

Lecture 3: October 31

*Lecturer: Yishay Mansour Scribe: Ilan Cohen, Yaron Margalit, Alex Zhicharevich*¹

3.1 The PAC Model

In this lecture we will talk about the PAC model. The PAC learning model is one of the important and famous learning model. PAC stands for *Probably Approximately Correct*, our goal is to learn a hypothesis from a hypothesis class such that in high confidence we will have a small error rate (approximately correct). We start the lecture with an intuitive example to explain the idea behind the PAC model and the differences between the PAC model and the Bayesian learning that we studied last week. Then, we continue to formal the PAC model, review Occam's Razor principle and give a few examples.

3.2 An intuitive example

Suppose we want to predict what is a 'typical person'. Our input is a sample of different descriptions of people based on their height and weight and their label: whether it is a typical person ('+') or not ('-'). Let H be our hypothesis class. In our example, H can be the set of all possible axis-aligned rectangles. One example of hypothesis can be: mark an example to be a typical person ('+') if its description is in the range of:

$$1.60 \leq \text{height} \leq 1.90, 60 \leq \text{weight} \leq 90$$

Assuming the real target function is also rectangle. Our goal is to find the best rectangle ' R ' that approximates the real rectangle target R .

In this example, it is hard to learn directly the underlying probability of typical people distribution as we did in the Bayesian inference. Finding the common distribution of heights and weight can be very difficult without any knowledge about that distribution. As in this example and many other examples, we will try to learn an accurate predictor which will optimize our interests. For instance, in this example our interests is to find the best rectangle. This learning model is different than the Bayesian inference where we assumed that the underlying distribution has a specific form and our goal is to estimate this distribution.

Generally, in the PAC model, we won't impose any assumptions on the underlying distribution of the data other than that such a distribution exists and the data is independently

¹Based on scribes written by Yoav Giora and Omer Meshar (April 20, 2002)

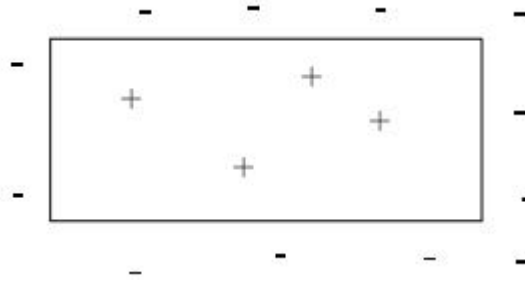


Figure 3.1: A rectangle with positive and negative examples

and identically distributed (i.i.d) according to the distribution. This assumption gives us hope that what we learn based on the training set, gives good results that are close to the real target function.

The learning problem we described can be viewed as follows:

- Goal: Learn rectangle R .
- Input: Examples based on data set and their label: $\langle(x, y), +/ -\rangle$.
- Output: R' , a good approximation rectangle of the real target rectangle R . (An example of R' , see Figure 3.1)

3.2.1 A Good Hypothesis

Our goal is to find a hypothesis that will have a small error rate, smaller than a defined ε . Let $R\Delta R'$ be the error of R' in respect to the real target rectangle R . This can be defined by two separate areas: $(R - R') \cup (R' - R)$ as shown in Figure 3.2.

Then, our goal can be defined as to find R' with probability of at least $1 - \delta$:

$$\Pr[\text{error}] = \mathcal{D}(R\Delta R') \leq \varepsilon$$

Assuming that R is the real target function.

3.2.2 Learning Strategy

Let $S = \{\langle(x_1, y_1), b_1\rangle, \dots, \langle(x_m, y_m), b_m\rangle\}$ be the sample data. Intuitive, we would like a rectangle that will be *consistent*, i.e., no error on the sample data. This is possible since we assumed that there exists a rectangle that labels correctly the data (the target rectangle

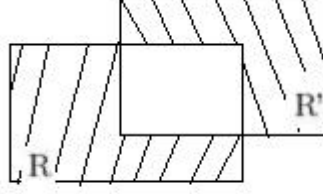


Figure 3.2: R and R' areas including two different error spaces. In our example, since $R' \subseteq R$ there exists only error of $(R - R')$.

R). We can choose any rectangle varies from the minimal size (R_{min}) up to the maximize size (R_{max}) as long it is *consistent* with the sample data. We will show later, that it doesn't matter which one of these rectangles the algorithm returns, it will be sufficient to be learnable.

