

Statistical Machine Learning (BE4M33SSU)

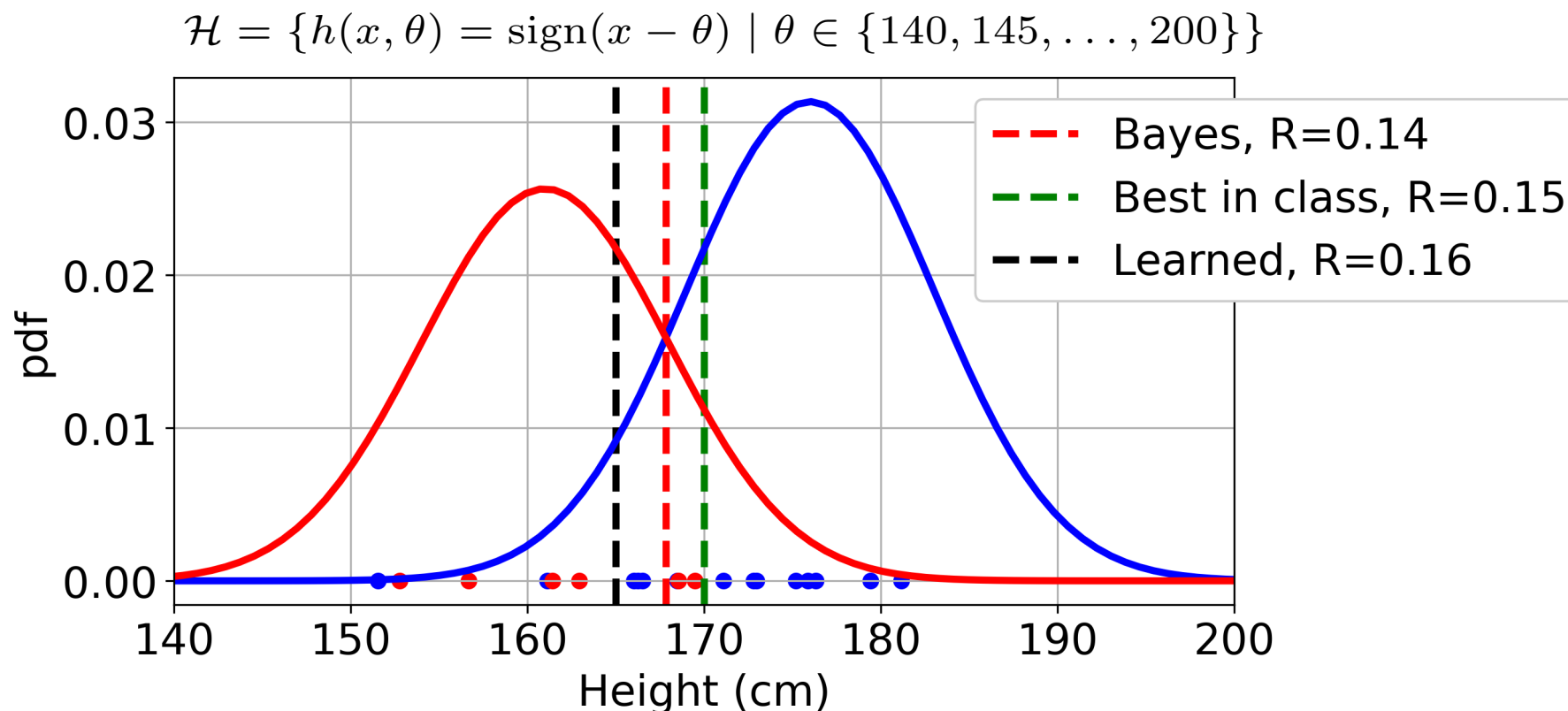
Lecture 3: Probably Approximately Correct Learning

Czech Technical University in Prague
V. Franc

Error decomposition

Errors:

1. Best (Bayes) attainable risk $R(p, h_*)$, where $h_*(x) = \arg \min_{h \in \mathcal{Y}^{\mathcal{X}}} R(p, h)$
2. Best risk in the class $R(p, h_{\mathcal{H}})$, where $h_{\mathcal{H}} = \arg \min_{h \in \mathcal{H}} R(p, h)$
3. Risk of the learned predictor $R(p, h_m)$, where $h_m = A(T_m)$



Error decomposition

Errors:

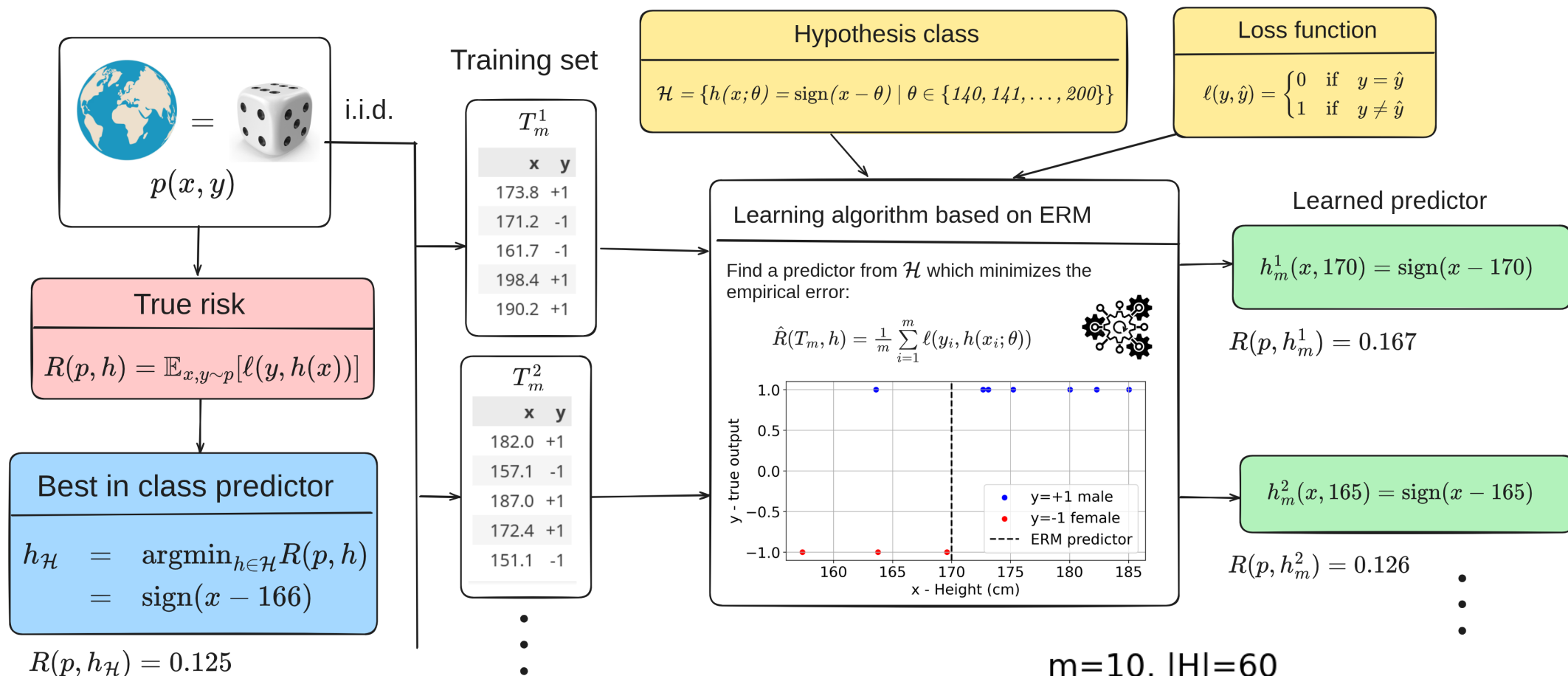
1. Best (Bayes) attainable risk $R(p, h_*)$, where $h_*(x) = \arg \min_{h \in \mathcal{Y}^{\mathcal{X}}} R(p, h)$
2. Best risk in the class $R(p, h_{\mathcal{H}})$, where $h_{\mathcal{H}} = \arg \min_{h \in \mathcal{H}} R(p, h)$
3. Risk of the learned predictor $R(p, h_m)$, where $h_m = A(T_m)$

Error decomposition:

$$\underbrace{R(p, h_m)}_{\text{learned predictor risk}} = \underbrace{\left(R(p, h_m) - R(p, h_{\mathcal{H}}) \right)}_{\text{estimation error}} + \underbrace{\left(R(p, h_{\mathcal{H}}) - R(p, h_*) \right)}_{\text{approximation error}} + \underbrace{R(p, h_*)}_{\text{Bayes risk}}$$

- ◆ The approximation error: depends on $\mathcal{H}$ chosen prior to learning.
- ◆ The estimation error: depends on $\mathcal{H}$, training data T_m and the algorithm A .
- ◆ Best (Bayes) attainable risk: irreducible error.

