

DTIC

4

LABORATORY FOR
COMPUTER SCIENCE



MASSACHUSETTS
INSTITUTE OF
TECHNOLOGY

AD-A216 389

MIT/LCS/TM-415

THE STRENGTH OF WEAK LEARNABILITY

Robert E. Schapire

DTIC
ELECTE
JAN 3 1990
S B D

October 1989

545 TECHNOLOGY SQUARE, CAMBRIDGE, MASSACHUSETTS 02139

DISTRIBUTION STATEMENT A

Approved for public release;
Distribution Unlimited

90 01 03 048

REPORT DOCUMENTATION PAGE

1a. REPORT SECURITY CLASSIFICATION Unclassified			1b. RESTRICTIVE MARKINGS		
2a. SECURITY CLASSIFICATION AUTHORITY			3. DISTRIBUTION/AVAILABILITY OF REPORT Approved for public release; distribution is unlimited.		
2b. DECLASSIFICATION/DOWNGRADING SCHEDULE					
4. PERFORMING ORGANIZATION REPORT NUMBER(S) MIT/LCS/TM-415			5. MONITORING ORGANIZATION REPORT NUMBER(S) N00014-89-J-1988		
6a. NAME OF PERFORMING ORGANIZATION MIT Lab for Computer Science		6b. OFFICE SYMBOL (If applicable)		7a. NAME OF MONITORING ORGANIZATION Office of Naval Research/Dept. of Navy	
6c. ADDRESS (City, State, and ZIP Code) 545 Technology Square Cambridge, MA 02139			7b. ADDRESS (City, State, and ZIP Code) Information Systems Program Arlington, VA 22217		
8a. NAME OF FUNDING/SPONSORING ORGANIZATION DARPA/DOD		8b. OFFICE SYMBOL (If applicable)		9. PROCUREMENT INSTRUMENT IDENTIFICATION NUMBER	
8c. ADDRESS (City, State, and ZIP Code) 1400 Wilson Blvd. Arlington, VA 22217			10. SOURCE OF FUNDING NUMBERS		
			PROGRAM ELEMENT NO.	PROJECT NO.	TASK NO.
			WORK UNIT ACCESSION NO.		
11. TITLE (Include Security Classification) The Strength of Weak Learnability					
12. PERSONAL AUTHOR(S) Schapire, Robert E.					
13a. TYPE OF REPORT Technical		13b. TIME COVERED FROM _____ TO _____		14. DATE OF REPORT (Year, Month, Day) October 1989	
15. PAGE COUNT 34					
16. SUPPLEMENTARY NOTATION					
17. COSATI CODES			18. SUBJECT TERMS (Continue on reverse if necessary and identify by block number)		
FIELD	GROUP	SUB-GROUP	Machine learning, learning from examples,		
			polynomial-time identification.		
19. ABSTRACT (Continue on reverse if necessary and identify by block number)					
The problem considered is that of improving the accuracy of an hypothesis output by a learning algorithm in the distribution-free(PAC) learning model. A concept class is learnable (or strongly learnable) if, given access to a source of examples from the unknown concept, the learner with high probability is able to output an hypothesis that is correct on all but an arbitrarily small fraction of the instances. The concept class is weakly learnable if the learner can produce an hypothesis that performs only slightly better than random guessing. In this paper, it is shown that these two notions of learnability are equivalent. A method is described for converting a weak learning algorithm into one that achieves arbitrarily high accuracy. This construction may have practical applications as a tool for efficiently converting a mediocre learning algorithm into one that performs extremely well. In addition, the construction has some interesting theoretical consequences, including a set of general upper bounds on the complexity of any strong learning algorithm as a function of the allowed error ϵ.					
20. DISTRIBUTION/AVAILABILITY OF ABSTRACT ☒ UNCLASSIFIED/UNLIMITED ☐ SAME AS RPT. ☐ DTIC USERS			21. ABSTRACT SECURITY CLASSIFICATION Unclassified		
22a. NAME OF RESPONSIBLE INDIVIDUAL Judy Little			22b. TELEPHONE (Include Area Code) (617) 253-5894		22c. OFFICE SYMBOL

The Strength of Weak Learnability

Robert E. Schapire

MIT Laboratory for Computer Science
Cambridge, MA 02139

October 16, 1989

Abstract

The problem considered is that of improving the accuracy of an hypothesis output by a learning algorithm in the distribution-free (PAC) learning model. A concept class is *learnable* (or *strongly learnable*) if, given access to a source of examples from the unknown concept, the learner with high probability is able to output an hypothesis that is correct on all but an arbitrarily small fraction of the instances. The concept class is *weakly learnable* if the learner can produce an hypothesis that performs only slightly better than random guessing. In this paper, it is shown that these two notions of learnability are equivalent.

A method is described for converting a weak learning algorithm into one that achieves arbitrarily high accuracy. This construction may have practical applications as a tool for efficiently converting a mediocre learning algorithm into one that performs extremely well. In addition, the construction has some interesting theoretical consequences, including a set of general upper bounds on the complexity of any strong learning algorithm as a function of the allowed error ϵ .

Keywords: Machine learning, learning from examples, polynomial-time identification.

1 Introduction

Since Valiant's pioneering paper [23], interest has flourished in the so-called distribution-free or probably approximately correct (PAC) model of learning. In this model, the learner tries to identify an unknown concept based on randomly chosen examples of the concept. Examples are chosen according to a fixed but unknown and arbitrary distribution on the

This paper prepared with support from ARO Grant DAAL03-86-K-0171, DARPA Contract N00014-89-J-1988, and a grant from the Siemens Corporation.

Author's net address: `rs@theory.lcs.mit.edu`.

space of instances. The learner's task is to find an hypothesis or prediction rule of his own that correctly classifies new instances as positive or negative examples of the concept. With high probability, the hypothesis must be correct for all but an arbitrarily small fraction of the instances.

Often, the inference task includes a requirement that the output hypothesis be of a specified form. In this paper, however, we will instead be concerned with a representation-independent model of learning in which the learner may output any hypothesis that classifies instances in polynomial time.

A class of concepts is *learnable* (or *strongly learnable*) if there exists a polynomial-time algorithm that achieves low error with high confidence for all concepts in the class. A weaker model of learnability, called *weak learnability*, drops the requirement that the learner be able to achieve arbitrarily high accuracy; a weak learning algorithm need only output an hypothesis that performs slightly better (by an inverse polynomial) than random guessing. The notion of weak learnability was introduced by Kearns and Valiant [17, 18] who left open the question of whether the notions of strong and weak learning are equivalent. This question was termed the *hypothesis boosting problem* since showing the notions are equivalent requires a method for boosting the low accuracy of a weak learning algorithm's hypotheses.

Kearns [15], considering the hypothesis boosting problem, gives a convincing argument discrediting the natural approach of trying to boost the accuracy of a weak learning algorithm by running the procedure many times and taking "majority vote" of the output hypotheses. Kearns and Valiant [14, 17] show that, under a uniform distribution on the instance space, monotone Boolean functions are weakly, but not strongly, learnable. This implies that strong and weak learnability are *not* equivalent when certain restrictions are placed on the instance space distribution. Thus, it did not seem implausible that the strong and weak learning models would prove to be inequivalent for unrestricted distributions as well.

Nevertheless, in this paper, the hypothesis boosting question is answered in the affirmative. The main result is a proof of the perhaps surprising equivalence of strong and weak learnability.

This result may have significant applications as a tool for proving that a concept class is learnable since, in the future, it will suffice to find an algorithm correct on only, say, 51% of the instances (for all distributions). Alternatively, in its negative contrapositive form, the result says that, if a concept class cannot be learned with accuracy 99.9%, then we cannot

hope to do even slightly better than guessing on the class (for some distribution).

The proof presented here is constructive; an explicit method is described for directly converting a weak learning algorithm into one that achieves arbitrary accuracy. The construction uses *filtering* to modify the distribution of examples in such a way as to force the weak learning algorithm to focus on the harder-to-learn parts of the distribution. Thus, the distribution-free nature of the learning model is fully exploited.

An immediate corollary of the main result is the equivalence of strong and *group* learnability. A group-learning algorithm need only output an hypothesis capable of classifying large groups of instances, all of which are either positive or negative. The notion of group learnability was considered by Kearns et al. [16], and was shown to be equivalent to weak learnability by Kearns and Valiant [14, 17]. The result also extends those of Haussler et al. [10] that prove the equivalence of numerous variations and relaxations on the basic PAC-learning model; both weak and group learnability are added to this general class of equivalent learning models. The relevance of the main result to a number of other learning models is also considered in this paper.

An interesting and unexpected consequence of the construction is a proof that any strong learning algorithm outputting hypotheses whose length (and thus whose time to evaluate) depends on the allowed error ϵ can be modified to output hypotheses of length only polynomial in $\log(1/\epsilon)$. Thus, any learning algorithm can be converted into one whose output hypotheses do not become significantly more complex as the error tolerance is lowered.

This bound on the size of the output hypothesis implies the hardness of learning any concept class not evaluable by a family of small circuits. For example, this shows that pattern languages — a class of languages considered previously by Angluin [1] and others — are unlearnable assuming only that $NP/poly \neq P/poly$. This is the first representation-independent hardness result not based on cryptographic assumptions. The bound also shows that, for any function not computable by polynomial-size circuits, there exists a distribution on the function's domain over which the function cannot be even roughly approximated by a family of small circuits.

In addition to the bound on hypothesis size, the construction implies a set of general bounds on the dependence on ϵ of the time, sample and space complexity needed to efficiently learn any learnable concept class. Most surprising is a proof that there exists for every learnable concept class an efficient algorithm requiring space only poly-logarithmic in $1/\epsilon$.



Dist

A-1

Avail and/or
Special

Codes

Because the size of the sample needed to learn with this accuracy is in general $\Omega(1/\epsilon)$, this means, for example, that far less space is required to learn than would be necessary to store the entire sample. Since most of the known learning algorithms work in exactly this manner — i.e., by storing a large sample and finding an hypothesis consistent with it — this implies a dramatic savings of memory for a whole class of algorithms (though possibly at the cost of requiring a larger sample).

Such general complexity bounds have implications for the *on-line* learning model as well. In this model, the learner is presented one instance at a time in a series of trials. As each is received, the learner tries to predict the true classification of the new instance, attempting to minimize the number of mistakes, or prediction errors.

Translating the bounds described above into the on-line model, it is shown that, for every learnable concept class, there exists an on-line algorithm whose space requirements are quite modest in comparison to the number of examples seen so far. In particular, the space needed on the first m trials is only poly-logarithmic in m . Such space efficient on-line algorithms are of particular interest because they capture the notion of an incremental algorithm forced by its limited memory to explicitly generalize or abstract from the data observed. Also, these results on the space-efficiency of batch and on-line algorithms extend the work of others interested in this problem, including Boucheron and Sallantin [6], Floyd [8], and Haussler [9]. In particular, these results solve an open problem proposed by Haussler, Littlestone and Warmuth [11].