The strategy A we will try is to request a "sufficiently large" number of m examples, to choose the hypothesis rectangle R' that is the tightest fit to the positive examples. That can be done by using four positive points indicating the left-most, right-most, down-most, up-most points in the sample data that are positive.

3.2.3 Sample Size

In this section, we try to find what is a "sufficiently large" number of examples that is needed to learn a good hypothesis. For that, we will fix our accuracy and confidence parameters (ε, δ) , and the strategy A . We will show that for any distribution D , we can assert sample size m that with high confidence (that is, probability at least $1 - \delta$), the returned rectangle R' from strategy A (i.e. the tightest fit rectangle) has an error of at most ε .

We will be aware of the fact that $R' \subseteq R$ and we construct $T'_1, \dots, T'_4$ areas that will indicate the error area, as defined in Figure 3.4. That is, $R \Delta R' = \cup T'_i$.

If the distribution of having each $T'_1, \dots, T'_4$ is $D(T'_i) \leq \frac{1}{4}\varepsilon$, then the error rate of R' is at most: $\Pr[error] = \mathcal{D}(R \Delta R') = \mathcal{D}(\cup T'_i) \leq \sum_{i=1}^4 \frac{\varepsilon}{4} = \varepsilon$ But this analysis is incorrect. This is because our analysis depends on the strip constructed T'_i that depended on the sample which defines the returned rectangle R' from strategy A . We need to create event which do not depend on the sample and use them on the proof.

A better construction will depend only on the real target rectangle R . We will construct strips $T_1, \dots, T_4$ such that $\forall i D(T_i) \leq \frac{1}{4}\varepsilon$ (Figure 3.5). From the construction we can see that T_i is independent of the sample and of R' . Note that we cannot certainly find the T_i but we can be sure that such T_i exists.

We would want to have $\forall i T'_i \subseteq T_i$. If that is the case, we obtained our requirement since

$$\Pr[error] = \mathcal{D}(R \Delta R') = \mathcal{D}(\cup T'_i) \leq \mathcal{D}(\cup T_i) \leq \varepsilon$$

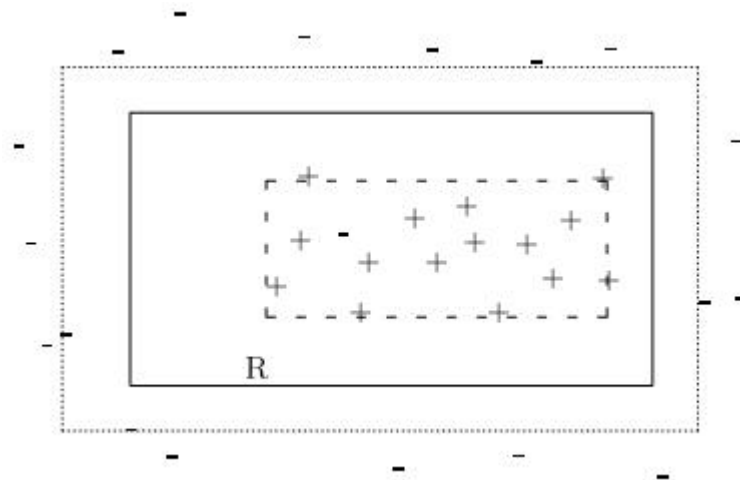


Figure 3.3: The smallest R_{min} (dashed line) and largest R_{max} (dotted line) rectangles border the rectangles which are consistent with the examples. The area between them represents the error region.

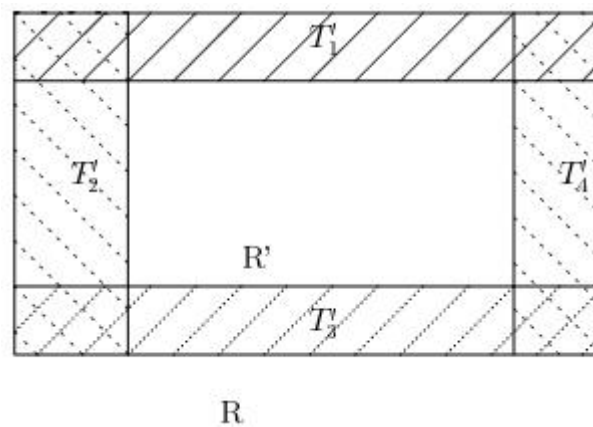


Figure 3.4: Breaking up the error into four rectangular strips $T'_1, \dots, T'_4$