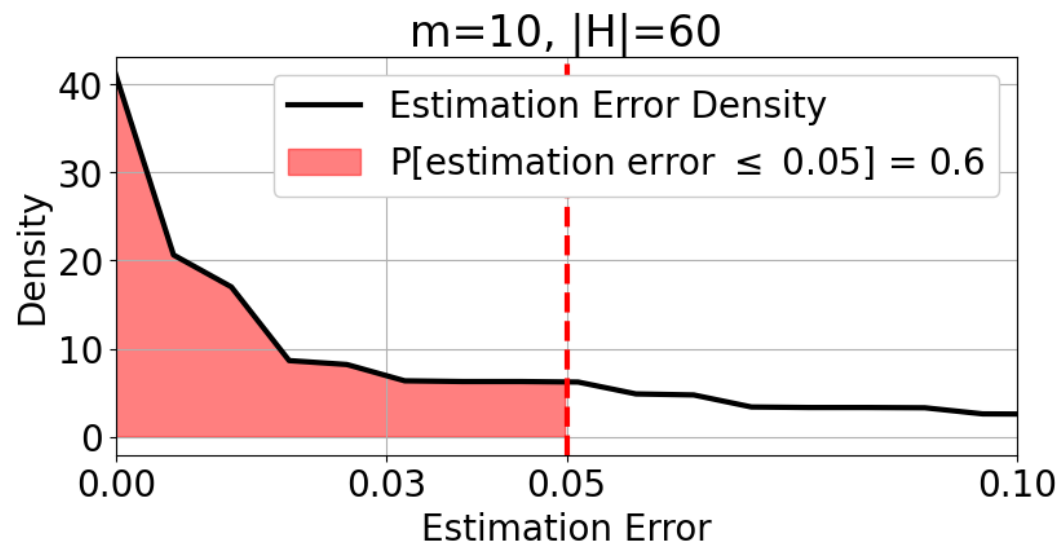
Probably Approximately Correct



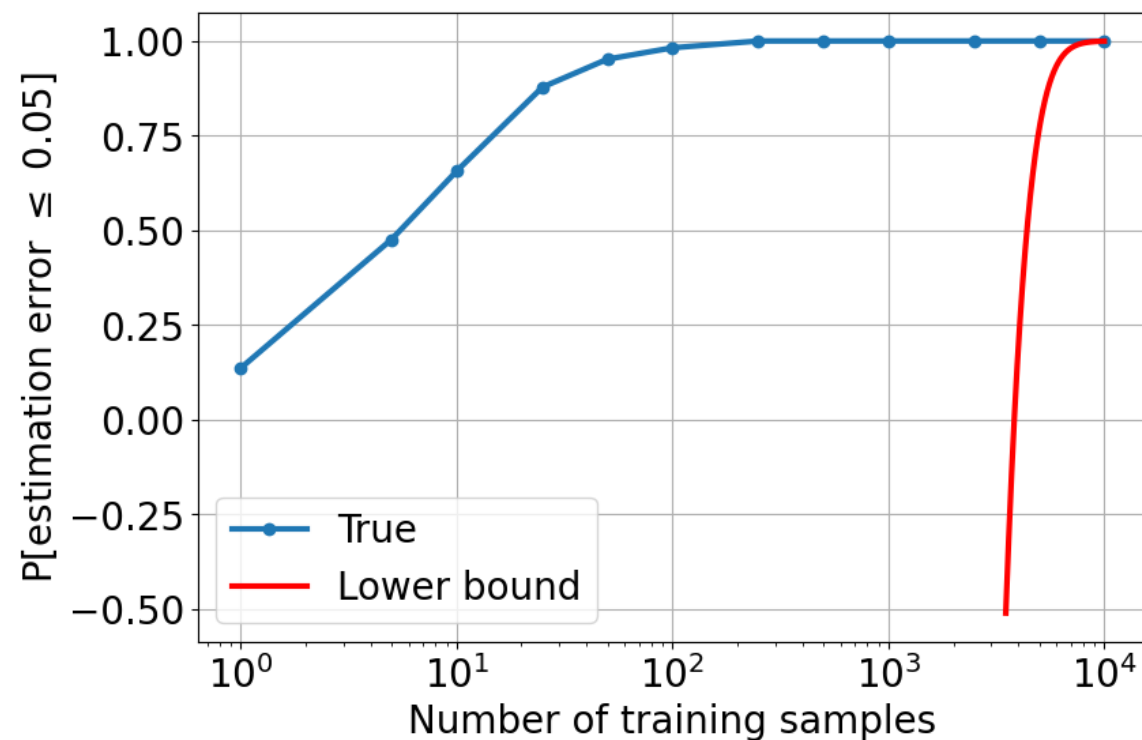
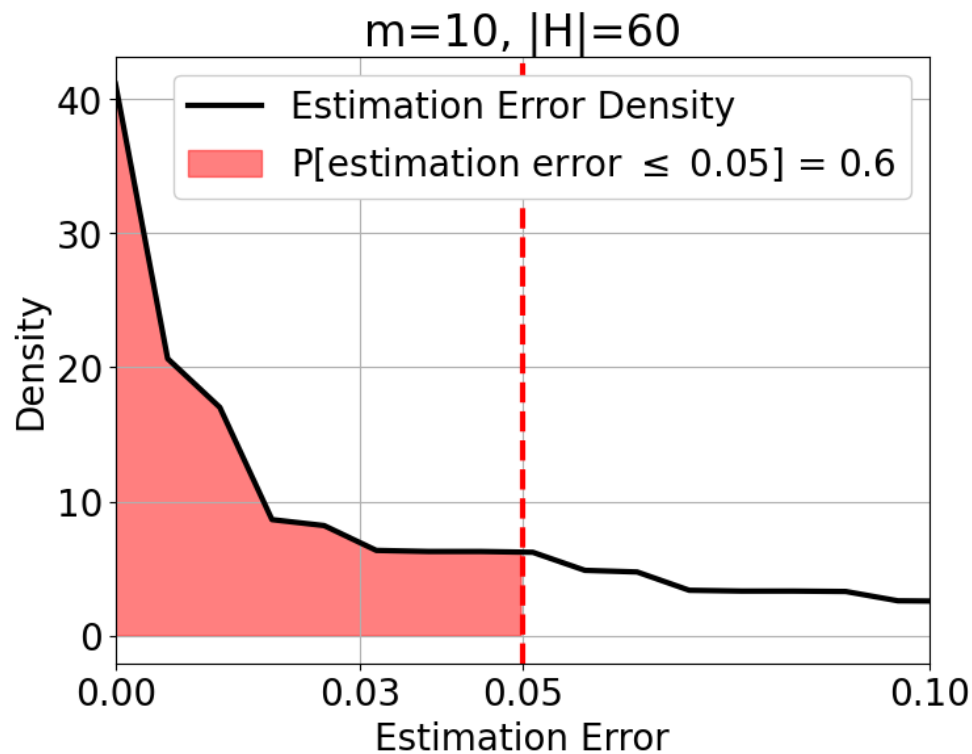
Given $\varepsilon > 0$, learned predictor h_m is approximately correct provided:

$$\underbrace{R(p, h_m) - R(p, h_{\mathcal{H}})}_{\text{estimation error}} \leq \varepsilon$$

approximately correct



Probably Approximately Correct



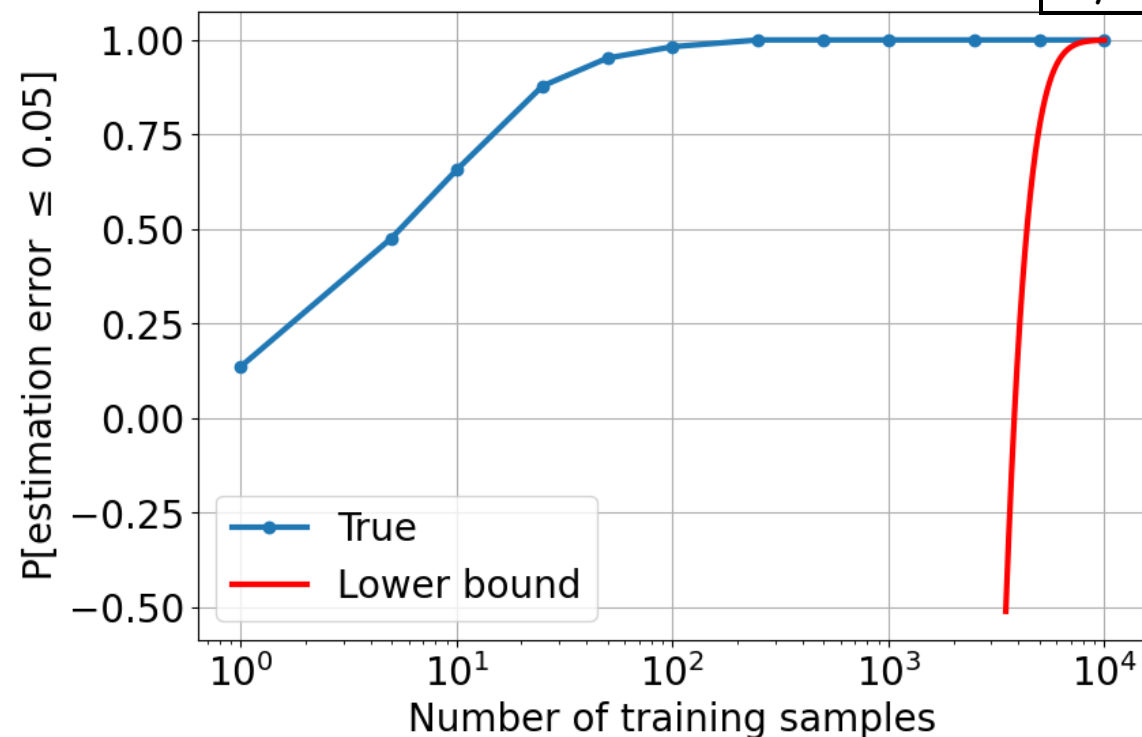
- ERM algorithm: $h_m = A(T_m) = \arg \min_{h \in \mathcal{H}} \left[\frac{1}{m} \sum_{i=1}^m \mathbb{I}[y_i \neq h(x_i)] \right]$
- We derive a lower bound valid for any distribution $p(x, y)$ and finite hypothesis class $\mathcal{H} = \{h_1, \dots, h_H\}$:

$$\mathbb{P}_{T_m \sim p^m} \left[\underbrace{R(p, h_m) - R(p, h_{\mathcal{H}}) \leq \varepsilon}_{\text{approximately correct}} \right] \geq \underbrace{1 - 2 |\mathcal{H}| e^{-\frac{1}{2} m \varepsilon^2}}_{\text{lower bound increasing with } m}$$