An interesting bound is also derived on the expected number of mistakes made on the first m trials. It is shown that, if a concept class is learnable, then there exists an on-line algorithm for the class for which this expectation is bounded by a polynomial in $\log m$. Thus, for large m , we expect an extremely small fraction of the first m predictions to be incorrect. This result answers another open question given by Haussler, Littlestone and Warmuth [11], and significantly improves a similar bound given in their paper (as well as their paper with Kearns [10]) of m^α for some constant $\alpha < 1$.

2 Preliminaries

We begin with a description of the distribution-free learning model. A *concept* c is a Boolean function on some domain of *instances*. A *concept class* C is a collection of concepts. Often,

$\mathcal{C}$ is decomposed into subclasses $\mathcal{C}_n$ indexed by a parameter n . That is, $\mathcal{C} = \bigcup_{n \geq 1} \mathcal{C}_n$, and all the concepts in $\mathcal{C}_n$ have a common domain X_n . We assume each instance in X_n has encoded length bounded by a polynomial in n , and we let $X = \bigcup_{n \geq 1} X_n$. Also, we associate with each concept c its *size* s , typically a measure of the length of c under some encoding scheme on the concepts in $\mathcal{C}$.

The learner is assumed to have access to a source EX of examples. Each time EX is called, one instance is randomly and independently chosen from X_n according to some fixed but unknown and arbitrary distribution D . The oracle returns the chosen instance v , along with a label indicating the value $c(v)$ of the instance under the unknown *target concept* $c \in \mathcal{C}_n$. Such a labeled instance is called an *example*. We assume EX runs in unit time.

Given access to EX , the learning algorithm runs for a time and finally outputs an *hypothesis* h , a prediction rule on X_n . In this paper, we make no restrictions on h other than that there exist a (possibly probabilistic) polynomial time algorithm that, given h and an instance v , computes $h(v)$, h 's prediction on v .

We write $\Pr_{v \in D}[\pi(v)]$ to indicate the probability of predicate π holding on instances v drawn from X_n according to distribution D . To accommodate probabilistic hypotheses, we will find it useful to regard $\pi(v)$ as a Bernoulli random variable. For example, $\Pr[h(v) \neq c(v)]$ is the chance that hypothesis h (which may be randomized) will misclassify some particular instance v . In contrast, the quantity $\Pr_{v \in D}[h(v) \neq c(v)]$ is the probability that h will misclassify an instance v chosen at random according to distribution D . Note that this last probability is taken over both the random choice of v , and any random bits used by h .

In general, assuming independence, we have

$$\Pr_{v \in D}[\pi(v)] = \sum_{v \in X_n} D(v) \Pr[\pi(v)]$$

where $D(v)$ is the probability of instance v being chosen under D . (Technically, this formula is valid only when X_n is discrete. If X_n is continuous, then the summation would need to be replaced by the appropriate integral, and D by a probability density function. To simplify the presentation, we will assume that X_n is discrete, and omit the extension of these results to continuous domains.)

The probability $\Pr_{v \in D}[h(v) \neq c(v)]$ is called the *error* of h on c under D ; if the error is no more than ϵ , then we say h is ϵ -close to the target concept c under D . The quantity $\Pr_{v \in D}[h(v) = c(v)]$ is the *accuracy* of h on c under D .

We say that a concept class $\mathcal{C}$ is *learnable*, or *strongly learnable*, if there exists an algorithm A such that for all $n \geq 1$, for all target concepts $c \in \mathcal{C}_n$, for all distributions D on X_n , and for all $0 < \epsilon, \delta \leq 1$, algorithm A , given parameters n, ϵ, δ , the size s of c , and access to oracle EX , runs in time polynomial in $n, s, 1/\epsilon$ and $1/\delta$, and outputs an hypothesis h that with probability at least $1 - \delta$ is ϵ -close to c under D . There are many other equivalent notions of learnability, including polynomial predictability [10].

Kearns and Valiant [17, 18] introduced a weaker form of learnability in which the error ϵ cannot necessarily be made arbitrarily small. A concept class $\mathcal{C}$ is *weakly learnable* if there exists a polynomial p and an algorithm A such that for all $n \geq 1$, for all target concepts $c \in \mathcal{C}_n$, for all distributions D on X_n , and for all $0 < \delta \leq 1$, algorithm A , given parameters n, δ , the size s of c , and access to oracle EX , runs in time polynomial in n, s and $1/\delta$, and outputs an hypothesis h that with probability at least $1 - \delta$ is $(\frac{1}{2} - \frac{1}{p(n,s)})$ -close to c under D . In other words, a weak learning algorithm produces a prediction rule that performs just slightly better than random guessing.

3 The Equivalence of Strong and Weak Learnability

The main result of this paper is a proof that learnability and weak learnability are equivalent notions.

Theorem 3.1 *A concept class $\mathcal{C}$ is weakly learnable if and only if it is learnable.*

That strong learnability implies weak learnability is trivial. The remainder of this section is devoted to a proof of the converse. We assume then that some concept class $\mathcal{C}$ is weakly learnable and show how to build a strong learning algorithm around a weak one.

We begin with a description of a technique by which the accuracy of any algorithm can be boosted by a small but significant amount. Later, we will show how this mechanism can be applied recursively to make the error arbitrarily small.

3.1 The Hypothesis Boosting Mechanism

Let A be an algorithm that produces with high probability an hypothesis α -close to the target concept c . We sketch an algorithm A' that simulates A on three different distributions, and outputs an hypothesis significantly closer to c .

Let EX be the given examples oracle, and let D be the distribution on X_n induced by EX . The algorithm A' begins by simulating A on the original distribution $D_1 = D$, using the given oracle $EX_1 = EX$. Let h_1 be the hypothesis output by A .

Intuitively, A has found some weak advantage on the original distribution; this advantage is expressed by h_1 . To force A to learn more about the “harder” parts of the distribution, we must somehow destroy this advantage. To do so, A' creates a new distribution D_2 under which an instance chosen according to D_2 has a roughly equal chance of being correctly or incorrectly classified by h_1 . The distribution D_2 is simulated by filtering the examples chosen according to D by EX . To simulate D_2 , a new examples oracle EX_2 is constructed. When asked for an instance, EX_2 first flips a fair coin: if the result is “heads,” then EX_2 requests examples from EX until one is chosen for which $h_1(v) = c(v)$; otherwise, EX_2 waits for an instance to be chosen for which $h_1(v) \neq c(v)$. (Later we show how to prevent EX_2 from having to wait too long in either of these loops for a desired instance.) The algorithm A is again simulated, this time providing A with examples chosen by EX_2 according to D_2 . Let h_2 be the output hypothesis.

Finally, D_3 is constructed by filtering from D those instances on which h_1 and h_2 agree. That is, a third oracle EX_3 simulates the choice of an instance according to D_3 by requesting instances from EX until one is found for which $h_1(v) \neq h_2(v)$. (Again, we will later show how to limit the time spent waiting in this loop for a desired instance.) For a third time, algorithm A is simulated with examples drawn this time by EX_3 , producing hypothesis h_3 .

At last, A' outputs its hypothesis h : given an instance v , if $h_1(v) = h_2(v)$ then h predicts the agreed upon value; otherwise, h predicts $h_3(v)$. (In other words, h takes “majority vote” of h_1 , h_2 and h_3 .) Later, we show that h ’s error is bounded by $g(\alpha) \equiv 3\alpha^2 - 2\alpha^3$. This quantity is significantly smaller than the original error α , as can be seen from its graph depicted in Figure 1. (The solid curve is the function g , and, for comparison, the dotted line shows a graph of the identity function.)

3.2 A Strong Learning Algorithm

An idea that follows naturally is to treat the previously described procedure as a subroutine for recursively boosting the accuracy of weaker hypotheses. The procedure is given a desired error bound ϵ and a confidence parameter δ , and constructs an ϵ -close hypothesis from weaker, recursively computed hypotheses. If $\epsilon \geq \frac{1}{2} - \frac{1}{p(n,\delta)}$ then an assumed weak learning

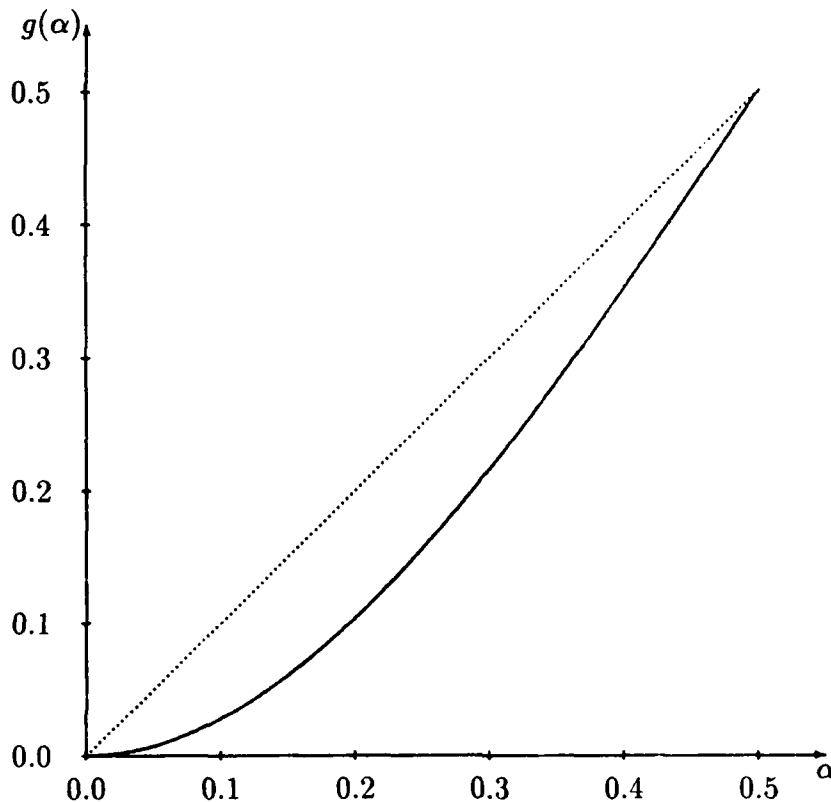


Figure 1: A graph of the function $g(\alpha) = 3\alpha^2 - 2\alpha^3$.

algorithm can be used to find the desired hypothesis; otherwise, an ϵ -close hypothesis is computed recursively by calling the subroutine with ϵ set to $g^{-1}(\epsilon)$.

Unfortunately, this scheme by itself does not quite work due to a technical difficulty: because of the way EX_2 and EX_3 are constructed, examples may be required from a very small portion of the original distribution. If this happens, the time spent waiting for an example to be chosen from this region may be great. Nevertheless, we will see that this difficulty can be overcome by explicitly checking that the errors of hypotheses h_1 and h_2 on D are not too small.