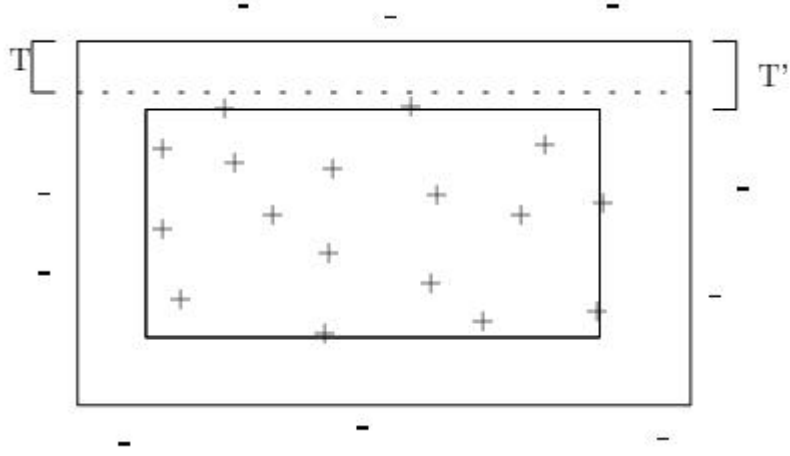


Figure 3.5: Adjusting strip size to have a weight of at most ε according to the real target function R

From the construction, if there is at least one sampled point that resides in T_i it implies that $T'_i \subseteq T_i$. This is true, since the rectangle from strategy A must include all sampled positive points in R . To achieve that we can ask: what is the probability of having a bad event, that is, what is the probability that we didn't receive points from the sample data that are located on the constructive strips, i.e., T_i . Formally,

$$\Pr[\text{bad event}] = \Pr[\exists i = 1..4 \forall x_i \in S, x_i \notin T_i]$$

By definition of T_1 ,

$$\Pr[x \notin T_1] = 1 - \frac{\varepsilon}{4}$$

Since our sample data is i.i.d from distribution D :

$$\Pr[\forall x_i \in S, x_i \notin T_1] = (1 - \frac{\varepsilon}{4})^m$$

The same analysis holds on each T_i strip. Hence, we get that on the entire region the error would not exceed the sum of probabilities for each of the strips. That is,

$$\Pr[\text{bad event}] \leq 4(1 - \frac{\varepsilon}{4})^m$$

From the inequality $(1 - x) \leq e^{-x}$, we obtain:

$$\Pr[\text{bad event}] \leq 4(1 - \frac{\varepsilon}{4})^m \leq 4e^{-\frac{\varepsilon}{4}m} < \delta$$

That is, if we want to have accuracy ε and confidence of at least $1 - \delta$, we have to choose the sample size m to satisfy:

$$4e^{-\frac{\varepsilon}{4}m} < \delta \iff m > \frac{4}{\varepsilon} \ln \frac{4}{\delta}$$

For this strategy A , and for every small (ε, δ) we like, we got the sample size that is needed for having a good learner.

3.2.4 Remarks

1. The analysis holds for any fixed probability distribution $\mathcal{D}$, we only required that the sample points are i.i.d from distribution $\mathcal{D}$ to obtain our bound.
2. The lower bound of the sample size behaves as we might expect. One might want to have better accuracy by decreasing ε or greater confidence by decreasing δ — our algorithm requires more examples to meet those requirements.
3. The parameter ε gives the degree of accuracy that we want to achieve. It determines what is a good hypothesis for achieving a good approximation in respect to the target function. In our example, the accuracy determines which of our hypothesis rectangles are good enough in respect to the real rectangle target. We pay attention that the accuracy does not depend on the data distribution.
4. The parameter δ gives the degree of confidence on having a good learner. Meaning, how sure we are that we've reached that level of accuracy. This can be related on, how typical the given sample data reflects the true distribution. Again, it does not depend on the data distribution.
5. An example of learning, we might have cases as shown in Figure 3.6, where the distribution $\mathcal{D}$ gives large weights to particular regions of the plain, creating a distorted image of the rectangle. In any case, under those conditions, since the learner is tested only through the distribution $\mathcal{D}$, and this distribution has small error between R and R' , the rectangle $\mathcal{R}'$ will be a good hypothesis (in respect to ε, δ).
6. The strategy A that we defined is efficient: In computational view, the only need is to search for the max and min points that defines our tightest-fit rectangle. In sample data size view, the number of examples that is required for achieving accuracy ε with confidence $1 - \delta$ is polynomial in $\frac{1}{\varepsilon}$ and $\ln \frac{1}{\delta}$.
7. In this example, as opposed to the Bayesian approach, we haven't been trying to model $\mathcal{D}$ or to guess which rectangle is more likely (prior). We have separated the distribution $\mathcal{D}$ from the target function (rectangle R), and directly try to predict hypothesis for this function.