Probably Approximately Correct

Setup:

- ◆ $h_m = \arg \min_{h \in \mathcal{H}} \left[\frac{1}{m} \sum_{i=1}^m \mathbb{I}[y_i \neq h(x_i)] \right]$
- ◆ Finite hypothesis class:
 $\mathcal{H} = \{h^i: \mathcal{X} \rightarrow \mathcal{Y} \mid i \in \{1, \dots, H\}\}$



- ◆ Distribution-free lower bound:

$$\mathbb{P}_{T_m \sim p^m} \left[\underbrace{R(p, h_m) - R(p, h_{\mathcal{H}}) \leq \varepsilon}_{\text{approximately correct}} \right] \geq 1 - 2 |\mathcal{H}| e^{-\frac{1}{2} m \varepsilon^2} = \underbrace{1 - \delta}_{\text{probably}}$$

- ◆ Given $\varepsilon > 0$, probability of failure $\delta > 0$, we can compute the sample complexity:

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

Probably Approximately Correct learning

- ◆ **Successful PAC learning algorithm:** An algorithm can learn a hypothesis that is likely ("probably") to be approximately correct, given a sufficient number of training examples.
- ◆ **Definition:** Algorithm is a successful PAC learner for hypothesis class $\mathcal{H}$ if there exists a function $m_{\text{pac}}^{\mathcal{H}}: (0, 1) \times (0, 1) \rightarrow \mathbb{N}$ such that: For every $\varepsilon \in (0, 1)$, $\delta \in (0, 1)$, and every distribution $p(x, y)$, when running the algorithm on $m \geq m_{\text{pac}}^{\mathcal{H}}(\varepsilon, \delta)$ examples T_m i.i.d. drawn from $p(x, y)$, then the algorithm returns $h_m = A(\mathcal{T}^m)$ such that

$$\mathbb{P}_{T_m \sim p^m} \left(\underbrace{R(p, h_m) - R(p, h_{\mathcal{H}})}_{\text{approximately correct}} \leq \varepsilon \right) \geq \underbrace{1 - \delta}_{\text{probably}}$$

◆ Key Concepts:

- **Approximately correct:** The learned predictor's risk is at most ε greater than the risk of the best possible predictor in the class $\mathcal{H}$.
- **Probably:** The probability of the algorithm failing to produce the approximately correct predictor is at most δ .
- **Sample complexity** $m_{\text{pac}}^{\mathcal{H}}(\varepsilon, \delta)$: The minimum number of examples required to guaranteed the desired accuracy ε and the confidence $1 - \delta$.
- **Distribution independence:** The guarantees hold for any data distribution $p(x, y)$.

ERM on Finite Hypothesis Space is a Successful PAC learner

- ◆ **Theorem.** Let $\mathcal{H} = \{h^i : \mathcal{X} \rightarrow \mathcal{Y} \mid i \in \{1, \dots, H\}\}$ be a finite hypothesis class. Then the Empirical Risk Minimization (ERM) algorithm

$$h_m = \arg \min_{h \in \mathcal{H}} \left[\frac{1}{m} \sum_{i=1}^m \mathbb{I}[y_i \neq h(x_i)] \right]$$

