.5 < p \leq 1$  in the above expressions for  $c_0^j$  and  $c_1^j$ , we can see that  $c_0^j < c_1^j$ , and hence  $\alpha_0^j$  is the optimal action and for an  $\epsilon$ -optimal scheme  $P_0^j[N_\infty] \rightarrow 1$ . Similarly when  $\text{mid}(\Delta^j) < \lambda^* < \sigma^j$ , we get,

$$\begin{aligned}
c_0^j &= p \\
c_1^j &= p - (2p-1) \cdot \frac{\lambda^* - \text{mid}(\Delta^j)}{\gamma^j - \text{mid}(\Delta^j)}.
\end{aligned}$$

Since  $p > 0.5$ , it follows that  $(2p-1) > 0$  and  $c_0^j > c_1^j$ , and therefore,  $\alpha_1^j$  is the optimal action, and from the  $\epsilon$ -optimality of the  $L_{RI}$  scheme, we get  $P_1^j[N_\infty] \rightarrow 1$ .

When  $\lambda^* = \text{mid}(\Delta^j)$ ,  $c_0^j = c_1^j$  and so there is no optimal action and hence  $\epsilon < P_0^j[N_\infty], P_1^j[N_\infty] < 1 - \epsilon$ .

The lemma immediately follows from the decision output rule ( $\Omega^j$ ) in the construction of the automaton  $\mathcal{A}^j$  based on the above values for  $P_k^j[N_\infty]$ .  $\square$

**Lemma 2** *For an Informative environment  $E$ , given the  $L_{RI}$  scheme with a parameter  $\theta$  which is arbitrarily close to unity, the following is true:*

$$\begin{aligned} \text{If } (\Omega^j = Left) \quad & \text{then } Pr(\lambda^* \otimes \Delta^j) \rightarrow 1 \\ \text{If } (\Omega^j = Right) \quad & \text{then } Pr(\lambda^* \otimes \Delta^j) \rightarrow 1 \\ \text{If } (\Omega^j = Inside) \quad & \text{then } Pr(\lambda^* \otimes \Delta^j) \rightarrow 1 \end{aligned}$$

**Proof:** By the first hypothesis  $\Omega^j = Left$ , we have  $Pr(\Omega^j = Left) = 1$ ,  $Pr(\Omega^j = Right) = 0$ , and  $Pr(\Omega^j = Inside) = 0$ . Therefore,

$$Pr(\lambda^* \otimes \Delta^j) = Pr(\lambda^* \otimes \Delta^j \mid \Omega^j = Left). \quad (6)$$

But by Baye's rule for conditional probabilities,

$$Pr(\lambda^* \otimes \Delta^j \mid \Omega^j = Left) = \quad (7)$$

$$\frac{Pr(\Omega^j = Left \mid \lambda^* \otimes \Delta^j).Pr(\lambda^* \otimes \Delta^j)}{Pr(\Omega^j = Left \mid \lambda^* \otimes \Delta^j).Pr(\lambda^* \otimes \Delta^j) + Pr(\Omega^j = Left \mid \lambda^* \otimes \Delta^j).Pr(\lambda^* \otimes \Delta^j)}$$

By Lemma 1, the product in the second term of the denominator tends to zero giving us

$$Pr(\lambda^* \otimes \Delta^j \mid \Omega^j = Left) \rightarrow 1.$$

When this is substituted in Equation (6), it leads to :  $Pr(\lambda^* \otimes \Delta^j) \rightarrow 1$ . By similar arguments (which we omit here in the interest of brevity), the lemma can be shown to hold for the other two hypotheses.  $\square$

**Theorem 2** *If the environment is Informative and if the partitions use the same  $L_{RI}$  scheme with parameters  $\theta$  as close to unity as needed, then the decision table given in Table 1 is complete.*

**Proof:** In any decision table with three input variables and three possible values for each of these variables, we expect a total of 27 potential entries. However, Lemma 3 imposes constraints on the output value combinations of the automata. A straightforward enumeration of these possibilities will show that the only possible combinations of outputs for the three  $L_{RI}$  automata are the seven entries shown in the decision table. Consequently Table 1 is complete.  $\square$

**Theorem 3** *If the algorithm uses the same  $L_{RI}$  scheme at all levels of the recursion with a parameter  $\theta$  arbitrarily close to unity and  $N_\infty$  sufficiently large, then for an Informative environment, the unknown  $\lambda^*$  is always contained in the new search-interval  $\Delta^{new}$  resulting from the application of the decision rules of Table 1.*