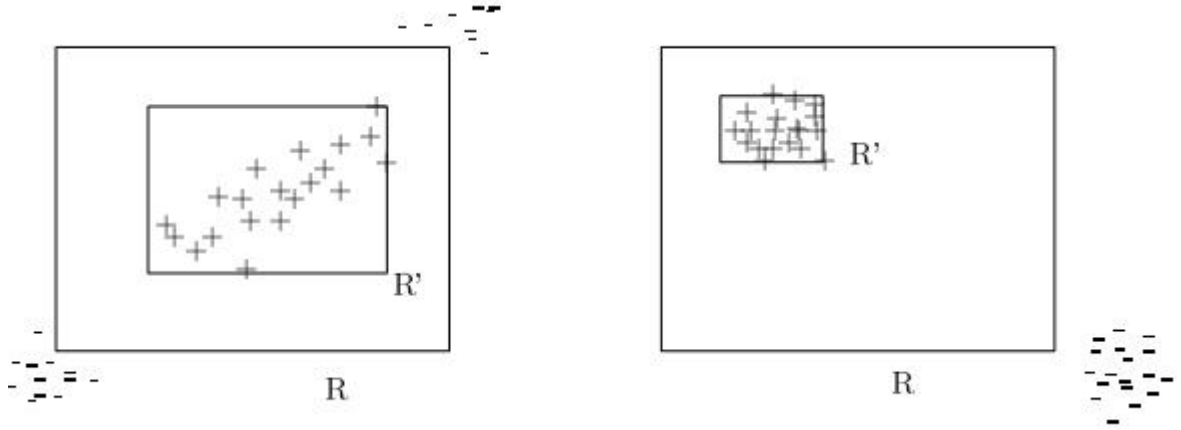


Figure 3.6: Two cases depending on the sample size

3.3 A formal Presentation of the PAC Model

3.3.1 Preliminaries

- The goal of the learning algorithm is to learn an unknown target function concept out of a known hypothesis class, for example to learn one particular rectangle out of the set of all possible rectangles. Unlike the Bayesian approach, no prior knowledge is needed.
- Learning occurs in a stochastic setting. Examples from the target rectangle are drawn randomly according to a fixed, unknown probability distribution and are i.i.d.
- We assume that the sample and test data are generated by the same unknown probability distribution.
- The solution should be efficient: the sample size required for obtaining small (ε) error with high $(1 - \delta)$ confidence is a function of $\frac{1}{\varepsilon}$ and $\ln \frac{1}{\delta}$. Also, we can process the sample within a time polynomial in the sample size.

3.3.2 Definition of the PAC Model

Let the set X be the *instance space* or *example space*. Let $\mathcal{D}$ be any fixed probability distribution over the instance space X . A *concept class* over X is a set

$$\mathcal{C} \subseteq \{c \mid c : X \rightarrow \{0, 1\}\}.$$

Let c_t be the *target concept*, $c_t \in \mathcal{C}$, and h a *hypothesis*, h may belong to a different concept class than $\mathcal{C}$, the *hypothesis* $\mathcal{H}$. We will define the *error* of h with respect to the

distribution $\mathcal{D}$ and the target concept c_t as follows:

$$\text{error}(h) = \Pr_{\mathcal{D}}[h(x) \neq c_t(x)] = D(h\Delta_{c_t}(x))$$

Let $EX(c_t, \mathcal{D})$ be a procedure (we will sometimes call it an *oracle*) that runs in a unit time, and on each call returns a labeled example $\langle x, c_t(x) \rangle$, where x is drawn independently from $\mathcal{D}$. In the PAC model, the oracle is the *only* source of examples for the learning algorithm.