Figure 2 shows a detailed sketch of the resulting strong learning algorithm **Learn**. The procedure takes an error parameter ϵ and a confidence parameter δ , and is also provided with an examples oracle EX . The procedure is required to return an hypothesis whose error is at

Input: error parameter ϵ
confidence parameter δ
examples oracle EX
(implicit) size parameters s and n

Return: hypothesis h that is ϵ -close to the target concept c with probability $\geq 1 - \delta$

Procedure:

```

if  $\epsilon \geq \frac{1}{2} - \frac{1}{p(n,s)}$  then return WeakLearn( $\delta, EX$ )
 $\alpha \leftarrow g^{-1}(\epsilon)$ 
 $EX_1 \leftarrow EX$ 
 $h_1 \leftarrow \text{Learn}(\alpha, \frac{1}{5}\delta, EX_1)$ 
 $\tau_1 \leftarrow \frac{1}{3}\epsilon$ 
let  $\hat{a}_1$  be an estimate of  $a_1 = \Pr_{v \in D}[h_1(v) \neq c(v)]$ :
    choose a sample sufficiently large that  $|a_1 - \hat{a}_1| \leq \tau_1$  with probability  $\geq 1 - \frac{1}{5}\delta$ 
if  $\hat{a}_1 \leq \epsilon - \tau_1$  then return  $h_1$ 
defun  $EX_2()$ 
    { flip coin
      if "heads," return the first instance  $v$  from  $EX$  for which  $h_1(v) = c(v)$ 
      else return the first instance  $v$  from  $EX$  for which  $h_1(v) \neq c(v)$  }
 $h_2 \leftarrow \text{Learn}(\alpha, \frac{1}{5}\delta, EX_2)$ 
 $\tau_2 \leftarrow \frac{1}{8}(1 - 2\alpha)\epsilon$ 
let  $\hat{e}$  be an estimate of  $e = \Pr_{v \in D}[h_2(v) \neq c(v)]$ :
    choose a sample sufficiently large that  $|e - \hat{e}| \leq \tau_2$  with probability  $\geq 1 - \frac{1}{5}\delta$ 
if  $\hat{e} \leq \epsilon - \tau_2$  then return  $h_2$ 
defun  $EX_3()$ 
    { return the first instance  $v$  from  $EX$  for which  $h_1(v) \neq h_2(v)$  }
 $h_3 \leftarrow \text{Learn}(\alpha, \frac{1}{5}\delta, EX_3)$ 
defun  $h(v)$ 
    {  $b_1 \leftarrow h_1(v), b_2 \leftarrow h_2(v)$ 
      if  $b_1 = b_2$  then return  $b_1$ 
      else return  $h_3(v)$  }
return  $h$ 

```

Figure 2: A strong learning algorithm **Learn**.

most ϵ with probability at least $1 - \delta$. In the figure, p is a polynomial and **WeakLearn**(δ, EX) is an assumed weak learning procedure that outputs an hypothesis $(\frac{1}{2} - \frac{1}{p(n,s)})$ -close to the target concept c with probability at least $1 - \delta$. As above, $g(\alpha)$ is the function $3\alpha^2 - 2\alpha^3$, and the variable α is set to the value $g^{-1}(\epsilon)$. Also, the quantities $\hat{a}_1$ and $\hat{e}$ are estimates of the errors of h_1 and h_2 under the given distribution D . These estimates are made with error tolerances τ_1 and τ_2 (defined in the figure), and are computed in the obvious manner

based on samples drawn from EX ; the required size of these samples can be determined, for instance, using Chernoff bounds [22]. The parameters s and n are assumed to be known globally.

Note that **Learn** is a procedure taking as one of its inputs a function (EX) and returning as output another function (h , a hypothesis, which is treated like a procedure). Furthermore, to simulate new example oracles, **Learn** must have a means of dynamically defining new procedures (as is allowed, for instance, by most Lisp-like languages). Therefore, in the figure, we have used the somewhat non-standard keyword **defun** to denote the definition of a new function; its syntax calls for a name for the procedure, followed by a parenthesized list of arguments, and the body indented in braces. Static scoping is assumed.

Learn works by recursively boosting the accuracy of its hypotheses. **Learn** typically calls itself three times using the three simulated example oracles described in the preceding section. On each recursive call, the required error bound of the constructed hypotheses comes closer to $\frac{1}{2}$; when this bound reaches $\frac{1}{2} - \frac{1}{p(n,s)}$, the weak learning algorithm **WeakLearn** can be used.

The procedure takes measures to limit the run time of the simulated oracles it provides on recursive calls. When **Learn** calls itself a second time to find h_2 , the expected number of iterations of EX_2 to find an example depends on the error of h_1 , which is estimated by $\hat{a}_1$. If h_1 already has the desired accuracy $1 - \epsilon$, then there is no need to find h_2 and h_3 since h_1 is a sufficiently good hypothesis; otherwise, if $a_1 = \Omega(\epsilon)$, then it can be shown that EX_2 will not loop too long to find an instance. Similarly, when **Learn** calls itself to find h_3 , the expected number of iterations of EX_3 depends on how often h_1 and h_2 disagree, which we will see is in turn a function of the error of h_2 on the original distribution D . If this error e (which is estimated by $\hat{e}$) is small, then h_2 is a good hypothesis and is returned by **Learn**. Otherwise, it will be shown that EX_3 also will not run for too long.

3.3 Correctness

We show in this section that the algorithm is correct in the following sense:

Theorem 3.2 *For $0 < \epsilon < \frac{1}{2}$ and for $0 < \delta \leq 1$, the hypothesis returned by calling **Learn**(ϵ, δ, EX) is ϵ -close to the target concept with probability at least $1 - \delta$.*

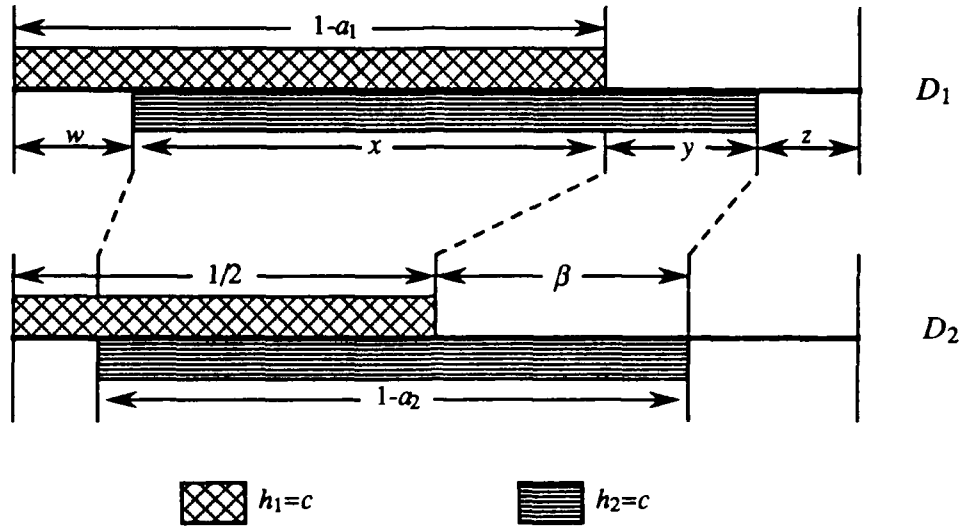


Figure 3: The distributions D_1 and D_2 .

Proof: Proof is by induction starting at the bottom of the recursion. (Technically, the induction is on $B(\epsilon, p(n, s))$, where B is a function defined in the next section.) The base case that $\epsilon \geq \frac{1}{2} - \frac{1}{p(n, s)}$ follows trivially from our assumptions about **WeakLearn**. Another easy case is that $\hat{a}_1$ or $\hat{e}$ is found to be smaller than $\epsilon - \tau_1$ or $\epsilon - \tau_2$, respectively. In either case, it follows immediately, due to the accuracy with which a_1 and e have been estimated, that the returned hypothesis is ϵ -close to the target concept.

Otherwise, all three hypotheses must be found and combined. Let a_i be the error of h_i under D_i . Here, D is the distribution of the provided oracle EX , and D_i is the distribution induced by oracle EX_i on the i th recursive call ($i = 1, 2, 3$). By inductive hypothesis, $a_i \leq \alpha$ with probability $1 - \frac{1}{5}\delta$.

In the special case that all hypotheses are deterministic, the distributions D_1 and D_2 can be depicted schematically as shown in Figure 3. The figure shows the portion of each distribution for which the hypotheses h_1 and h_2 agree with the target concept c . For each distribution, the top crosshatched bar represents the relative fraction of the instance space for which h_1 agrees with c ; the bottom striped bar represents those instances for which h_2 agrees with c . Although only valid for deterministic hypotheses, this figure may be helpful for motivating one's intuition in what follows.

Let $p_i(v) = \Pr[h_i(v) \neq c(v)]$ be the chance that some fixed instance v is misclassified by h_i . (Recall that hypotheses may be randomized, and therefore it is necessary to consider the *probability* that a particular fixed instance is misclassified.) Similarly, let $q(v) = \Pr[h_1(v) \neq h_2(v)]$ be the chance that v is classified differently by h_1 and h_2 . Also define w , x , y , and z as follows:

$$\begin{aligned} w &= \Pr_{v \in D}[h_2(v) \neq h_1(v) = c(v)] \\ x &= \Pr_{v \in D}[h_1(v) = h_2(v) = c(v)] \\ y &= \Pr_{v \in D}[h_1(v) \neq h_2(v) = c(v)] \\ z &= \Pr_{v \in D}[h_1(v) = h_2(v) \neq c(v)] \end{aligned}$$

Clearly,

$$w + x = \Pr_{v \in D}[h_1(v) = c(v)] = 1 - a_1, \quad (1)$$

and since c , h_1 and h_2 are Boolean,

$$y + z = \Pr_{v \in D}[h_1(v) \neq c(v)] = a_1. \quad (2)$$

In terms of these variables, we can express explicitly the chance that EX_i returns instance v :

$$D_1(v) = D(v) \quad (3)$$

$$D_2(v) = \frac{D(v)}{2} \left(\frac{p_1(v)}{a_1} + \frac{1 - p_1(v)}{1 - a_1} \right) \quad (4)$$

$$D_3(v) = \frac{D(v)q(v)}{w + y} \quad (5)$$

Equation (3) is trivial. To see that (4) holds, note that the chance that the initial coin flip comes up "tails" is $\frac{1}{2}$, and the chance that instance v is the first instance misclassified by h_1 is $D(v)p_1(v)/a_1$. The case that the coin comes up "heads" is handled in a similar fashion, as is the derivation of equation (5).

