

# Statistical Machine Learning (BE4M33SSU)

## VC dimension

Czech Technical University in Prague  
V. Franc

# Vapnik-Chervonenkis Dimension

Threshold classifiers:  $\mathcal{H}_1 = \{h(x) = \text{sign}(x - \theta) \mid \theta \in \mathbb{R}\}$

Oriented threshold classifiers:  $\mathcal{H}_2 = \{\{h(x) = \text{sign}(x - \theta) \mid \theta \in \mathbb{R}\} \cup \{h(x) = \text{sign}(\theta - x) \mid \theta \in \mathbb{R}\}\}$

| #inputs<br>$m$ | possible label<br>configurations $2^m$                       | $\mathcal{H}_1 = \left\{ \begin{array}{c} \ominus \oplus \\ \longleftrightarrow \\ \hline \end{array} \right\}$ | $\mathcal{H}_2 = \left\{ \begin{array}{c} \ominus \oplus \\ \longleftrightarrow \\ \hline \end{array} \cup \begin{array}{c} \oplus \ominus \\ \longleftrightarrow \\ \hline \end{array} \right\}$ |
|----------------|--------------------------------------------------------------|-----------------------------------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| 1              | <div> <div></div> <div></div> </div>                         | <div> <div></div> <div></div> </div> <p>VCdim = 1</p>                                                           | <div> <div></div> <div></div> </div>                                                                                                                                                              |
| 2              | <div> <div></div> <div></div> <div></div> <div></div> </div> | <div> <div></div> <div></div> <div></div> <div></div> </div> <p>VCdim = 2</p>                                   | <div> <div></div> <div></div> <div></div> <div></div> </div>                                                                                                                                      |
| 3              | <div> <div></div> <div></div> <div></div> <div></div> </div> |                                                                                                                 | <div> <div></div> <div></div> <div></div> <div></div> <div></div> <div></div> <div></div> </div>                                                                                                  |

# Vapnik-Chervonenkis Dimension

- ◆ The VC dimension quantifies the complexity (capacity) of a hypothesis class  $\mathcal{H} \subseteq \{-1, +1\}^{\mathcal{X}}$ .

**Definition (Shattering):** Let  $\mathcal{H} \subseteq \{-1, +1\}^{\mathcal{X}}$  and let  $\{x_1, \dots, x_m\} \subset \mathcal{X}^m$  be a set of  $m$  input points. The set  $\{x_1, \dots, x_m\}$  is shattered by  $\mathcal{H}$  if, for every labeling  $y \in \{-1, +1\}^m$ , there exists a hypothesis  $h \in \mathcal{H}$  such that

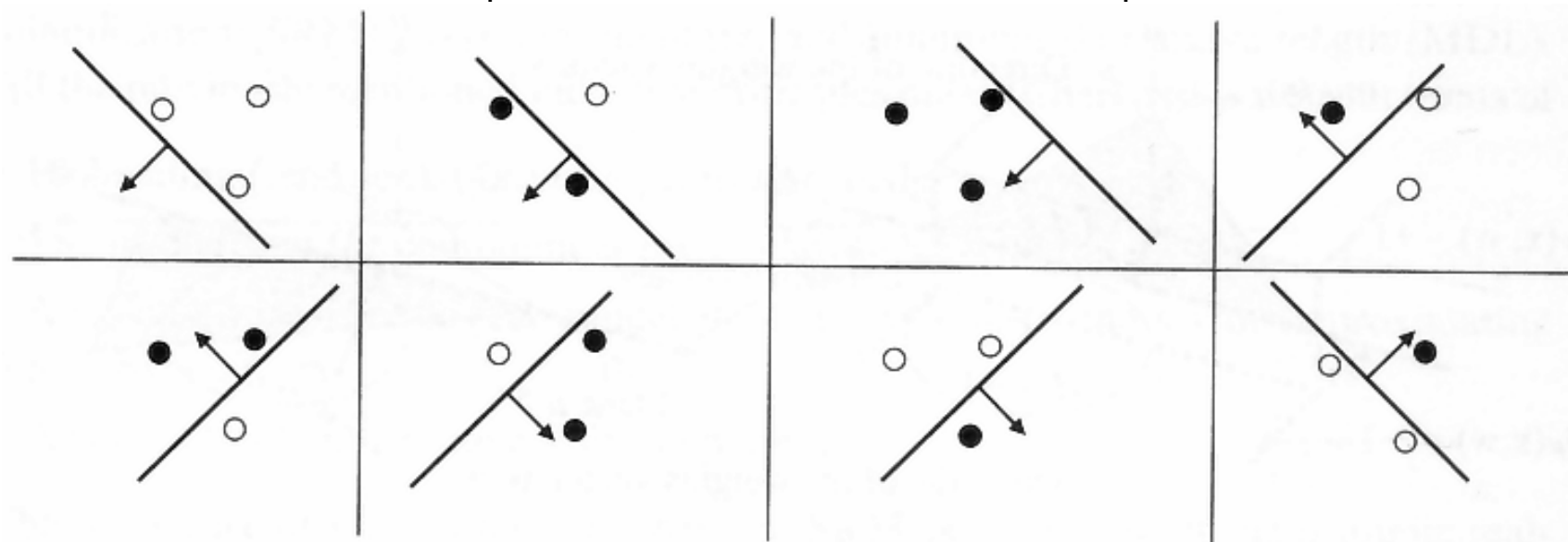
$$h(x_i) = y_i, \quad \forall i \in \{1, \dots, m\}.$$