Definition Let $\mathcal{C}$ and H be concept classes over X . We say that $\mathcal{C}$ is *PAC learnable* by H if there exists an algorithm A with the following property: for every concept $c_t \in \mathcal{C}$, for every distribution $\mathcal{D}$ on X , and for all $0 < \varepsilon, \delta < \frac{1}{2}$, if A is given access to $EX(c_t, \mathcal{D})$ and inputs ε and δ , then with probability at least $1 - \delta$, A outputs a hypothesis concept $h \in H$: If $c_t \in H$ (Realizable case), then h is satisfying

$$\text{error}(h) \leq \varepsilon$$

If $c_t \notin H$ (Unrealizable), then h is satisfying

$$\text{error}(h) \leq \varepsilon + \min_{h' \in H} \text{error}(h')$$

we say that $\mathcal{C}$ is *efficiently PAC learnable*, if A runs in time polynomial in $\frac{1}{\varepsilon}$, $\ln \frac{1}{\delta}$, n and m when n is the size of the input and m the size of the target function, (for example, the number of bits that are needed to characterize it). Implicit, we mean that it is "easier" to learn a "simpler" target function (This idea is further explained later in the lecture).

3.4 Finite Hypothesis Class

In this section, we will show how to learn a good hypothesis from a finite hypothesis class $\mathcal{H}$. The prove "trick" will use the idea of analyzing a bad hypothesis: we define a hypothesis h to be ε -Bad if $\text{error}(h) > \varepsilon$.

3.4.1 The Case $c_t \in \mathcal{H}$

Generally, after having processed $m(\varepsilon, \delta)$ instances, we will request from our algorithm A to find an h which is *consistent* with the sample S i.e., classifies all the instances the same as the target concept (Formally, $\forall x \in S, h(x) = c_t(x)$). The algorithm will succeed to learn if h is not ε -Bad, that is, if $\text{error}(h) < \varepsilon$. We note that at least one *consistent* hypothesis exists because $c_t \in \mathcal{H}$.

Now, we will try to bound the probability of algorithm A to return h that is ε -Bad. From our defined algorithm A , having a ε -Bad hypothesis is bound by the probability that there exists a concept h which is both consistent and ε -Bad. First, we will look at a fixed h ε -Bad:

$$\Pr[h \text{ is } \varepsilon\text{-Bad} \ \& \ h(x_i) = c_t(x_i) \text{ for } 1 \leq i \leq m] \leq (1 - \varepsilon)^m < e^{-\varepsilon m}$$

First inequality derived directly from $\Pr[h(x) = c_t(x)] \leq 1 - \varepsilon$ and the fact that the examples are sampled i.i.d from distribution $\mathcal{D}$

Now, we bound the failure probability of algorithm A , that is by:
 $\Pr [A \text{ returns an "}\varepsilon\text{-Bad" hypothesis}] =$

$$\begin{aligned} & \Pr[\exists h \in \varepsilon\text{-Bad} \ \& \ h(x_i) = c_t(x_i) \text{ for } 1 \leq i \leq m] \\ & \leq \sum_{h \in \varepsilon\text{-Bad} \ \& \ h \in H} \Pr[h(x_i) = c_t(x_i) \text{ for } 1 \leq i \leq m] \\ & \leq |\{h : h \text{ is } \varepsilon\text{-Bad} \ \& \ h \in H\}| (1 - \varepsilon)^m \\ & \leq |\mathcal{H}| (1 - \varepsilon)^m < |\mathcal{H}| e^{-\varepsilon m}. \end{aligned}$$

The first inequality comes directly from the union bound. The second inequality is true by the previous analysis for ε -Bad, and the third inequality is derived from the fact:

$$|\{h : h \text{ is } \varepsilon\text{-Bad} \ \& \ h \in H\}| \subseteq H$$

In order to satisfy the condition for PAC learning, we bound the failure probability by δ :

$$|\mathcal{H}| e^{-\varepsilon m} \leq \delta,$$

which implies a bound on the sample size,

$$m \geq \frac{1}{\varepsilon} \ln \frac{|\mathcal{H}|}{\delta}.$$

The algorithm A is general and achieves the desired results. As expected it tells us that the larger the hypothesis class, the larger the sample we require. However it is not clear how to implement it efficiently, and in general it will not be efficient. Also, it only works for finite classes and even in this lecture we saw an example that doesn't apply for this analysis — the class of all rectangles.