From equation (4), we have that

$$\begin{aligned} 1 - a_2 &= \sum_{v \in X_n} D_2(v)(1 - p_2(v)) \\ &= \frac{1}{2a_1} \sum_{v \in X_n} D(v)p_1(v)(1 - p_2(v)) + \frac{1}{2(1 - a_1)} \sum_{v \in X_n} D(v)(1 - p_1(v))(1 - p_2(v)) \\ &= \frac{y}{2a_1} + \frac{x}{2(1 - a_1)} \end{aligned} \quad (6)$$

Combining equations (1), (2) and (6), we see that the values of w , x and z can be written explicitly in terms of y , a_1 and a_2 . (Note that (6) could also have been derived from Figure 3 in the case of deterministic hypotheses: if β is as shown in the figure, then it is not hard to see that $y = 2a_1\beta$ and $x = 2(1 - a_1)(1 - a_2 - \beta)$. These imply (6).)

Finally, using equation (5), we are ready to compute the error of the output hypothesis h :

$$\begin{aligned}
\Pr_{v \in D}[h(v) \neq c(v)] &= \Pr_{v \in D}[(h_1(v) = h_2(v) \neq c(v)) \vee (h_1(v) \neq h_2(v) \wedge h_3(v) \neq c(v))] \\
&= z + \sum_{v \in X_n} D(v)q(v)p_3(v) \\
&= z + \sum_{v \in X_n} (w + y)D_3(v)p_3(v) \\
&= z + a_3(w + y) \\
&= a_1 - a_3 + a_1a_3 + 2a_2a_3 - 2a_1a_2a_3 + \frac{y}{a_1}(a_3 - a_1) \\
&\leq 3\alpha^2 - 2\alpha^3 = g(\alpha) = \epsilon
\end{aligned}$$

as desired. The inequality here follows with some care from the facts that each $a_i \leq \alpha$, and that $y \leq a_1$.

Finally, note that the confidence parameter δ has been “spread around” so that the overall chance of anything “going wrong” is at most δ . ■

3.4 Analysis

In this section, we argue that **Learn** runs in polynomial time. Here and throughout this section, unless stated otherwise, polynomial refers to polynomial in n , s , $1/\epsilon$ and $1/\delta$. Our approach will be to first derive a bound on the *expected* running time of the procedure, and to then use a part of the confidence δ to bound with high probability the actual running time of the algorithm. Thus, we will have shown that the procedure is probably fast and correct, completing the proof of Theorem 3.1. (Although technically we only show that **Learn** halts probabilistically, by the results of Haussler et al. [10], the procedure can easily be converted into a learning algorithm that halts deterministically in polynomial time.)

We will be interested in bounding several quantities. First, we are of course interested in bounding the expected running time $T(\epsilon, \delta)$ of **Learn**(ϵ, δ, EX). This running time in turn depends on the time $U(\epsilon, \delta)$ to evaluate an hypothesis returned by **Learn**, and on the expected number of examples $M(\epsilon, \delta)$ needed by **Learn**. In addition, let $t(\delta)$, $u(\delta)$ and $m(\delta)$

be analogous quantities for $\text{WeakLearn}(\delta, EX)$. By assumption, t , u and m are polynomially bounded. Also, all of these functions depend implicitly on n and s .

As a technical point, we note that the expectations denoted by T and M are taken only over “good” runs of Learn . That is, the expectations are computed given the assumption that every sub-hypothesis and every estimator is successfully computed with the desired accuracy. By Theorem 3.2, this will be the case with probability $1 - \delta$.

It is also important to point out that T (respectively, t) is the expected running time of Learn (WeakLearn) when called with an oracle EX that provides examples *in unit time*. Our analysis will take into account the fact that the simulated oracles supplied to Learn or WeakLearn at lower levels of the recursion do not in general run in unit time.

We will see that T , U and M are all exponential in the depth of the recursion induced by calling Learn . We therefore begin by bounding this depth. Let $B(\epsilon, p)$ be the smallest integer i for which $g^i\left(\frac{1}{2} - \frac{1}{p}\right) \leq \epsilon$. On each recursive call, ϵ is replaced by $g^{-1}(\epsilon)$. Thus the depth of the recursion is bounded by $B(\epsilon, p(n, s))$. We have:

Lemma 3.1 *The depth of the recursion induced by calling $\text{Learn}(\epsilon, \delta, EX)$ is at most $B(\epsilon, p(n, s)) = O(\log(p(n, s)) + \log \log(1/\epsilon))$.*

Proof:

We can say $B(\epsilon, p(n, s)) \leq b + c$ if $g^b\left(\frac{1}{2} - \frac{1}{p(n, s)}\right) \leq \frac{1}{4}$ and $g^c\left(\frac{1}{4}\right) \leq \epsilon$. Clearly, $g(x) \leq 3x^2$ and so $g^i(x) \leq (3x)^{2^i}$. Thus, it suffices to choose $c = \lceil \lg \log_{4/3}(1/\epsilon) \rceil$. Similarly, if $\frac{1}{4} \leq x \leq \frac{1}{2}$ then $\frac{1}{2} - g(x) = \left(\frac{1}{2} - x\right)(1 + 2x - 2x^2) \geq \frac{11}{8}\left(\frac{1}{2} - x\right)$. Thus, the proof is completed by choosing $b = \lceil \log_{11/8}\left(\frac{1}{4}p(n, s)\right) \rceil$. ■

For the remainder of this analysis, we let $p = p(n, s)$ and, where clear from context, let $B = B(\epsilon, p)$.

We show next that U is polynomially bounded. This is important because we require that the returned hypothesis be polynomially evaluable.

Lemma 3.2 *The time to evaluate an hypothesis returned by $\text{Learn}(\epsilon, \delta, EX)$ is $U(\epsilon, \delta) = O(3^B \cdot u(\delta/5^B))$.*

Proof: If $\epsilon \geq \frac{1}{2} - \frac{1}{p}$, then Learn returns an hypothesis computed by WeakLearn . In this case, $U(\epsilon, \delta) = u(\delta)$. Otherwise, the hypothesis returned by Learn involves the computation of at most three sub-hypotheses. Thus,

$$U(\epsilon, \delta) \leq 3 \cdot U(g^{-1}(\epsilon), \frac{1}{5}\delta) + O(1).$$

Solving this easy recurrence, we arrive at the stated time bound. ■

When an example is requested of a simulated oracle on one of **Learn**'s recursive calls, that oracle must itself draw several examples from its own oracle EX . For instance, on the third recursive call, the simulated oracle must draw instances until it finds one on which h_1 and h_2 disagree. Naturally, the running time of **Learn** depends on how many examples must be drawn in this manner by the simulated oracle. The next lemma bounds this quantity.

Lemma 3.3 *Let r be the expected number of examples drawn from EX by any oracle EX_i simulated by **Learn** when asked to provide a single example. Then $r \leq 4/\epsilon$.*

Proof: When **Learn** calls itself the first time (to find h_1), the examples oracle EX it was passed is left unchanged. In this case, $r = 1$.

The second time **Learn** calls itself, the constructed oracle EX_2 loops each time it is called until it receives a desirable example. Depending on the result of the initial coin flip, we expect EX_2 to loop $1/a_1$ or $1/(1 - a_1)$ times. Note that if $a_1 \leq \epsilon - 2\tau_1 = \frac{1}{3}\epsilon$ then, based on its estimate of a_1 , **Learn** would have simply returned h_1 instead of making a second or third recursive call. Thus, we can assume $\frac{1}{3}\epsilon \leq a_1 \leq \frac{1}{2}$, and so $r \leq 3/\epsilon$ in this case.

Finally, when **Learn** calls itself the third time, we expect the constructed oracle EX_3 to loop $1/(w + y)$ times before finding a suitable example. (Here, the variables w , x , y and z are as defined in the proof of Theorem 3.2.) It remains then only to show that $w + y \geq \frac{1}{4}\epsilon$. Note that the error e of h_2 on the original distribution D is $w + z$. Thus, using this fact and equations (1), (2) and (6), we can solve explicitly for w and y in terms of e , a_1 and a_2 , and so find that

$$w + y = a_1 + \frac{\epsilon - 4a_1a_2(1 - a_1)}{1 - 2a_1} \geq a_1 + \frac{\epsilon - 4a_1\alpha(1 - a_1)}{1 - 2a_1} \quad (7)$$

To lower bound $w + y$, we will find the minimum of this second function on the interval $[0, \alpha]$. Differentiating (with respect to a_1) we see that the function has at most one critical point less than $\frac{1}{2}$, and we note further that such a critical point cannot be minimal since the function tends to $-\infty$ as $a_1 \rightarrow -\infty$. This means that the function's minimum on any closed subinterval of $(-\infty, \frac{1}{2})$ is achieved at one endpoint of the subinterval. In particular, for the subinterval of interest to us, the function achieves its minimum either when $a_1 = 0$ or when $a_1 = \alpha$.

We can assume that $e \geq \epsilon - 2\tau_2 = \left(\frac{3}{4} + \frac{1}{2}\alpha\right)\epsilon$; otherwise, if e were smaller than this quantity, then **Learn** would have returned h_2 rather than going on to compute h_3 . Thus, if

$a_1 = 0$ then $w + y \geq e \geq \frac{3}{4}\epsilon$, and if $a_1 = \alpha$ then $w + y \geq \frac{1}{4}\alpha(4 - 7\alpha + 2\alpha^2) \geq \frac{1}{4}\alpha \geq \frac{1}{4}\epsilon$. ■

Lemma 3.4 *The expected number of examples $M(\epsilon, \delta)$ needed by $\text{Learn}(\epsilon, \delta, EX)$ is $O\left(\frac{36^B}{\epsilon^2} \cdot (Bp^2 + p^2 \log(1/\delta) + m(\delta/5^B))\right)$.*

Proof: In the base case that $\epsilon \geq \frac{1}{2} - \frac{1}{p}$, Learn simply calls WeakLearn , so we have $M(\epsilon, \delta) = m(\delta)$. On each of the recursive calls, the simulated oracle is required to provide $M(g^{-1}(\epsilon), \frac{1}{5}\delta)$ examples. To provide one such example, the simulated oracle must itself draw at most an average of $4/\epsilon$ examples from EX . Thus, each recursive call demands at most $(4/\epsilon) \cdot M(g^{-1}(\epsilon), \frac{1}{5}\delta)$ examples on average.

In addition, Learn requires some examples for making its estimates $\hat{a}_1$ and $\hat{e}$. Using Chernoff bounds [22], we can show that a sample of size $O(p^2 \log(1/\delta)/\epsilon^2)$ suffices.