**Definition (VC dimension):** Let  $\mathcal{H} \subseteq \{-1, +1\}^{\mathcal{X}}$ . The Vapnik–Chervonenkis dimension of  $\mathcal{H}$ , denoted  $d_{\text{VC}}(\mathcal{H})$ , is the cardinality of the largest set of points from  $\mathcal{X}$  that can be shattered by  $\mathcal{H}$ .

# Vapnik-Chervonenkis Dimension

**Theorem:** The VC-dimension of the hypothesis class of all two-class linear classifiers in  $d$ -dimensional feature space  $\mathcal{H} = \{h(x; \mathbf{w}, b) = \text{sign}(\langle \mathbf{w}, \phi(x) \rangle + b) \mid (\mathbf{w}, b) \in \mathbb{R}^{d+1}\}$  is  $d_{\text{VC}}(\mathcal{H}) = d + 1$ .

Example for  $n = 2$ -dimensional feature space



*Quiz 1:* Let  $\mathcal{H} = \{h(x) = \hat{y}, (\hat{x}, \hat{y}) = \arg \min_{(x', y') \in T_m} \|x' - x\|\}$  be a space of Nearest-Neighbor classifiers. What is the VC dimension of  $\mathcal{H}$ ?

*Quiz 2:* Let  $\mathcal{H} = \{h_1, \dots, h_K\}$  be a finite hypothesis class. What is the VC dimension of  $\mathcal{H}$ ?