3.4.2 The Case $c_t \notin \mathcal{H}$

In this case, we need to relax our goal, since a small error hypothesis might not even exist (in H). As we defined before, our goal is to learn according to the "best" hypothesis in the hypothesis class (i.e. with the minimal error). Let h^* be the hypothesis with the minimal error, i.e., for every $h \in \mathcal{H}$:

$$0 < error(h^*) \leq error(h).$$

Let β to be $\beta = error(h^*) = \min_{h \in H} error(h)$. Our goal is to find hypothesis h that will achieve:

$$error(h) \leq \beta + \varepsilon$$

Note that this is in fact a generalization of our previous system (in which $\beta = 0$).

After having processed $m(\varepsilon, \delta)$ instances, we define $\widehat{error}(h)$ to be the empirical error of h on sample S . Formally:

$$\widehat{error}(h) = \frac{1}{m} \sum_{i=1}^m I(h(x_i) \neq c_t(x_i)),$$

where I is the indicator function. We also need to correct the algorithm from the realizable case, since now it may be impossible to choose a consistent h . Algorithm A will return a hypothesis $\bar{h}$ that will have the minimal empirical error. That is, $\bar{h} = \operatorname{argmin}_{h \in H} \widehat{error}(h)$ (If there exists more than one hypothesis, the algorithm will pick one of them).

We will now bound the sample size required for obtaining a good hypothesis from algorithm A . We want to choose a sample size m such that the difference between the true and the estimated error is small for every hypothesis. That is, with probability at least $1 - \delta$:

$$\forall h \in H, |\widehat{error}(h) - error(h)| \leq \frac{\varepsilon}{2}.$$

If we have a good error estimation, we can derive easily that using the hypothesis $\bar{h}$ from algorithm A , we obtain:

$$error(\bar{h}) \leq \widehat{error}(\bar{h}) + \frac{\varepsilon}{2} \leq \widehat{error}(h^*) + \frac{\varepsilon}{2} \leq error(h^*) + \varepsilon$$

The first and third inequality holds since the difference between the true and the estimated error is small (up to $\frac{\varepsilon}{2}$). The second inequality holds since the empirical error of the hypothesis obtained from algorithm A is the minimal one.

To conclude our lower bound for the sample size we need to find the probability of failure in estimating $error(h)$, we will use Chernoff bounds that we introduced in the first lecture:

$$\Pr[|\widehat{error}(h) - error(h)| \geq \frac{\varepsilon}{2}] \leq 2e^{-2(\frac{\varepsilon}{2})^2 m}$$

We can use Chernoff bounds since the data m is i.i.d from the same distribution $\mathcal{D}$. We can get a lower bound on the sample size by requiring the probability of failure over $\mathcal{H}$ to be smaller than δ :

$$\Pr[\exists h \in \mathcal{H} : |\widehat{error}(h) - error(h)| \geq \frac{\varepsilon}{2}] \leq 2|\mathcal{H}|e^{-2(\frac{\varepsilon}{2})^2 m} \leq \delta,$$

hence

$$m \geq \frac{2}{\varepsilon^2} \ln \frac{2|\mathcal{H}|}{\delta}$$

We have thus found a lower bound on the sample size for PAC learning when $c_t \notin \mathcal{H}$. Note that the sample size depends on $\frac{1}{\varepsilon^2}$ instead of $\frac{1}{\varepsilon}$ when we had $c_t \in \mathcal{H}$. This results from the difference between needing a single counter-example (in which $error(h) = 0$), to the need of having the average over many examples to disqualify a function, by estimating the error on the data.

3.4.3 Example - Learning Boolean Disjunctions

To demonstrate the PAC model we will look at the example of learning boolean disjunction functions. The problem is defined as follows: Given a set of boolean variables $T = \{x_1, \dots, x_n\}$ and a set of literals $L = x_1, \bar{x}_1, \dots, x_n, \bar{x}_n$. we need to learn an Or function over the literals, for example: $x_1 \vee \bar{x}_3 \vee x_5$. C will be the set of all possible disjunctions. We can notice that $|C| = 3^n$, since for each variable x_i the target disjunction c_t may contain x_i , $\bar{x}_i$, or neither.

We will use $H=C$.