is a successful PAC learner with sample complexity

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

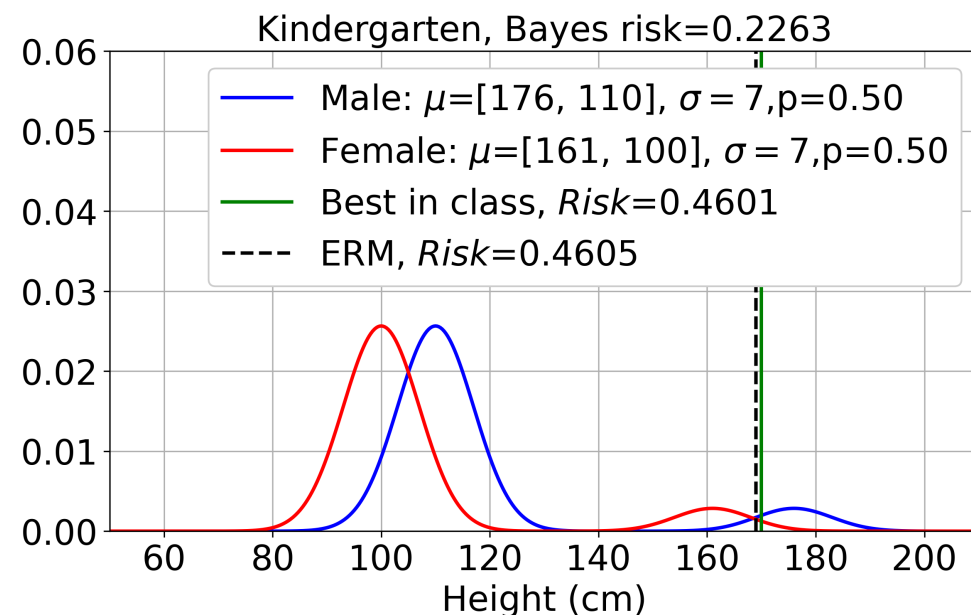
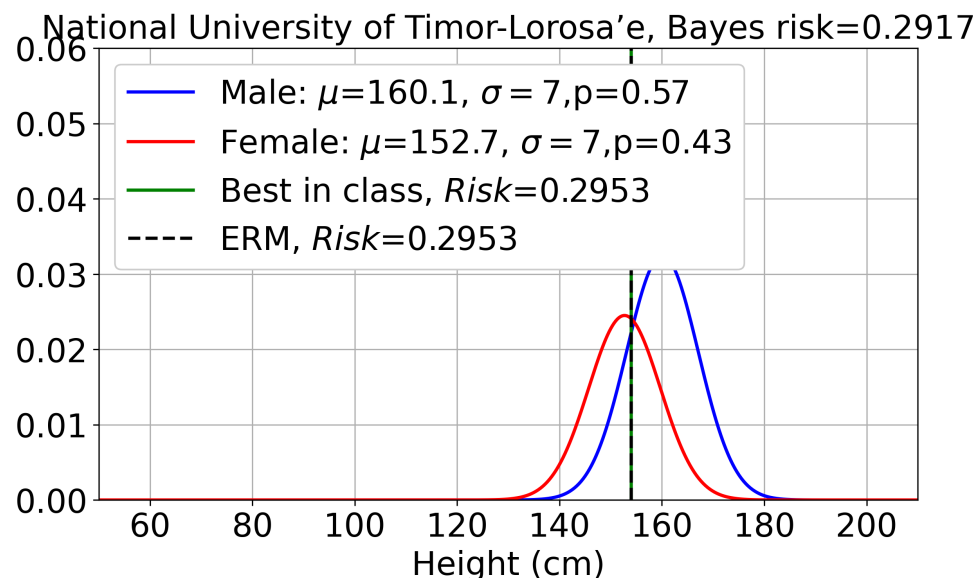
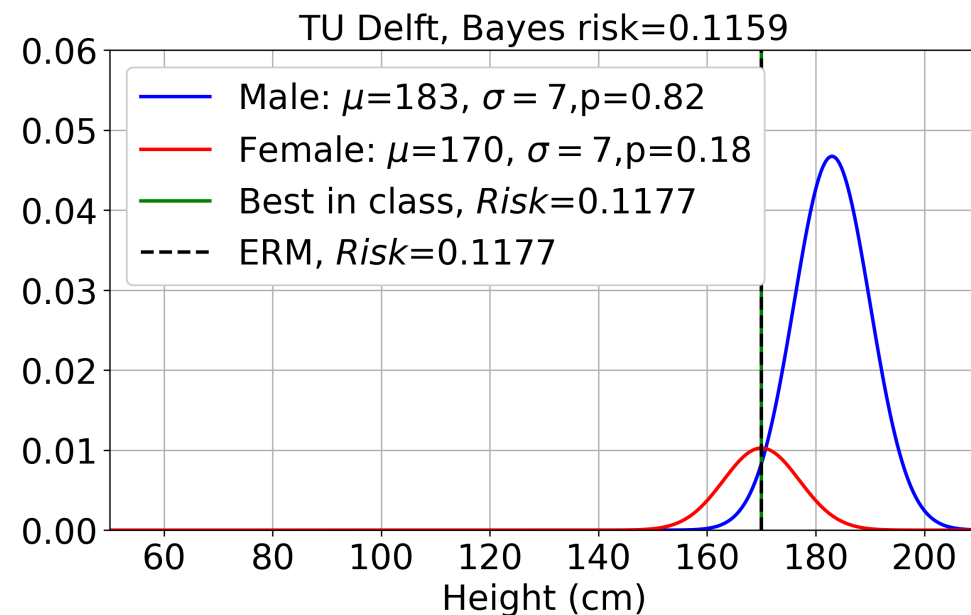
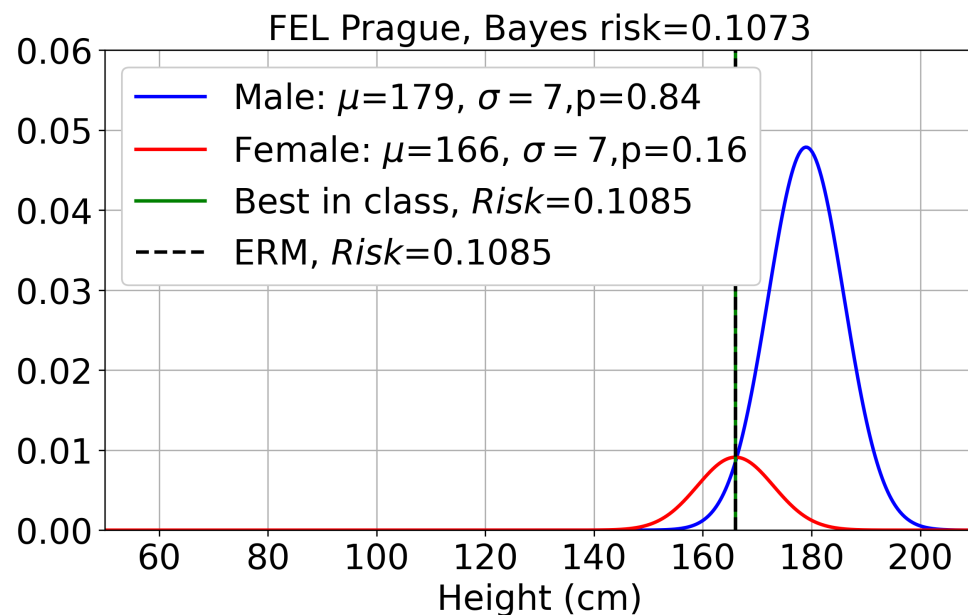
- ◆ The theorem is a consequence of the bound we introduced earlier (however, have not yet proved):