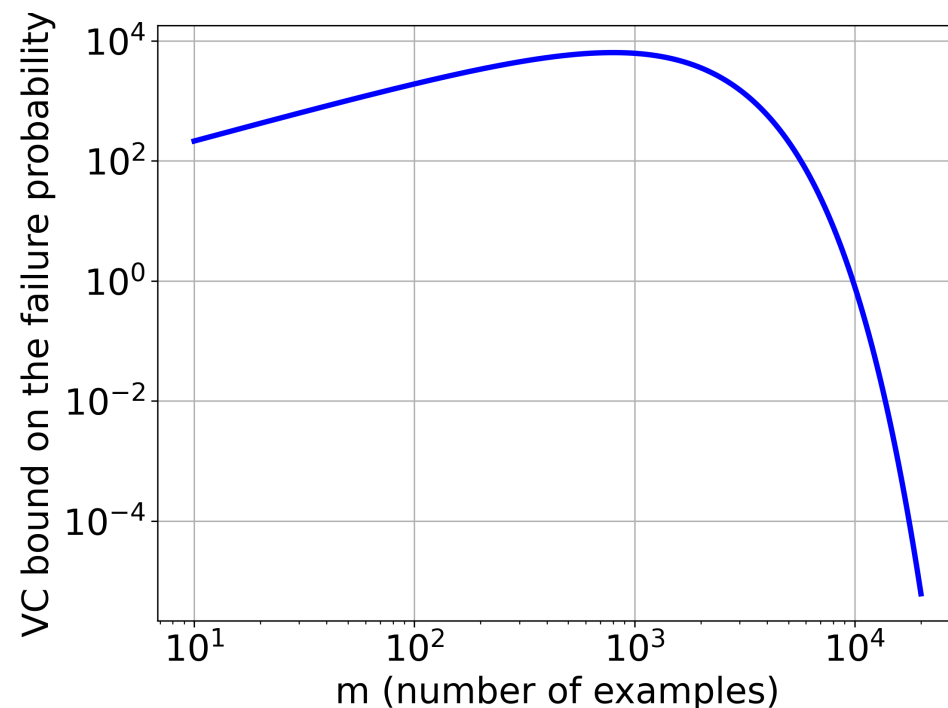
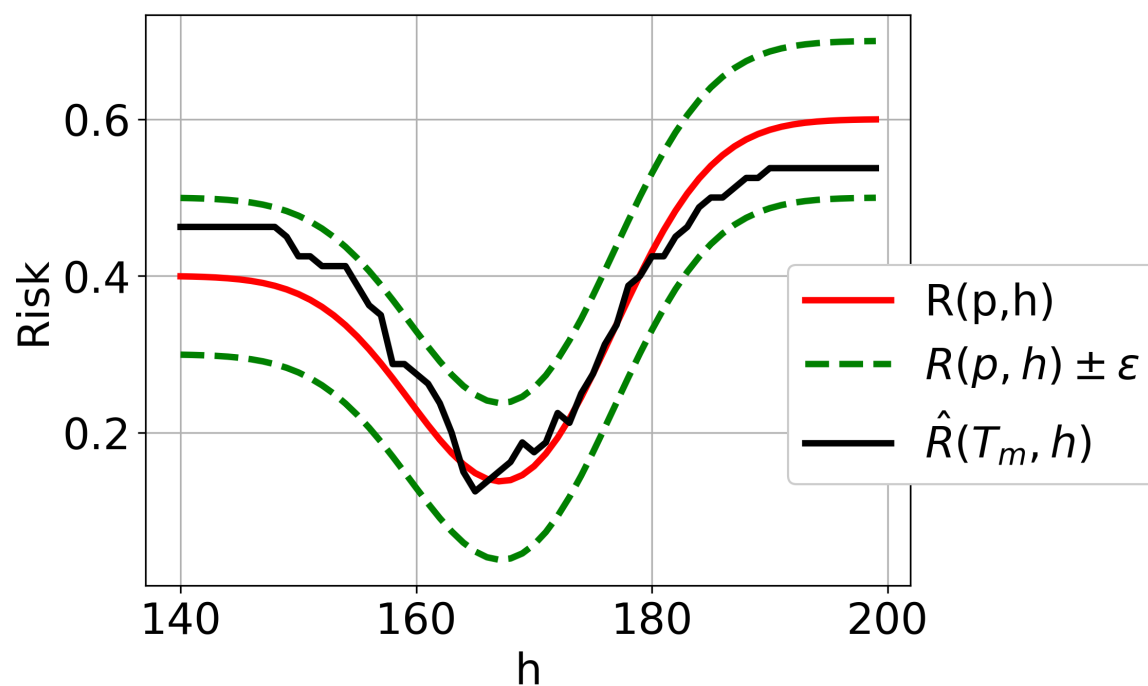
# VC Dimension Generalization Bound

**Theorem:** Let  $\mathcal{H} \subset \{+1, -1\}^{\mathcal{X}}$  be a hypothesis class with VC dimension  $d_{\text{VC}}(\mathcal{H}) < \infty$  and  $T_m = \{(x_1, y_1), \dots, (x_m, y_m)\} \in (\mathcal{X} \times \mathcal{Y})^m$  a training set i.i.d. drawn from a distribution  $p(x, y)$ . Then for any  $\varepsilon > 0$  it holds

$$\mathbb{P}\left(\max_{h \in \mathcal{H}} |R(p, h) - \hat{R}(T_m, h)| \geq \varepsilon\right) \leq 4 \left(\frac{2 e m}{d_{\text{VC}}(\mathcal{H})}\right)^{d_{\text{VC}}(\mathcal{H})} e^{-\frac{m \varepsilon^2}{8}}$$

where  $R(p, h) = \mathbb{E}_{(x, y) \sim p}[\mathbb{I}[y \neq h(x)]]$  is the true error and  $\hat{R}(T_m, h) = \frac{1}{m} \sum_{i=1}^m \mathbb{I}[y_i \neq h(x_i)]$  is the empirical error.