We thus arrive at the recurrent equation:

$$M(\epsilon, \delta) \leq \frac{12}{\epsilon} \cdot M(g^{-1}(\epsilon), \frac{1}{5}\delta) + O\left(\frac{p^2 \log(1/\delta)}{\epsilon^2}\right).$$

Making use of the fact that $g^{-1}(\epsilon) \geq \sqrt{\frac{1}{3}}\epsilon$ and that $B(g^{-1}(\epsilon), p) = B(\epsilon, p) - 1$, we can solve this recurrence and arrive at the stated bound. ■

Lemma 3.5 *The expected execution time of $\text{Learn}(\epsilon, \delta, EX)$ is given by $T(\epsilon, \delta) = O\left(3^B \cdot t(\delta/5^B) + \frac{108^B \cdot u(\delta/5^B)}{\epsilon^2} \cdot (Bp^2 + p^2 \log(1/\delta) + m(\delta/5^B))\right)$.*

Proof: As in the previous lemmas, the base case that $\epsilon \geq \frac{1}{2} - \frac{1}{p}$ is easily handled. In this case, $T(\epsilon, \delta) = t(\delta)$.

Otherwise, Learn takes time $3 \cdot T(g^{-1}(\epsilon), \frac{1}{5}\delta)$ on its three recursive calls. In addition, Learn spends time drawing examples to make the estimates $\hat{a}_1$ and $\hat{e}$, and overhead time is also spent by the simulated examples oracles passed on the three recursive calls. A typical example that is drawn from Learn 's oracle EX is evaluated on zero, one or two of the previously computed sub-hypotheses. For instance, an example drawn for the purpose of estimating $\hat{a}_1$ is evaluated once by h_1 ; an example drawn for the simulated oracle EX_3 is evaluated by both h_1 and h_2 . Thus, Learn 's overhead time is proportional to the product of the total number of examples needed by Learn and the time it takes to evaluate a sub-hypothesis on one of these examples. Therefore, the following recurrence holds:

$$T(\epsilon, \delta) \leq 3 \cdot T(g^{-1}(\epsilon), \frac{1}{5}\delta) + O\left(U(g^{-1}(\epsilon), \frac{1}{5}\delta) \cdot M(\epsilon, \delta)\right) \quad (8)$$

Applying Lemmas 3.2 and 3.4, this recurrence implies the stated bound. ■

The main result of this section follows immediately:

Theorem 3.3 *Let $0 < \epsilon < \frac{1}{2}$ and let $0 < \delta \leq 1$. With probability at least $1 - \delta$, the execution of $\text{Learn}(\epsilon, \frac{1}{2}\delta, EX)$ halts in polynomial time and outputs an hypothesis ϵ -close to the target concept.*

Proof: The chance that Learn does not output an hypothesis ϵ -close to c is at most $\frac{1}{2}\delta$. The chance that such an hypothesis is output after time $(2/\delta) \cdot T(\epsilon, \frac{1}{2}\delta)$ is also at most $\frac{1}{2}\delta$. ■

3.5 Space Complexity

Although not of immediate consequence to the proof of Theorem 3.3, it is worth pointing out that Learn 's space requirements are relatively modest, as proved in this section.

Let $S(\epsilon, \delta)$ be the space used by $\text{Learn}(\epsilon, \delta, EX)$; let $Q(\epsilon, \delta)$ be the space needed to store an output hypothesis; and let $R(\epsilon, \delta)$ be the space needed to evaluate such an hypothesis. Let $s(\delta)$, $q(\delta)$ and $r(\delta)$ be analogous quantities for $\text{WeakLearn}(\delta, EX)$. Then we have:

Lemma 3.6 *The space $Q(\epsilon, \delta)$ required to store an hypothesis output by $\text{Learn}(\epsilon, \delta, EX)$ is at most $O(3^B \cdot q(\delta/5^B))$. The space $R(\epsilon, \delta)$ needed to evaluate such an hypothesis is $O(B + r(\delta/5^B))$. Finally, the total space $S(\epsilon, \delta)$ required by Learn is $O(3^B \cdot q(\delta/5^B) + s(\delta/5^B) + B \cdot r(\delta/5^B))$.*

Proof: For $\epsilon \geq \frac{1}{2} - \frac{1}{p}$, the bounds are trivial. Otherwise, the following recurrences are easy to derive bounding Q and R :

$$Q(\epsilon, \delta) \leq 3 \cdot Q(g^{-1}(\epsilon), \frac{1}{5}\delta) + O(1)$$

$$R(\epsilon, \delta) \leq R(g^{-1}(\epsilon), \frac{1}{5}\delta) + O(1)$$

For bounding S , note that the space required by Learn is dominated by the storage of the sub-hypotheses, by their recursive computation, and by the space needed to evaluate them. Thus,

$$S(\epsilon, \delta) \leq S(g^{-1}(\epsilon), \frac{1}{5}\delta) + O\left(Q(g^{-1}(\epsilon), \frac{1}{5}\delta) + R(g^{-1}(\epsilon), \frac{1}{5}\delta)\right)$$

which implies the desired bound. ■

4 Improving Learn's Time and Sample Complexity

In this section, we describe a modification to the construction of Section 3 that significantly improves **Learn**'s time and sample complexity. In particular, we will improve these complexity measures by roughly a factor of $1/\epsilon$, giving bounds that are linear in $1/\epsilon$ (ignoring log factors). These improved bounds will have some interesting consequences, described in later sections.

In the original construction of **Learn**, much time and many examples are squandered by the simulated oracles EX_i waiting for a desirable instance to be drawn. Lemma 3.3 showed that the expected time spent waiting is $O(1/\epsilon)$. The modification described below will reduce this to $O(1/\alpha) = O(1/\sqrt{\epsilon})$. (Here, $\alpha = g^{-1}(\epsilon)$ as before.)

Recall that the running time of oracle EX_2 depends on the error a_1 of the first sub-hypothesis h_1 . In the original construction, we ensured that a_1 not be too small by estimating its value, and, if smaller than ϵ , returning h_1 instead of continuing the normal execution of the subroutine. Since this approach only guarantees that $a_1 \geq \Omega(\epsilon)$, there does not seem to be any way of ensuring that EX_2 run for $o(1/\epsilon)$ time. To improve EX_2 's running time then, we will instead modify h_1 by deliberately *increasing* its error. Ironically, this intentional injection of error will have the effect of improving **Learn**'s worst case running time by limiting the time spent by either EX_2 or EX_3 waiting for a suitable instance.

4.1 The Modifications

Specifically, here is how **Learn** is modified. Call the new procedure **Learn'**. Following the recursive computation of h_1 , **Learn'** estimates the error a_1 of h_1 , although less accurately than **Learn**. Let $\hat{a}_1$ be this estimate, and choose a sample large enough that $|a_1 - \hat{a}_1| \leq \frac{1}{4}\alpha$ with probability at least $1 - \frac{1}{5}\delta$. Since $0 \leq a_1 \leq \alpha$, we can assume without loss of generality that $\frac{1}{4}\alpha \leq \hat{a}_1 \leq \frac{3}{4}\alpha$.

Next, **Learn'** defines a new hypothesis h'_1 as follows: given an instance v , h'_1 first flips a coin biased to turn up "heads" with probability exactly

$$p = \frac{\frac{3}{4}\alpha - \hat{a}_1}{1 - \frac{1}{4}\alpha - \hat{a}_1}.$$

If the outcome is "tails," then h'_1 evaluates $h_1(v)$ and returns the result. Otherwise, if "heads," h'_1 predicts the wrong answer, $\neg c(v)$. Since h'_1 will only be used during the training

phase, we can assume that the correct classification of v is available, and thus that h'_1 can be simulated.

This new hypothesis h'_1 is now used in place of h_1 by EX_2 and EX_3 . The rest of the subroutine is unmodified. In particular, the final returned hypothesis h is unchanged — that is, h_1 , not h'_1 , is used by h .

4.2 Correctness

To see that **Learn'** is correct, note first that the error of h'_1 is exactly $a'_1 = (1 - p)a_1 + p$ since the chance of error is a_1 on “tails,” and is 1 on “heads.” By our choice of p , it can be verified that $\frac{1}{2}\alpha \leq a'_1 \leq \alpha$.

Let h' be the same hypothesis as h , except with h'_1 used in lieu of h_1 . Note that h' , h'_1 , h_2 and h_3 are related to one another in exactly the same way that h , h_1 , h_2 and h_3 are related in the original proof of Theorem 3.2. That is, if we imagine that h'_1 is returned on the first recursive call of the original procedure **Learn**, then it is not impossible that h_2 and h_3 would be returned on the second and third recursive calls, in which case h' would be the returned hypothesis. Therefore, by the same argument used to prove Theorem 3.2, the error of h' is at most $g(\alpha) = \epsilon$.

It is not surprising, and is not hard to verify, that increasing h_1 's chance of error on instance v cannot possibly decrease the chance that h misclassifies v . Therefore, since $\Pr[h'_1(v) \neq c(v)] \geq \Pr[h_1(v) \neq c(v)]$ for any v , the error of h cannot be greater than the error h' , which is bounded by ϵ .

4.3 Analysis

Next, we show that **Learn'** runs faster using fewer examples than **Learn**. We use essentially the same analysis as in Section 3.4. The following three lemmas are modified versions of Lemmas 3.3, 3.4 and 3.5. The proofs of the other lemmas apply immediately to **Learn'** with little or no modification, and so are omitted.

Lemma 4.1 *Let r be the expected number of examples drawn from EX by any oracle EX_i simulated by **Learn'** when asked to provide a single example. Then $r \leq 4/\alpha$.*

Proof: As in the original proof, $r = 1$ for EX_1 . We expect the second oracle to loop at most $1/a'_1$ times on average. Since $a'_1 \geq \frac{1}{2}\alpha$, r is at most $2/\alpha$ in this case.

Finally, to bound the number of iterations of EX_3 , we will show that $w + y \geq \frac{1}{4}\alpha$ using equation (7) as in the original proof. To lower bound $w + y$, we find the minimum of the last formula of (7) (with a_1 replaced by a'_1 of course) on the interval $[\frac{1}{2}\alpha, \alpha]$. As noted previously, the function must achieve its minimum at one endpoint of the interval. Assuming as in the original proof that $e \geq (\frac{3}{4} + \frac{1}{2}\alpha)\epsilon$, we see that when $a'_1 = \alpha$, $w + y \geq \frac{1}{4}\alpha(4 - 7\alpha + 2\alpha^2) \geq \frac{1}{4}\alpha$. Similarly, when $a'_1 = \frac{1}{2}\alpha$, $w + y \geq \frac{1}{2}\alpha + \alpha^3 + \frac{1}{4}\alpha^2/(1 - \alpha) \geq \frac{1}{2}\alpha$. This completes the proof. ■

Lemma 4.2 *The expected number of examples $M(\epsilon, \delta)$ needed by $\text{Learn}'(\epsilon, \delta, EX)$ is $O\left(\frac{36^B}{\epsilon} \cdot (Bp^2 + p^2 \log(1/\delta) + m(\delta/5^B))\right)$.*

Proof: The proof is nearly the same as for Lemma 3.4. In addition to incorporating the superior bound given by Lemma 4.1 on the number of examples needed by the simulated oracles, we must also consider the number of examples needed to estimate a_1 and e . The first, a_1 , can be estimated using a sample of size $O(\log(1/\delta)/\alpha^2) = O(\log(1/\delta)/\epsilon)$. By estimating e in a slightly different manner, we can also achieve a better bound on the sample size needed. Specifically, we can choose a sample large enough that, with probability $1 - \frac{1}{5}\delta$, $\hat{e} \leq \epsilon - \tau_2$ if $e \leq \epsilon - 2\tau_2$, and $\hat{e} \geq \epsilon - \tau_2$ if $e \geq \epsilon$. Such an estimate has all of the properties needed by Learn' , but only requires a sample of size $O(p^2 \log(1/\delta)/\epsilon)$ as can be derived using Chernoff bounds [22].