$$\mathbb{P}_{T_m \sim p^m} \left[\underbrace{R(p, h_m) - R(p, h_{\mathcal{H}})}_{\text{approximately correct}} \leq \varepsilon \right] \geq 1 - 2 |\mathcal{H}| e^{-\frac{1}{2} m \varepsilon^2} = \underbrace{1 - \delta}_{\text{probably}}$$

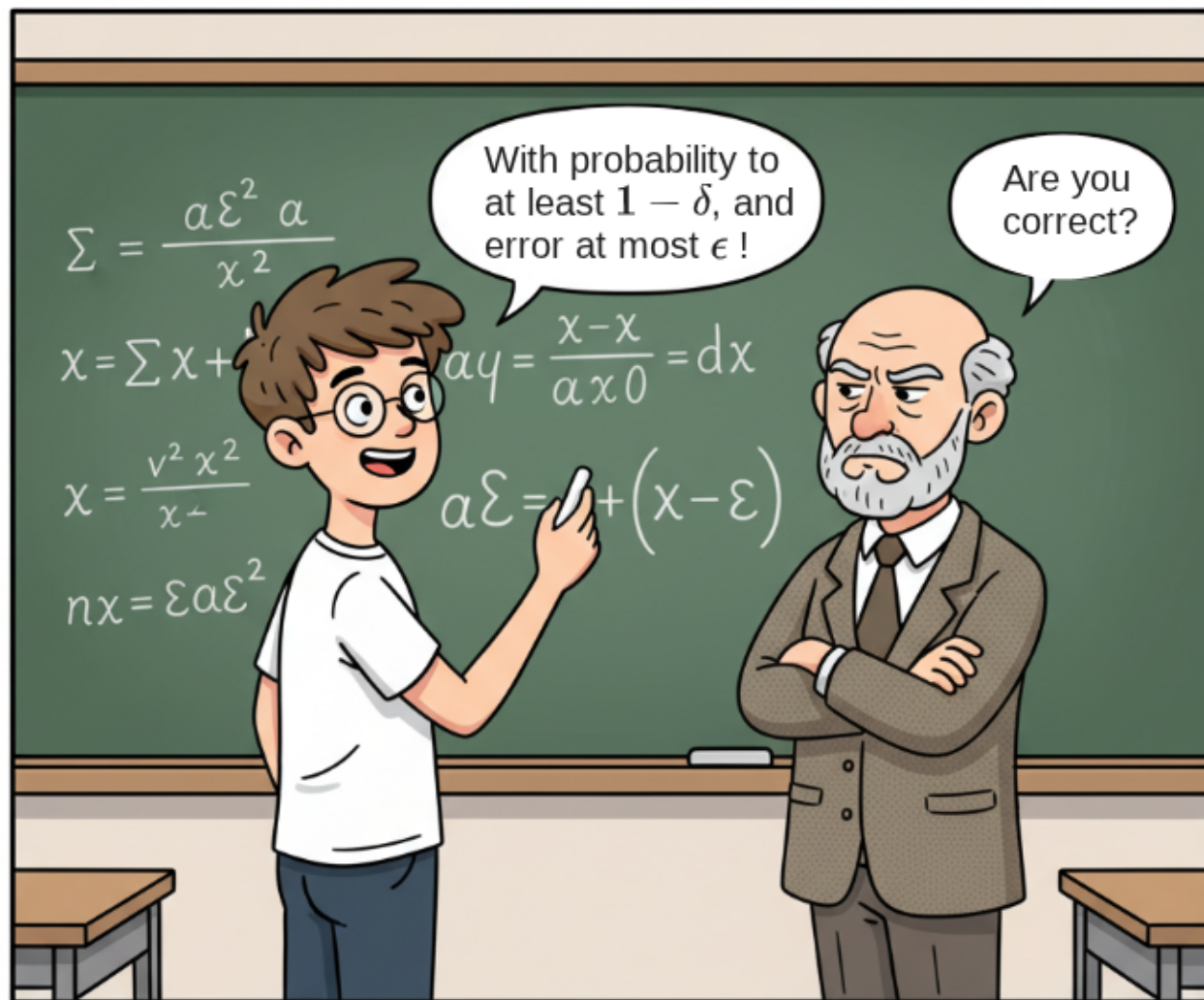
Probably Approximately Correct learning

Example: $\varepsilon = 0.05$, $\delta = 0.1$, $\mathcal{H} = \{h(x; \theta) = \text{sign}(x - \theta) \mid \theta \in \{140, 141, \dots, 200\}\}$

The sample complexity: $m_{\text{pac}}^{\mathcal{H}}(\varepsilon, \delta) = \frac{2}{\varepsilon^2} \ln \left(\frac{2|\mathcal{H}|}{\delta} \right) = \frac{2}{0.05^2} \ln \left(\frac{2 \cdot 60}{0.1} \right) \approx 5,673$



Probably Approximately Correct Joke



Joke: Chat GPT-5
Image: Nano Banana