**Example:**  $\mathcal{H} = \{h(x; \theta) = \text{sign}(x - \theta) \mid \theta \in \mathbb{R}\}$ ,  $d_{\text{VC}}(\mathcal{H}) = 1$ ,  $\varepsilon = 0.1$



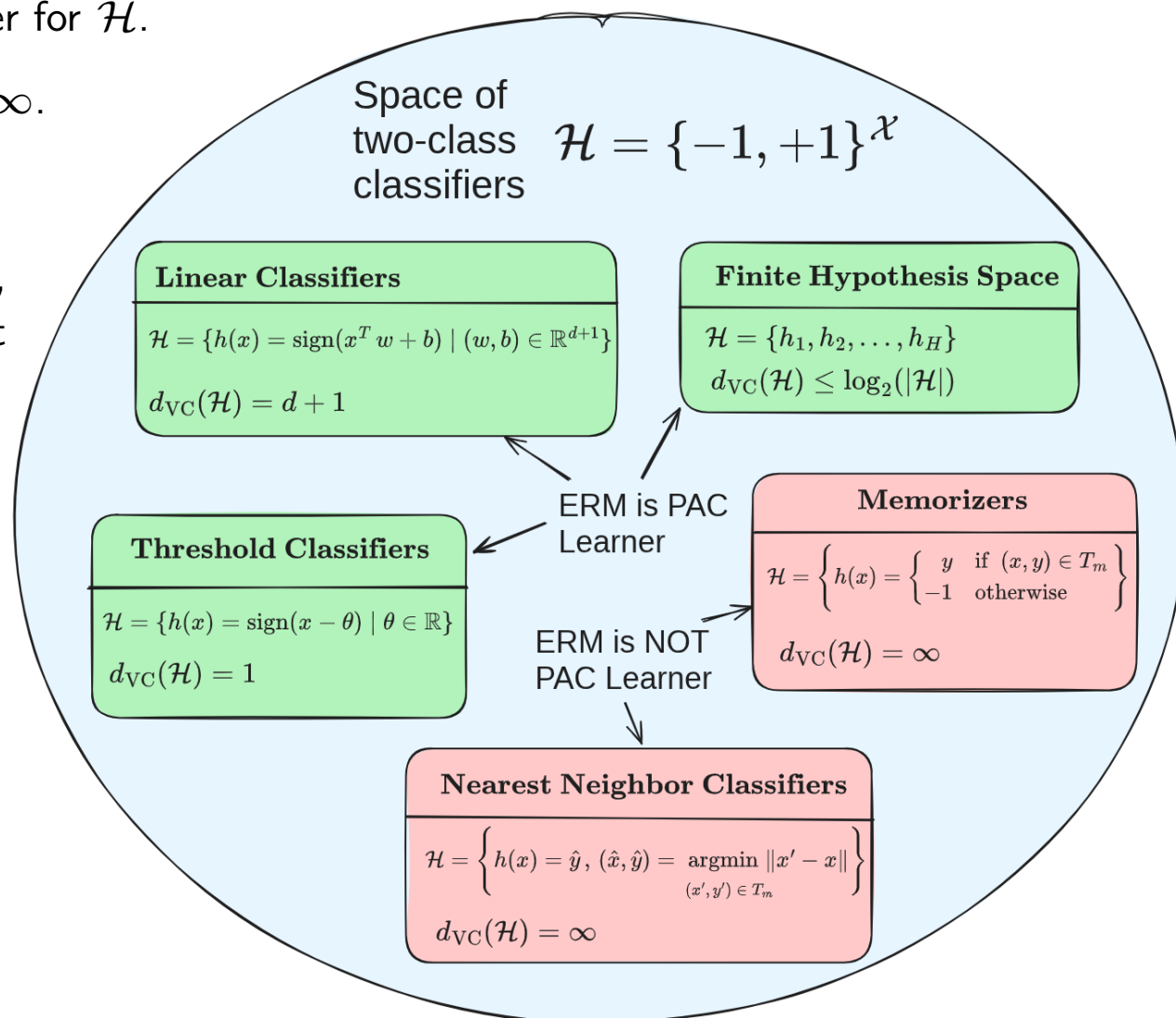
# Fundamental Theorem of PAC Learning

**Theorem:** Let  $\mathcal{H} \subset \{-1, +1\}^{\mathcal{X}}$  be a hypothesis class of functions from  $\mathcal{X}$  to  $\{-1, +1\}$  and let  $\ell(y, y') = \mathbb{I}[y \neq y']$  be the 0/1-loss function. Then, the following statements are equivalent:

- ◆ Uniform Law of Large numbers holds for  $\mathcal{H}$ .
- ◆ ERM algorithm is a successful PAC learner for  $\mathcal{H}$ .
- ◆  $\mathcal{H}$  has finite VC dimension,  $d_{\text{VC}}(\mathcal{H}) < \infty$ .

Assume the VC dimension of  $\mathcal{H}$  is finite,  $d_{\text{VC}}(\mathcal{H}) < \infty$ . Then, there is a constant  $C$  such that the sample complexity is

$$m_{\text{pac}}^{\mathcal{H}}(\varepsilon, \theta) \leq C \frac{d_{\text{VC}}(\mathcal{H}) + \log(\frac{1}{\delta})}{\varepsilon^2}$$



# VC Dimension Generalization Bound

**Theorem.** Let  $\mathcal{H} \subseteq \{-1, +1\}^{\mathcal{X}}$  be a hypothesis class of binary classifiers, and let  $\ell(y, y') = \mathbb{I}[y \neq y']$  denote the 0–1 loss. Assume that  $\mathcal{H}$  has a finite VC dimension,  $d_{\text{VC}}(\mathcal{H}) < \infty$ . Then, for any  $\delta \in (0, 1)$ , with probability at least  $1 - \delta$  over an i.i.d. training sample  $T_m = ((x_i, y_i) \in \mathcal{X} \times \mathcal{Y} \mid i = 1, \dots, m)$ , the following inequality holds for all  $h \in \mathcal{H}$  simultaneously:

$$R(h, p) \leq \underbrace{\widehat{R}(T_m, h)}_{\text{empirical risk}} + \underbrace{4 \sqrt{\frac{d_{\text{VC}}(\mathcal{H}) \log\left(\frac{2em}{d_{\text{VC}}(\mathcal{H})}\right) + \log\left(\frac{4}{\delta}\right)}{m}}_{\text{complexity term}}$$

## Practical implications of the VC bound:

- ◆ Minimize the empirical risk  $\widehat{R}(T_m, h)$  (fit the data well).
- ◆ Control the complexity term:
  - Use as many training examples  $m$  as possible.
  - Incorporate prior knowledge to restrict the complexity of  $\mathcal{H}$  (simpler models  $\Rightarrow$  smaller VC dimension).