ELIM algorithm for learning boolean disjunctions

We can notice that by receiving a negative example we can eliminate all literals that give 1 to the variables in the sample. For example if there was a sample 001 (assume 0 was the value of x_1 and x_2 and 1 was the value of x_3) which was negative we can eliminate the literals $\bar{x}_1, \bar{x}_2, x_3$ because if one of them was in the target disjunction it would give a positive value. The algorithm uses that fact and eliminates the inconsistent literals.

So the algorithm initializes a set $L = x_1, \bar{x}_1, \dots, x_n, \bar{x}_n$ and for each negative sample Z it assigns $L = L - \{\bar{z}_i | z_i \in Z\}$

From previous sections we can see that the algorithm learns when the sample size:

$$m > \frac{1}{\varepsilon} \ln \frac{|\mathcal{H}|}{\delta} = \frac{1}{\varepsilon} \ln \frac{3^n}{\delta} = \frac{n \ln 3}{\varepsilon} + \frac{1}{\varepsilon} \ln \frac{1}{\delta}$$

3.4.4 Example - Learning parity

Consider the concept class of xor of n variables. More formally: A set of variables : $x_1, \dots, x_n$. The concepts class is the set of all XOR functions over the variables. Example for a function in C is : $x_1 \otimes x_7 \otimes x_9$.

We notice that $|C| = 2^n$ and allowing literals will not increase it since XOR of two negated variables equals to the XOR of the variable, so allowing literals will increase the size of C in a factor of 2 since it only matters how many negated literals appears in the clause. The algorithm will regard the samples as a linear equation problem (in Z_2 field) and then solve it using Gaus elimination method. Each sample consists of a bit vector (satisfying the value of the variable) and the classification c_t of this example. This creates a linear equation, for example: (01101, 1) represents the equation: $0 * a_1 + 1 * a_2 + 1 * a_3 + 0 * a_4 + 1 * a_5 = 1(mod 2)$ (deduced directly from the properties of xor). Each a_i is a binary indicator of x_i in the formula. A solution to all the examples exists (since we want to learn a xor function). Thus, solving the linear equation problem will produce a solution which is consistent with the whole sample The needed sample size m is

$$m > \frac{1}{\varepsilon} \ln \frac{|\mathcal{H}|}{\delta} = \frac{n}{\varepsilon} \ln 2 + \frac{1}{\varepsilon} \ln \frac{1}{\delta}.$$

3.5 Occam Razor

“Entities should not be multiplied unnecessarily”

(William of Occam, c. 1320)

The above quote is better known as Occam’s Razor or the principle of Ontological Parsimony. Over the centuries, people from different fields have given it different interpretations. One interpretation used by experimental scientists is: given two explanations of the data, the simpler explanation is preferable. This principle is consistent with the goal of machine learning: to discover the simplest hypothesis that is consistent with the sample data. However, the question still remains: why should one assume that the simplest hypothesis based on past examples will perform well on future examples. After all, the real value of a hypothesis is in predicting the examples that it has not yet seen. It will be shown that, under very general assumptions, Occam’s Algorithm produces hypotheses that with high probability will be predictive of future observations. As a consequence, when hypotheses of minimum or near minimum complexity can be produced in time which is polynomial in the size of the sample data, it leads to a polynomial PAC-learning algorithm. Here, a simple hypothesis, means a short coded hypothesis. We saw that we can achieve a polynomial PAC-learning algorithm, for a finite class of hypotheses, H , if we choose the number of input examples m , to be:

$$m \geq \frac{1}{\epsilon} \ln \frac{|H|}{\delta}$$

We can allow H to grow with m , providing that it grows slower than m .

Definition: An (α, β) *Occam-algorithm* for a functions class C , using a hypotheses class H , is an algorithm which has two parameters: a constant $\alpha \geq 0$ and a compression factor $0 \leq \beta < 1$. Given a sample S of size m , which is labeled according to c_t , i.e. $\forall i \in \{1, \dots, m\} \ c_t(x_i) = \sigma_i$, the algorithm outputs an hypothesis $h \in H$, such that:

1. The hypothesis h is consistent with the sample, i.e. $\forall i \in \{1, \dots, m\} \ h(x_i) = \sigma_i$.
2. The size of h is at most $n^\alpha m^\beta$, where n is the size of c_t , m is the sample size, and α and β are the algorithm's parameters.