**Proof:** Consider the first row of Table 1 where we see that  $\Omega^1 = Left$ ,  $\Omega^2 = Left$  and  $\Omega^3 = Left$ . Appealing to Lemma 2 for each of the automata outputs,

we get  $Pr(\lambda^* \otimes \Delta^1) \rightarrow 1$ ,  $Pr(\lambda^* \otimes \Delta^2) \rightarrow 1$  and  $Pr(\lambda^* \otimes \Delta^3) \rightarrow 1$ . When we consider the fact that  $\Delta^1 \prec \Delta^2 \prec \Delta^3$ , the above three reduce to the equivalent predicate,  $Pr(\lambda^* \otimes \Delta^1) \rightarrow 1$ . By the definition of  $\otimes$  we have the two possibilities  $-\lambda^* \otimes \Delta^1$  and  $\lambda^* \otimes \Delta^1$ . But, since  $\lambda^* \otimes \Delta$  and  $\Delta = \Delta^1 \cup \Delta^2 \cup \Delta^3$ , we can rule out  $\lambda^* \otimes \Delta^1$ . Therefore,  $Pr(\lambda^* \otimes \Delta^1) \rightarrow 1$ . Thus, the partition  $\Delta^1$  that remains after pruning still contains  $\lambda^*$  with as high a probability as we want. The same arguments can be repeated for each of the other entries, and are omitted in the interest of brevity.  $\square$

**Theorem 4** Given a Deceptive environment E,

$$\begin{aligned} \text{If } (\lambda^* \otimes \Delta^j), \text{ then } Pr(\Omega^j = Right) &\rightarrow 1 \\ \text{If } (\lambda^* \oplus \Delta^j), \text{ then } Pr(\Omega^j = Left) &\rightarrow 1 \\ \text{If } (\lambda^* \ominus \Delta^j), \text{ then } Pr(\Omega^j = \{Left, Inside or Right\}) &\rightarrow 1. \end{aligned}$$

**Proof:** Since E is a Deceptive environment, the probability of it giving correct responses is  $p < 0.5$ . Now construct a dual  $E^*$  for this environment such that the probability of its correct feedback ( $p' = 1 - p$ )  $> 0.5$ . Clearly,  $E^*$  is an Informative environment, and hence we can apply Lemma 1 to it. Consequently, if we are given that  $\lambda^* \otimes \Delta^j$ , by Lemma 1, we know that for  $E^*$ ,  $Pr(\Omega^j = Left) \rightarrow 1$ . Lemma 4 however tells us that if  $\alpha_k^j$  was the optimal action for E, then  $\alpha_{1-k}^j$  is the optimal option for  $E^*$  and vice versa. Therefore, we conclude that for E,  $Pr(\Omega^j = Right) \rightarrow 1$ . Similar arguments hold good for the other two hypotheses.  $\square$

**Theorem 5** Suppose  $\lambda^*$  is such that  $\lambda^* \otimes \Delta^1$ ,  $\lambda^* \ominus \Delta^2$  and  $\lambda^* \otimes \Delta^3$ . Then under a Deceptive environment, the decision output vector  $\Omega$  for the three automata  $A^j, j = 1, 2, 3$ , will be inconsistent with the decision table of Table 1. Conversely, if for the given environment the decision output vector  $\Omega$  of the automata is inconsistent with the decision table for large  $N_\infty$  and  $\theta \rightarrow 1$ , then the environment is Deceptive.

**Proof:** Applying Theorem 4 to each of  $\lambda^* \otimes \Delta^1$ ,  $\lambda^* \ominus \Delta^2$  and  $\lambda^* \otimes \Delta^3$ , we have,

$$\begin{aligned} Pr(\Omega^1 = Left) &\rightarrow 1 \\ Pr(\Omega^2 = \{Left, Inside or Right\}) &\rightarrow 1 \\ Pr(\Omega^3 = Right) &\rightarrow 1 \end{aligned}$$

By inspecting the decision table (Table 1), we see that there are no entries for this output vector  $\Omega = [Left, \{Left, Inside or Right\}, Right]$  and hence the entry is inconsistent with Table 1. The converse is proved by contradiction by alluding to the completeness result of Theorem 2 for an Informative environment. Therefore, whenever the decision output vector  $\Omega$  is inconsistent with Table 1, we can safely conclude that the environment is Deceptive.  $\square$

**Theorem 6** Let  $\mathcal{I} = [0, 1)$  be the original search interval in which  $\lambda^*$  is to be found. Let  $\mathcal{I}' = [-1, 2)$  be the initial search interval used by CPL-ATS. Then,

*CPL-ATS always determines whether or not an environment is Deceptive after a single epoch.*

**Proof:** First of all, we can see that  $\lambda^* \in \mathcal{I}'$  because,  $\mathcal{I} \subset \mathcal{I}'$ . When we divide  $\mathcal{I}'$  into three equi-partitions we get,  $\Delta^1 = [-1, 0)$ ,  $\Delta^2 = [0, 1)$  and  $\Delta^3 = [1, 2)$ . Since  $\lambda^* \in \mathcal{I} = \Delta^2$ , we have,  $\lambda^* \ominus \Delta^1$ ,  $\lambda^* \ominus \Delta^2$ ,  $\lambda^* \ominus \Delta^3$ , which is the pre-condition for Theorem 5. Hence, by appealing to Theorem 5 we see that if the environment was Deceptive, we would get an inconsistent decision vector. If not, by Theorem 2 we would get a consistent decision vector. Thus, after *one* epoch we conclude decisively about the nature of the environment.  $\square$