# Structural Risk Minimization

## Algorithm:

1. Construct a nested sequence of hypothesis classes:

$$\mathcal{H}_1 \subset \mathcal{H}_2 \subset \dots \subset \mathcal{H}_K$$

2. For each  $i \in \{1, \dots, K\}$ , apply ERM:

$$h_i = \arg \min_{h \in \mathcal{H}_i} \hat{R}(T_m, h)$$

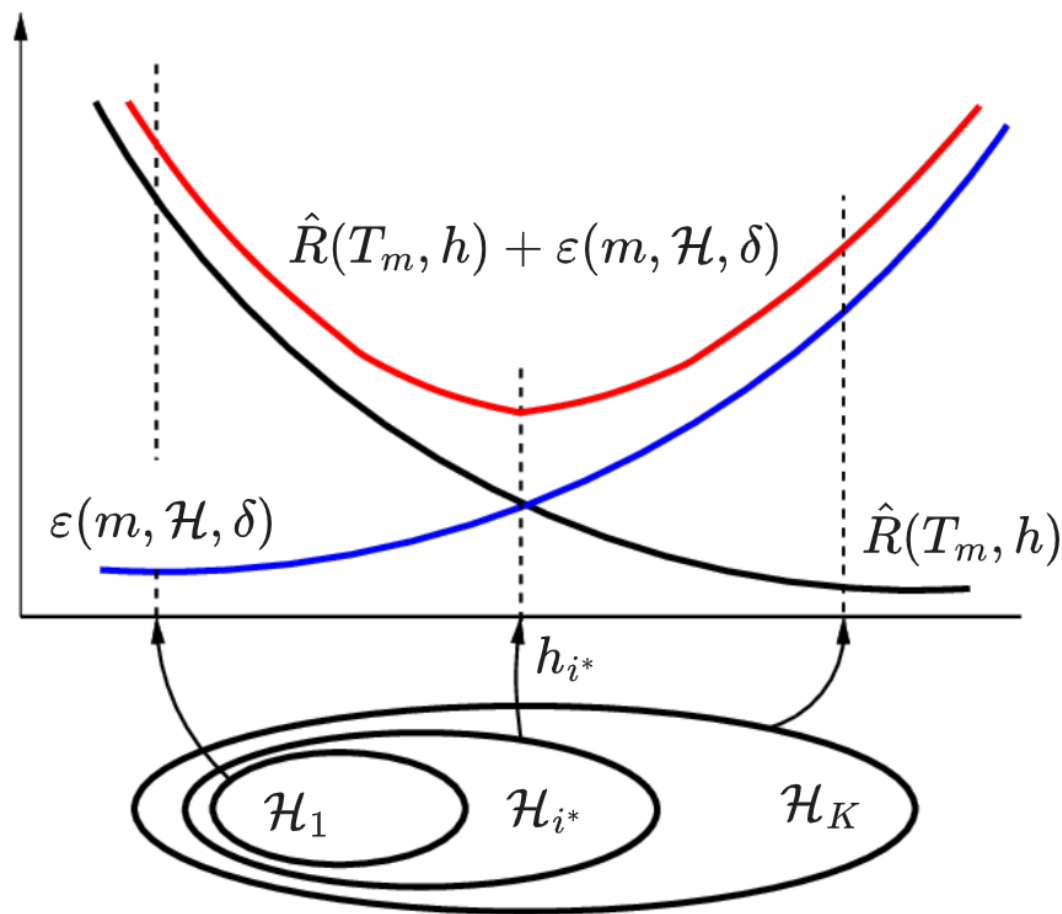
3. Select the best model using the VC generalization bound:

$$i^* = \arg \min_{i=1, \dots, K} \left( \hat{R}(T_m, h_i) + \varepsilon(m, \mathcal{H}_i, \delta) \right)$$

where

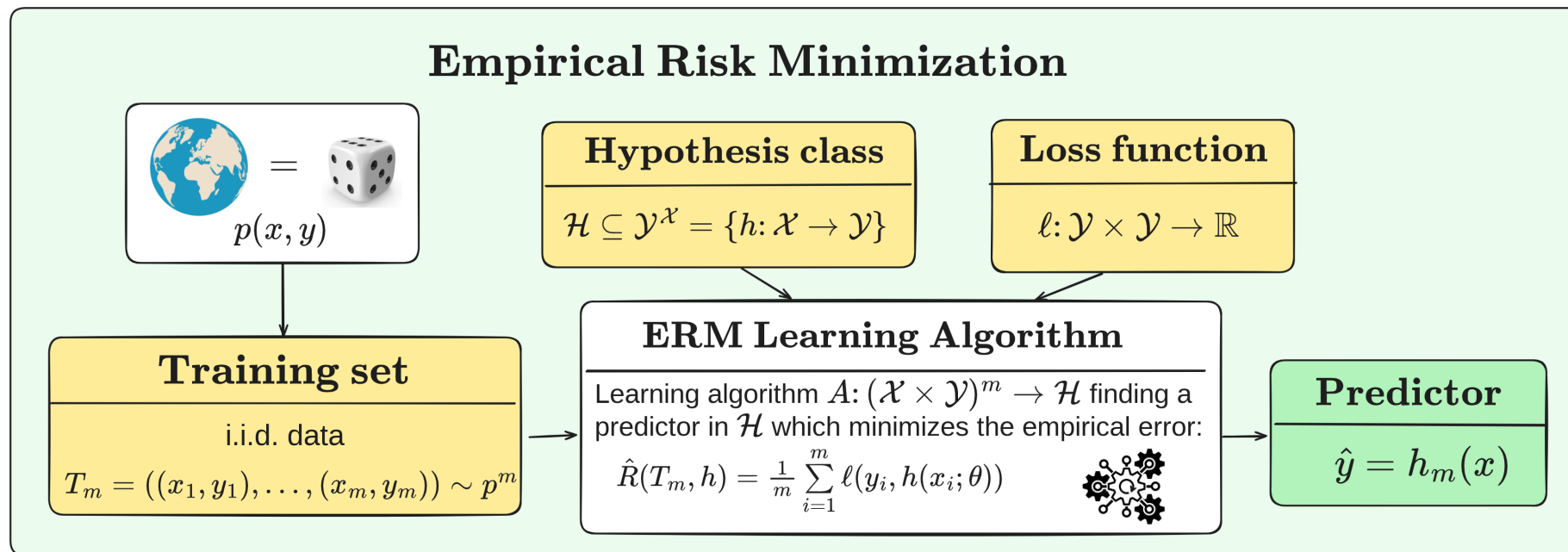
$$\varepsilon(m, \mathcal{H}_i, \delta) = 4 \sqrt{\frac{d_{\text{VC}}(\mathcal{H}) \log\left(\frac{2em}{d_{\text{VC}}(\mathcal{H})}\right) + \log\left(\frac{4}{\delta}\right)}{m}}$$