Thus, we arrive at the recurrence

$$M(\epsilon, \delta) \leq \frac{12}{g^{-1}(\epsilon)} \cdot M(g^{-1}(\epsilon), \frac{1}{5}\delta) + O\left(\frac{p^2 \log(1/\delta)}{\epsilon}\right)$$

which implies the stated bound. ■

Lemma 4.3 *The expected execution time of $\text{Learn}'(\epsilon, \delta, EX)$ is given by $T(\epsilon, \delta) = O\left(3^B \cdot t(\delta/5^B) + \frac{108^B \cdot u(\delta/5^B)}{\epsilon} \cdot (Bp^2 + p^2 \log(1/\delta) + m(\delta/5^B))\right)$.*

Proof: This bound follows from the recurrence (8), using the superior bound on M given by Lemma 4.2. ■

5 Variations on the Learning Model

Next, we consider how the main result relates to some other learning models.

5.1 Group Learning

An immediate consequence of Theorem 3.1 concerns *group learnability*. In the group learning model, the learner produces a hypothesis that need only correctly classify large groups of instances, all of which are either positive or negative examples. Kearns and Valiant [14, 17] prove the equivalence of group learning and weak learning. Thus, by Theorem 3.1, group learning is also equivalent to strong learning.

5.2 Miscellaneous PAC Models

Haussler et al. [10] describe numerous variations on the basic PAC model, and show that all of them are equivalent. For instance, they consider randomized versus deterministic algorithms, algorithms for which the size s of the target concept is known or unknown, and so on. It is not hard to see that all of their equivalence proofs apply to weak learning algorithms as well (with one exception described below), and so that any of these weak learning models are equivalent by Theorem 3.1 to the basic PAC-learning model.

The one reduction from their paper that does not hold for weak learning algorithms concerns the equivalence of the one- and two-oracle learning models. In the one-oracle model (used exclusively in this paper), the learner has access to a single source of positive and negative examples. In the two-oracle model, the learner has access to one oracle that returns only positive examples, and another returning only negative examples. The authors show that these models are equivalent for strong learning algorithms. However, their proof apparently cannot be adapted to show that one-oracle weak learnability implies two-oracle weak learnability (although their proof of the converse is easily and validly adapted). This is because their proof assumes that the error ϵ can be made arbitrarily small, clearly a bad assumption for weak learning algorithms. Nevertheless, this is not a problem since we have shown that one-oracle weak learnability implies one-oracle strong learnability, which in turn implies two-oracle strong (and therefore weak) learnability. Thus, despite the inapplicability of Haussler et al.'s original proof, all four learning models are equivalent.

5.3 Fixed Hypotheses

Much of the PAC-learning research has been concerned with the form or representation of the hypotheses output by the learning algorithm. Clearly, the construction described in

Section 3 does not in general preserve the form of the hypotheses used by the weak learning algorithm. It is natural to ask whether there exists any construction preserving this form. That is, if concept class $\mathcal{C}$ is weakly learnable by an algorithm using hypotheses from a class $\mathcal{H}$ of representations, does there then exist a strong learning algorithm for $\mathcal{C}$ that also only outputs hypotheses from $\mathcal{H}$?

In general, the answer to this question is “no” (modulo some relatively weak complexity assumptions). As a simple example, consider the problem of learning k -term DNF formulas using only hypotheses represented by k -term DNF. (A formula in disjunctive normal form (DNF) is one written as a disjunction of terms, each of which is a conjunction of literals, a literal being either a variable or its complement.) Pitt and Valiant [19] show that this learning problem is infeasible if $\text{RP} \neq \text{NP}$.

Nevertheless, the weak learning problem is solved by the algorithm sketched below. (A similar algorithm is given by Kearns [15].) First, choose a “large” sample. If significantly more than half of the examples in the sample are negative (positive), then output the “always predict negative (positive)” hypothesis, and halt. Otherwise, we can assume that the distribution is roughly evenly split between positive and negative examples. Select and output the disjunction of k or fewer literals that misclassifies none of the positive examples, and the fewest of the negative examples. Working through the details, it is not hard to show that this output formula is correct for nearly all of the positive examples and for at least $\Omega(1/n^k)$ of the negative examples. Since the distribution is roughly evenly divided between positive and negative examples, this implies that the output hypothesis is roughly $(\frac{1}{2} - \Omega(\frac{1}{n^k}))$ -close to the target formula.

5.4 Queries

A number of researchers have considered learning scenarios in which the learner is not only able to passively observe randomly selected examples, but is also able to ask a “teacher” various sorts of questions or *queries* about the target concept. For instance, the learner might be allowed to ask if some particular instance is a positive or negative example. Angluin [2] describes several kinds of query that might be useful to the learner. The purpose of this section is simply to point out that the construction of Section 3 is applicable even in the presence of most kinds of query. That is, a weak learning algorithm that depends on the availability of certain kinds of query can be converted, using the same construction, into a

strong learning algorithm using the same query types.

5.5 Many-Valued Concepts

In this paper, we have only considered Boolean valued concepts, i.e., concepts that classify every instance as either a positive or a negative example. Of course, in the “real world,” most learning tasks require classification into one of several categories (for instance, character recognition). How does the result generalize to handle many-valued concepts?

First of all, for learning a k -valued concept, it is not immediately clear how to define the notion of weak learnability. An hypothesis that guesses randomly on every instance will be correct only $1/k$ of the time, so one natural definition would require only that the weak learning algorithm classify instances correctly slightly more than $1/k$ of the time. Unfortunately, under this definition, strong and weak learnability are inequivalent for k as small as three. As an informal example, consider learning a concept taking the values 0, 1 and 2, and suppose that it is “easy” to predict when the concept has the value 2, but “hard” to predict whether the concept’s value is 0 or 1. Then to weakly learn such a concept, it suffices to find an hypothesis that is correct whenever the concept is 2, and that guesses randomly otherwise. For any distribution, this hypothesis will be correct half of the time, achieving the weak learning criterion of accuracy significantly better than $\frac{1}{3}$. However, boosting the accuracy further is clearly infeasible.

Thus, a better definition of weak learnability is one requiring that the hypothesis be correct on slightly more than half of the distribution, regardless of k . Using this definition, the construction of Section 3 is easily modified to handle many-valued concepts.

6 General Complexity Bounds for PAC Learning

The construction derived in Sections 3 and 4 yields some unexpected relationships between the allowed error ϵ and various complexity measures that might be applied to a strong learning algorithm. One of the more surprising of these is a proof that, for every learnable concept class, there exists an efficient algorithm whose output hypotheses can be evaluated in time polynomial in $\log(1/\epsilon)$. Furthermore, such an algorithm’s space requirements are also only poly-logarithmic in $1/\epsilon$ — far less, for instance, than would be needed to store the entire sample. In addition, its time and sample size requirements grow only linearly in $1/\epsilon$

(disregarding log factors).

Theorem 6.1 *If $\mathcal{C}$ is a learnable concept class, then there exists an efficient learning algorithm for $\mathcal{C}$ that:*

- *requires a sample of size $\frac{p_1(n, s, \log(1/\epsilon), \log(1/\delta))}{\epsilon}$,*
- *runs in time $\frac{p_2(n, s, \log(1/\epsilon), \log(1/\delta))}{\epsilon}$,*
- *uses space $p_3(n, s, \log(1/\epsilon), \log(1/\delta))$, and*
- *outputs hypotheses of size $p_4(n, s, \log(1/\epsilon), \log(1/\delta))$*

for some polynomials p_1, p_2, p_3 and p_4 .

Proof: Given a strong learning algorithm A for $\mathcal{C}$, convert A into a weak learning algorithm A' that outputs hypotheses $\frac{1}{4}$ -close to the target concept. Now let A'' be the procedure obtained by applying the construction of **Learn'** with A' plugged in for **WeakLearn**. Furthermore, assume without loss of generality, using the results of Haussler et al. [10], that A'' halts deterministically in time polynomial in $\log(1/\delta)$. Then the lemmas of Sections 3 and 4 imply that A'' has all of the stated properties. ■

The remainder of this section is a discussion of some of the consequences of Theorem 6.1.

6.1 Improving the Performance of Existing Algorithms

These bounds can be applied immediately to a number of existing learning algorithms, yielding improvements in time and/or space complexity (at least in terms of ϵ). For instance, the computation time of Blumer et al.'s [4] algorithm for learning half-spaces of R^n , which involves the solution of a linear programming problem of size proportional to the sample, can be improved by a polynomial factor of $1/\epsilon$. The same is also true of Baum's [3] algorithm for learning unions of half-spaces, which involves finding the convex hull of a significant fraction of the sample.

There are many more algorithms for which the theorem implies improved space efficiency. This is especially true of the many known PAC algorithms that work by choosing a large sample and then finding an hypothesis consistent with it. For instance, this is how Rivest's [20] decision list algorithm works, as do most of the algorithms described by Blumer et al. [4], as

well as Helmbold, Sloan and Warmuth's [13] construction for learning nested differences of learnable concepts. Since the entire sample must be stored, these algorithms are not terribly space efficient, and so can be dramatically improved by applying Theorem 6.1. Of course, these improvements typically come at the cost of requiring a somewhat larger sample (by a polynomial factor of $\log(1/\epsilon)$). Thus, there appears to be a trade-off between sample size and space (or time) complexity.

6.2 Data Compression

Blumer et al. [5, 4] have considered the relationship between learning and data compression. They have shown that, if any sample can be "compressed" — i.e., represented by a prediction rule significantly smaller than the original sample — then this compression algorithm can be converted into a PAC-learning algorithm.

In some sense, the bound given in Theorem 6.1 on the size of the output hypothesis implies the converse. In particular, suppose $\mathcal{C}_n$ is a learnable concept class and that we have been given m examples $(v_1, c(v_1)), (v_2, c(v_2)), \dots, (v_m, c(v_m))$ where each $v_i \in X_n$ and c is a concept in $\mathcal{C}_n$ of size s . The data compression problem is to find a small representation for the data, i.e., an hypothesis h that is significantly smaller than the original data set with the property that $h(v_i) = c(v_i)$ for each v_i . An hypothesis with this last property is said to be *consistent* with the sample.