3.5.1 Occam Algorithm and PAC

We now show that the existence of an Occam-algorithm for $\mathcal{C}$ implies polynomial PAC learnability.

Theorem 3.1 *Let A be an (α, β) Occam Algorithm for C , using H . Then A is PAC with*

$$m \geq \left(\frac{n^\alpha}{\epsilon} \ln 2\right)^{1/(1-\beta)} + \frac{2}{\epsilon} \ln \frac{1}{\delta}$$

Remark: Notice, that when $\beta \rightarrow 1$, then $m \rightarrow \infty$. This makes sense, since the closer β is to 1, the poorer the compression we get, and hence we have less “information” about the function.

Proof: Fix m and n . A returns an hypothesis h s.t. $size(h) \leq n^\alpha m^\beta$. The number of hypothesis of this size is $2^{n^\alpha m^\beta}$.

We have shown that for PAC we need $m \geq \frac{1}{\epsilon} \ln \frac{|H|}{\delta}$. Since $|H| \leq 2^{n^\alpha m^\beta}$, we get $m \geq \frac{1}{\epsilon} n^\alpha m^\beta \ln 2 + \frac{1}{\epsilon} \ln \frac{1}{\delta}$.

$$m \geq \max\left\{\frac{2}{\epsilon} n^\alpha m^\beta \ln 2, \frac{2}{\epsilon} \ln \frac{1}{\delta}\right\} \Rightarrow m \geq \frac{2}{\epsilon} n^\alpha m^\beta \ln 2 \Rightarrow m \geq \left(\frac{2}{\epsilon} n^\alpha \ln 2\right)^{\frac{1}{1-\beta}}$$

□

3.5.2 Example: Boolean OR with k Literals

As before we want to learning boolean disjunctions for n variables, but we know that the expression as at most k literals, i.e. the hypothesis C size is $n^k (\ll 3^n)$, we will show Occam algorithm that creates hypothesis h with size $O(k \ln n \ln m)$ we will use the greedy algorithm for set cover.

Algorithm 1: Greedy Algorithm for Set Cover**Input:** $S_1, S_2 \dots S_t \subseteq V = \{1 \dots m\}$ **Output:** $S_{i_1}, \dots, S_{i_k} \Rightarrow \bigcup S_{i_j} = V$ SETCOVERGREEDY(U)

- (1) $S \leftarrow \emptyset, j \leftarrow 0, V_0 \leftarrow V$
- (2) **while** $V_j \neq \emptyset$
- (3) Choose $i, S_i = \operatorname{argmax}_{S_r} \{|S_r \cap V_j|\}$
- (4) $S \leftarrow S \cup \{i\}$
- (5) $V_{j+1} \leftarrow V_j - S_i$
- (6) $j \leftarrow j + 1$
- (7) **return** S

Algorithm analyze

Let S_{opt} coverage for V with k sets, then the greedy algorithm return a cover with at most $1 + k \ln m$. The cover S_{opt} is also cover for V_j , and has just k sets, hence

$$\exists t \in S_{opt} \Rightarrow |V_j \cap S_t| \geq \frac{|V_j|}{k}$$

$$|V_{j+1}| \leq |V_j| - \frac{|V_j|}{k} = (1 - \frac{1}{k})|V_j| \text{ (We choose the maximum set)}$$

$$|V_{j+1}| \leq (1 - \frac{1}{k})^j |V_0| \leq e^{-\frac{j}{k}} m,$$

when the bound is less the 1 then $V_j = \emptyset$. Therefore, after $1 + k \ln m$ steps the algorithm will stop.

Algorithm Boolean OR with k Literals

Input S examples Run algorithm ELIM the output is L .

Notice $|L| = O(n)$, and we need $m = o(n)$

So our objective is to choose small group of literals that satisfy S (the positive subset).

Reduction to set cover:

$T = \{x : x \in S\}$ (All the positive examples)

$T_i = \{x : x \in T, x_i \in x\}$ (all the examples that literal x_i satisfy)

We know that there exists a group of k literals that satisfy the examples, so if we run greedy set cover, it guaranteed that we will get $k \ln m$ literals that satisfy all the examples.

There exists $2n$ literals that can be encoded with $O(\log 2n)$, hence

$\text{size}(h) \simeq k \ln n \ln m$, $\ln m \leq m^\beta$, $\ln n \leq n^\alpha$, so we got Occam algorithm for α, β , small as we want.