4. Output  $h_{i^*}$ .





# Summary of Key Concepts



## Error decomposition

Predictor Error = Estimation Error + Approximation Error + Bayes Error

## "Too complex" hypothesis space

E.g. Memorizer  
 ERM is not PAC learner

## Uniform Law of Large Numbers

ULLN holds for  $\mathcal{H} \iff$   
 ERM is PAC learner.

## PAC learning

Successful PAC learner: finds approximately correct predictor with high probability.

$$m \geq m_{\text{pac}}^{\mathcal{H}}(\varepsilon, \delta) :$$

$$\mathbb{P}[\text{estimation error} \leq \varepsilon] \geq 1 - \delta$$

## Finite hypothesis space

$$\mathcal{H} = \{h_1, h_2, \dots, h_{\mathcal{H}}\}$$

ERM is PAC learner

$$m_{\text{pac}}^{\mathcal{H}}(\varepsilon, \delta) = \frac{2}{\varepsilon^2} \log \left( \frac{2|\mathcal{H}|}{\delta} \right)$$

## VC dimension

$$\text{VCdim}: \{-1, +1\}^{\mathcal{X}} \rightarrow \mathbb{N}$$

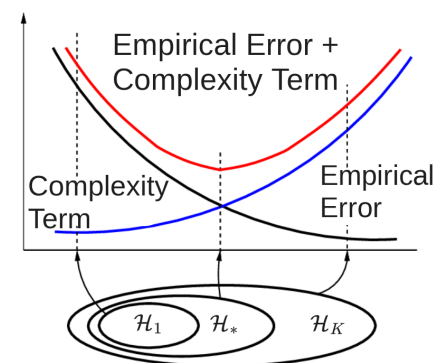
Fundamental Theorem:

$$\text{VCdim}(\mathcal{H}) < \infty$$

$$\iff \text{ULLN holds for } \mathcal{H}$$

$$\iff \text{ERM is PAC learner}$$

## Structured Risk Minimization



## Summary of Key Concepts

- ◆ **Vapnik–Chervonenkis (VC) dimension:** measures the complexity of a hypothesis class containing two-class classifiers  $\mathcal{H} \subset \{-1, +1\}^{\mathcal{X}}$ .
  - Defined as the largest number of points that can be *shattered* by  $\mathcal{H}$ .
- ◆ **Fundamental Theorem of PAC Learning:** A finite VC dimension implies that the Uniform Law of Large Numbers (ULLN) holds, and that Empirical Risk Minimization (ERM) is a PAC learner.
- ◆ **Structural Risk Minimization (SRM):** Minimize the empirical risk while simultaneously controlling the complexity of the hypothesis class.