Theorem 6.1 implies the existence of an efficient algorithm that outputs consistent hypotheses only poly-logarithmic in the size m of the sample. This is proved by the following theorem:

Theorem 6.2 *Let $\mathcal{C}$ be a learnable concept class. Then there exists an efficient algorithm that, given $0 < \delta \leq 1$ and m examples of a concept $c \in \mathcal{C}_n$ of size s , outputs with probability at least $1 - \delta$ a deterministic hypothesis consistent with the sample of size polynomial in n , s and $\log m$.*

Proof: Haussler et al. [10] show how to convert any learning algorithm into one that finds hypotheses consistent with a set of data points. The idea is to choose $\epsilon < 1/m$ and to run the learning algorithm on a (simulated) uniform distribution over the data set. Since ϵ is less than the weight placed on any element of the sample, the output hypothesis must have error zero. Applying this technique to a learning algorithm A satisfying the conditions of

Theorem 6.1, we see that the output hypothesis has size only polynomial in n , s and $\log m$, and so is far smaller than the original sample for large m .

Technically, this technique requires that the learning algorithm output deterministic hypotheses. However, probabilistic hypotheses can also be handled by choosing a somewhat smaller value for ϵ , and by “hard-wiring” the computed probabilistic hypothesis with a sequence of random bits. More precisely, set $\epsilon = 1/2m^2$, and run A over the same distribution as before. Then with high probability, the computed hypothesis h has chance of error at most $1/2m$ on any one of the m examples in the sample. Now note that h can be regarded as a deterministic function of an instance v and a sequence of random bits r . If r is such a sequence, then the chance that $h(\cdot, r)$ correctly classifies all of the m examples is at least $\frac{1}{2}$. Thus, choosing and testing random sequences r , we can quickly find one for which the deterministic hypothesis $h(\cdot, r)$ is consistent with the sample. Finally, note that the size of this output hard-wired hypothesis is bounded by $|h| + |r|$, and that $|r|$ is bounded by the time it takes to evaluate h , which is poly-logarithmic in m . ■

Naturally, the notion of size in the preceding theorem depends on the underlying model of computation, which has deliberately been left unspecified. However, the theorem has some immediate corollaries when the learning problem is discrete, i.e. when every instance in the domain X_n is encoded using a finite alphabet by a string of length presumably bounded by a polynomial in n .

Corollary 6.1 *Let $\mathcal{C}$ be a learnable discrete concept class. Then there exists an efficient algorithm that, given $0 < \delta \leq 1$ and a sample as in Theorem 6.2, outputs a deterministic consistent hypothesis of size polynomial in n and s , and independent of m .*

Proof: The size m of the sample is clearly bounded by $|X_n|$. Since $\log |X_n|$ is bounded by a polynomial in n , the corollary follows immediately. ■

Applying “Occam’s Razor” of Blumer et al. [5], this implies the following strong general bound on the sample size needed to efficiently learn $\mathcal{C}$. Although the bound is better than that given by Theorem 6.1 (at least in terms of ϵ), it should be pointed out that this improvement requires the sacrifice of space efficiency since the entire sample must be stored.

Theorem 6.3 *Let $\mathcal{C}$ be a learnable discrete concept class. Then there exists an efficient learning algorithm for $\mathcal{C}$ requiring a sample of size $O\left(\frac{p(n, s) + \log(1/\delta)}{\epsilon}\right)$ for some polynomial p .*

6.3 Hard Functions are Hard to Learn

Theorem 6.1's bound on the size of the output hypothesis also implies that any hard-to-evaluate concept class is unlearnable. Although this result does not sound surprising, it was previously unclear how it might be proved: since a learning algorithm's hypotheses are technically permitted to grow polynomially in $1/\epsilon$, the learnability of such classes did not seem out of the question.

This result yields the first representation-independent hardness results not based on cryptographic assumptions. For instance, assuming $P/\text{poly} \neq NP/\text{poly}$, the class of polynomial-size, nondeterministic Boolean circuits is not learnable. (The set P/poly (NP/poly) consists of those languages accepted by a family of polynomial-size deterministic (nondeterministic) circuits.) Furthermore, since learning pattern languages was recently shown [21] to be as hard as learning NP/poly , this result shows that pattern languages are also unlearnable under this relatively weak structural assumption.

Theorem 6.4 *Suppose C is learnable, and assume that $X_n = \{0,1\}^n$. Then there exists a polynomial p such that for all concepts $c \in C_n$ of size s , there exists a circuit of size $p(n, s)$ exactly computing c .*

Proof: Consider the set of 2^n pairs $\{(v, c(v)) \mid v \in X_n\}$. By Corollary 6.1, there exists an algorithm that, with positive probability, will output an hypothesis consistent with this set of elements of size only polynomial in n and s . Since this hypothesis is polynomially evaluable, it can be converted using standard techniques into a circuit of the required size.

■

6.4 Hard Functions are Hard to Approximate

By a similar argument, the bound on hypothesis size implies that any function not computable by small circuits cannot even be weakly approximated by a family of small circuits, for some distribution on the inputs.

Let f be a Boolean function on $\{0,1\}^n$, D a distribution on $\{0,1\}^n$ and C a circuit on n variables. Then C β -approximates f under D if the probability is at most β that $C(v) \neq f(v)$ on an assignment v chosen randomly from $\{0,1\}^n$ according to D .

Theorem 6.5 Suppose some function f cannot be computed by any family of polynomial-size circuits. Then there exists a family of distributions $D_1, D_2, \dots$, where D_n is over the set $\{0,1\}^n$, such that for all polynomials p and q , there exist infinitely many n for which there exists no n -variable circuit of size at most $q(n)$ that $(\frac{1}{2} - \frac{1}{p(n)})$ -approximates f under D_n .

Proof: Throughout this proof, we will assume without loss of generality that $p(n) = q(n) = n^k$ for some integer $k \geq 1$.

Suppose first that for some k there exists for every distribution D on $\{0,1\}^n$ a circuit of size at most n^k that $(\frac{1}{2} - \frac{1}{n^k})$ -approximates f under D . Then f can, in a sense, be weakly learned. More precisely, there exists an (exponential time) procedure that, by searching exhaustively the set of all circuits of size n^k , will find one that $(\frac{1}{2} - \frac{1}{n^k})$ -approximates f under some given distribution D . Therefore, by Theorem 3.1, f is strongly learnable in a similar sense in exponential time. Applying Theorem 6.4 (whose validity depends only on the size of the output hypothesis, and not on the running time), this implies that f can be exactly computed by a family of polynomial-size circuits, contradicting the theorem's hypothesis.

Thus, for all $k \geq 1$, there exists an integer n and a distribution D on $\{0,1\}^n$ such that no circuit of size at most n^k is able to $(\frac{1}{2} - \frac{1}{n^k})$ -approximate f under D . To complete the proof, it suffices to show that this implies the theorem's conclusion.

Let $\mathcal{D}_n^k$ be the set of distributions D on $\{0,1\}^n$ for which no circuit of size at most n^k $(\frac{1}{2} - \frac{1}{n^k})$ -approximates f under D . It is easy to verify that $\mathcal{D}_n^k \supseteq \mathcal{D}_n^{k+1}$ for all k, n . Also, since every function can be computed by exponential size circuits, there must exist a constant $c > 0$ for which $\mathcal{D}_n^{cn} = \emptyset$ for all n . Let $n[k]$ be the smallest n for which $\mathcal{D}_n^k \neq \emptyset$. By the preceding argument, $n[k]$ must exist. Furthermore, $n[k] \geq k/c$, which implies that the set $N = \{n[k] \mid k \geq 1\}$ cannot have finite cardinality.

To eliminate repeated elements from N , let $k_1 < k_2 < \dots$ be such that $n[k_i] \neq n[k_j]$ for $i \neq j$, and such that $\{n[k_i] \mid i \geq 1\} = N$. Let D_i be defined as follows: if $i = n[k_j]$ for some j , then let D_i be any distribution in $\mathcal{D}_i^{k_j}$ (which cannot be empty by our definition of $n[k]$); otherwise, if $i \notin N$, then define D_i arbitrarily. Then $D_1, D_2, \dots$ is the desired family of "hard" distributions. For if k is any integer, then for all $k_i \geq k$, $D_{n[k_i]} \in \mathcal{D}_{n[k_i]}^{k_i} \subseteq \mathcal{D}_{n[k_i]}^k$. This proves the theorem. ■

Informally, Theorem 6.5 states that any language not in the complexity class P/poly cannot be even weakly approximated by any other language in P/poly under some "hard"

family of distributions. In fact, the theorem can easily be modified to apply to other circuit classes as well, including monotone P/poly, and monotone or non-monotone NC^k for fixed k . (The class NC^k consists of all languages accepted by polynomial-size circuits of depth at most $O(\log^k n)$, and a monotone circuit is one in which no negated variables appear.) In general, the theorem applies to all circuit classes closed under the transformation on hypotheses resulting from the construction of Sections 3 and 4.

6.5 On-Line Learning

Finally, we consider implications of Theorem 6.1 for on-line learning algorithms. In the *on-line* learning model, the learner is presented one (randomly selected) instance at a time in a series of *trials*. Before being told its correct classification, the learner must try to predict whether the instance is a positive or negative example. An incorrect prediction is called a *mistake*. In this model, the learner's goal is to minimize the number of mistakes.

Previously, Haussler, Littlestone and Warmuth [11] have shown that a concept class $\mathcal{C}$ is learnable if and only if there exists an on-line learning algorithm for $\mathcal{C}$ with the properties that:

- the probability of a mistake on the m th trial is at worst linear in $m^{-\beta}$ for some constant $0 < \beta \leq 1$, and (equivalently)
- the expected number of mistakes on the first m trials is at worst linear in m^α for some constant $0 \leq \alpha < 1$.

(This result is also described in their paper with Kearns [10].) Noting several examples of learning algorithms for which this second bound only grows poly-logarithmically in m , the authors ask if *every* learnable concept class has an algorithm attaining such a bound. Theorem 6.6 below answers this open question affirmatively, showing that in general the expected number of mistakes on the first m trials need only grow as a polynomial in $\log m$. Thus, we expect only a minute fraction of the first m predictions to be incorrect.

(This result should not be confused with those presented in another paper by Haussler, Littlestone and Warmuth [12]. In this paper, the authors describe a general algorithm applicable to a wide collection of concept classes, and they show that the expected number of mistakes made by this algorithm on the first m trials is linear in $\log m$. However, their algorithm requires exponential computation time, even if it is known that the concept class is learnable. In contrast, Theorem 6.6 states that, if a concept class is learnable, then there

exists an efficient algorithm making poly-logarithmic in m mistakes on average on the first m trials.)

Haussler, Littlestone and Warmuth [11] also consider the space efficiency of on-line learning algorithms. They define a *space-efficient* learning algorithm to be one whose space requirements on the first m trials do not exceed a polynomial in n , s and $\log m$. Thus, a space efficient algorithm is one using far less memory than would be required to store explicitly all of the preceding observations. The authors describe a number of space-efficient algorithms (though are unable to find one for learning unions of axis-parallel rectangles in the plane), and so are lead to ask whether there exist space-efficient algorithms for *all* learnable concept classes. Surprisingly, this open question can also be answered affirmatively, as proved by the theorem below.

Lastly, Theorem 6.6 gives a bound on the computational complexity of on-line learning (in terms of ϵ). In particular, the total computation time required to process the first m examples is only proportional to $m \log^c m$, for some constant c . Thus, in a sense, the “amortized” or “average” computation time on the m th trial is only poly-logarithmic in m . (In fact, a more careful analysis would show that this is also true of the worst case computation time on the m th trial.)

Theorem 6.6 *Let $\mathcal{C}$ be a learnable concept class. Then there exists an efficient on-line learning algorithm for $\mathcal{C}$ with the properties that:*

- *the probability of a mistake on the m th trial is at most $\frac{p_1(n, s, \log m)}{m}$,*
- *the expected number of mistakes on the first m trials is at most $p_2(n, s, \log m)$,*
- *the total computation time required on the first m trials is at most $m \cdot p_3(n, s, \log m)$,*
and
- *the space used on the first m trials is at most $p_4(n, s, \log m)$,*

for some polynomials p_1, p_2, p_3, p_4 .

Proof: Since $\mathcal{C}$ is learnable, there exists an efficient (batch) algorithm satisfying the properties of Theorem 6.1. Let A be such an algorithm, but with $\epsilon/2$ substituted for both ϵ and δ . Then the chance that A 's output hypothesis incorrectly classifies a randomly chosen instance is at most ϵ . (This technique is also used by Haussler et al. [10].)

Fix n and s , and let $m(\epsilon)$ be the number of examples needed by A . From Theorem 6.1, $m(\epsilon) \leq (p/\epsilon) \cdot \lg^c(1/\epsilon)$ for some constant c and some value p implicitly bounded by a polynomial in n and s . Let $\epsilon(m) = (p/m) \cdot \lg^c(m/p)$. Then it can be verified that $m(\epsilon(m)) \leq m$ for $m \geq 2p$. Thus, for sufficiently large m , $\epsilon(m)$ gives a bound on the best error rate achievable from a sample of size m .

To convert A into an on-line learning algorithm in a manner that preserves time and space efficiency, imagine breaking the sequence of trials into blocks of increasing size: the first block consists of the first $2p$ trials, and each new block has twice the size of the last. Thus, in general, the i th block has size $s_i = 2^i p$, and consists of trials $a_i = 2(2^{i-1} - 1)p + 1$ through $b_i = 2(2^i - 1)p$.

On the trials of the i th block, algorithm A is simulated to compute the i th hypothesis h_i . Specifically, A is simulated with ϵ set to $\epsilon(s_i)$, which thus bounds the probability that h_i misclassifies a new instance. (Note that there are enough instances available in this block for A to compute an hypothesis of the desired accuracy.) On the next block, as the $(i+1)$ st hypothesis is being computed, h_i is used to make predictions; at the end of this block, h_i is discarded as h_{i+1} takes its place.

Thus, if the m th trial occurs in the i th block (i.e., if $a_i \leq m \leq b_i$), then the probability of a mistake is bounded by $\epsilon(s_{i-1})$, the error rate of h_{i-1} . From the definition of $\epsilon()$, this implies the desired bound on the probability of a mistake on the m th trial, and, in turn, on the expected number of mistakes on the first m trials.

Finally, note that on the i th block, space is needed only to store the hypothesis from the last block h_{i-1} , and to simulate A 's computation of block i 's hypothesis. By Theorem 6.1, both of these quantities grow polynomially in $\log(1/\epsilon)$. By our choice of ϵ , this implies the desired bound on the algorithm's space efficiency. The time complexity of the procedure is bounded in a similar fashion. ■

7 Conclusions and Open Problems

We have shown that a model of learnability in which the learner is only required to perform slightly better than guessing is as strong as a model in which the learner's error can be made arbitrarily small. The proof of this result was based on the filtering of the distribution in a manner causing the weak learning algorithm to eventually learn nearly the entire distribution.

We have also shown this proof implies a set of general bounds on the complexity of PAC-learning (both batch and on-line), and have discussed some of the applications of these bounds.

It is hoped that these results will open the way on a new method of algorithm design for PAC-learning. As previously mentioned, the vast majority of currently known algorithms work by finding a hypothesis consistent with a large sample. An alternative approach suggested by the main result is to seek instead a hypothesis covering slightly more than half the distribution. Perhaps, such an hypothesis is easier to find, at least from the point of view of the algorithm designer. This approach leads to algorithms with a flavor similar to the one described for k -term DNF in Section 5.3, and it is possible to find similar algorithms for a number of other concept classes that are already known to be learnable (for example, k -decision lists [20] and rank r decision trees [7]). To what extent will this approach be fruitful for other classes not presently known to be learnable? This is an open question.

Another open problem concerns the robustness of the construction described in this paper. Intuitively, it seems that there should be a close relationship between reducing the error of the hypothesis, and overcoming noise in the data. Is this a valid intuition? Can our construction be modified to handle noise?

Finally, turning away from the theoretical side of machine learning, we can ask how well would our construction perform in practice? Often, a learning program (for instance, a neural network) is designed, implemented, and found empirically to achieve a "good" error rate, but no way is seen of improving the program further to enable it to achieve a "great" error rate. Suppose our construction is implemented on top of this learning program. Would it help? This is not a theoretical question, but one that can only be answered experimentally, and one that obviously depends on the domain and the underlying learning program. Nevertheless, it seems plausible that the construction might in some cases give good results in practice.

Acknowledgments

Thanks to Sally Goldman, Michael Kearns, Joe Killian and Ron Rivest for their helpful comments and suggestions.

References

- [1] Dana Angluin. Finding patterns common to a set of strings. *Information and Control*, 39:337-

- [2] Dana Angluin. *Types of queries for concept learning*. Technical Report YALEU/DCS/TR-479, Yale University Department of Computer Science, June 1986.
- [3] Eric B. Baum. On learning a union of half spaces. May 1989. (Unpublished draft).
- [4] Anselm Blumer, Andrzej Ehrenfeucht, David Haussler, and Manfred K. Warmuth. Learnability and the Vapnik-Chervonenkis dimension. *Journal of the ACM*. To appear.
- [5] Anselm Blumer, Andrzej Ehrenfeucht, David Haussler, and Manfred K. Warmuth. Occam's razor. *Information Processing Letters*, 24:377-380, April 1987.
- [6] Stéphane Boucheron and Jean Sallantin. Some remarks about space-complexity of learning, and circuit complexity of recognizing. In *Proceedings of the 1988 Workshop on Computational Learning Theory*, pages 125-138, Cambridge, MA, August 1988.
- [7] Andrzej Ehrenfeucht and David Haussler. Learning decision trees from random examples. In *Proceedings of the 1988 Workshop on Computational Learning Theory*, pages 182-194, Cambridge, MA, August 1988.
- [8] Sally Floyd. Space-bounded learning and the Vapnik-Chervonenkis dimension. In *Proceedings of the Second Annual Workshop on Computational Learning Theory*, pages 349-364, Santa Cruz, CA, July 1989.
- [9] David Haussler. *Space Efficient Learning Algorithms*. Technical Report UCSC-CRL-88-2, University of California Santa Cruz, September 1985. (revised March 1988).
- [10] David Haussler, Michael Kearns, Nick Littlestone, and Manfred K. Warmuth. Equivalence of models for polynomial learnability. In *Proceedings of the 1988 Workshop on Computational Learning Theory*, pages 42-55, Morgan-Kaufmann, Cambridge, MA, August 1988.
- [11] David Haussler, Nick Littlestone, and Manfred K. Warmuth. Expected mistake bounds for on-line learning algorithms. April 1987. (Unpublished).
- [12] David Haussler, Nick Littlestone, and Manfred K. Warmuth. Predicting $\{0, 1\}$ -functions on randomly drawn points. In *Proceedings of the Twenty-Ninth Annual Symposium on Foundations of Computer Science*, pages 100-109, White Plains, NY, 1988.
- [13] David Helmbold, Robert Sloan, and Manfred K. Warmuth. Learning nested differences of intersection-closed concept classes. In *Proceedings of the Second Annual Workshop on Computational Learning Theory*, pages 41-56, Santa Cruz, CA, July 1989.
- [14] Michael Kearns. *The Computational Complexity of Machine Learning*. PhD thesis. Harvard University Center for Research in Computing Technology, May 1989. Technical Report TR-13-89.
- [15] Michael Kearns. Thoughts on hypothesis boosting. December 1988. (Unpublished).
- [16] Michael Kearns, Ming Li, Leonard Pitt, and Leslie Valiant. On the learnability of boolean formulae. In *Proceedings of the Nineteenth Annual ACM Symposium on Theory of Computing*, pages 285-295, New York, New York, May 1987.

- [17] Michael Kearns and Leslie G. Valiant. Cryptographic limitations on learning boolean formulae and finite automata. In *Proceedings of the Twenty-First Annual ACM Symposium on Theory of Computing*, pages 433-444, Seattle, Washington, May 1989.
- [18] Michael Kearns and Leslie G. Valiant. *Learning Boolean Formulae or Finite Automata is as Hard as Factoring*. Technical Report TR 14-88, Harvard University Aiken Computation Laboratory, 1988.
- [19] Leonard Pitt and Leslie G. Valiant. Computational limitations on learning from examples. *Journal of the ACM*, 35(4):965-984, 1988.
- [20] Ronald L. Rivest. Learning decision lists. *Machine Learning*, 2(3):229-246, 1987.
- [21] Robert E. Schapire. Pattern languages are not learnable. 1989. (Unpublished).
- [22] Robert H. Sloan. Some notes on Chernoff bounds. 1987. (Unpublished).
- [23] Leslie G. Valiant. A theory of the learnable. *Communications of the ACM*, 27(11):1134-1142, November 1984.

OFFICIAL DISTRIBUTION LIST

Director 2 copies
Information Processing Techniques Office
Defense Advanced Research Projects Agency
1400 Wilson Boulevard
Arlington, VA 22209

Office of Naval Research 2 copies
800 North Quincy Street
Arlington, VA 22217
Attn: Dr. Gary Koop, Code 433

Director, Code 2627 6 copies
Naval Research Laboratory
Washington, DC 20375

Defense Technical Information Center 12 copies
Cameron Station
Alexandria, VA 22314

National Science Foundation 2 copies
Office of Computing Activities
1800 G. Street, N.W.
Washington, DC 20550
Attn: Program Director

Dr. E.B. Royce, Code 38 1 copy
Head, Research Department
Naval Weapons Center
China Lake, CA 93555