BE4M33SSU exam edition.

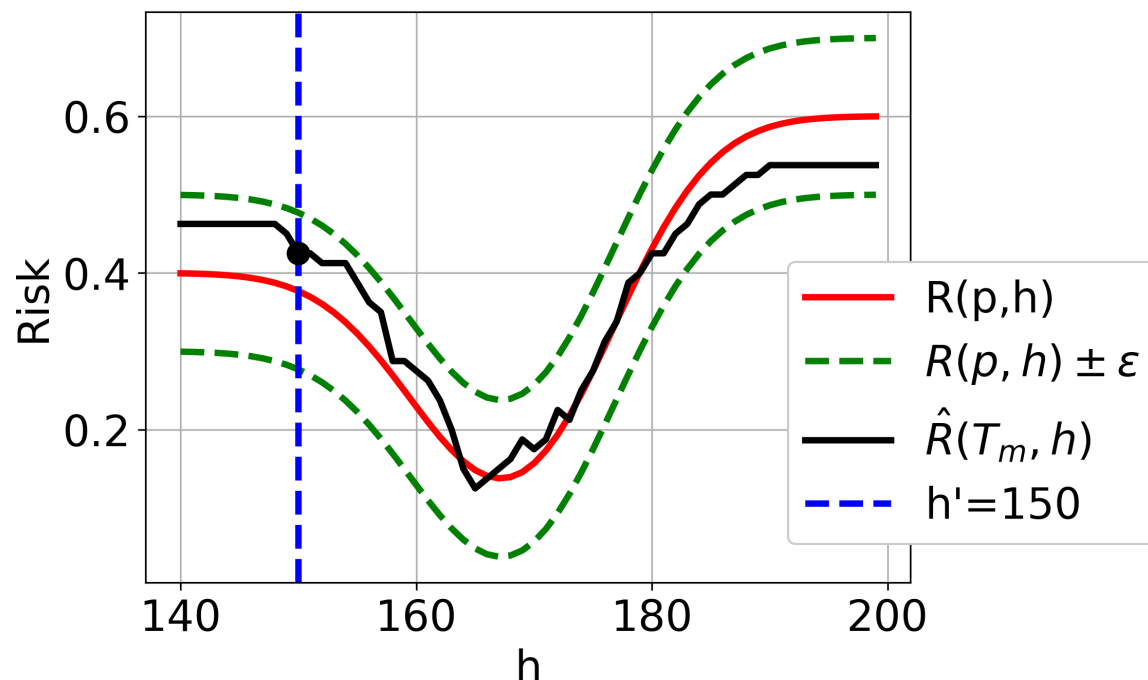
Uniform Law of Large Numbers

♦ **LLN:** $m \geq \frac{1}{2\varepsilon^2} \log\left(\frac{2}{\delta}\right) \Rightarrow \mathbb{P}\left(\underbrace{|R(p, h) - \hat{R}(T_m, h)|}_{\text{generalization error of } h \text{ is high}} > \varepsilon\right) \leq \delta$

LLN applies for any $h: \mathcal{X} \rightarrow \mathcal{Y}$ fixed prior to observing the data T_m

♦ **ULLN:** $m \geq m_{\text{ul}}^{\mathcal{H}}(\varepsilon, \delta) \Rightarrow \mathbb{P}\left(\underbrace{\max_{h \in \mathcal{H}} |R(p, h) - \hat{R}(T_m, h)|}_{\text{generalization error of some } h \in \mathcal{H} \text{ is high}} > \varepsilon\right) \leq \delta$

ULLN applies only for some $\mathcal{H}$, e.g., when $\mathcal{H}$ is finite $m_{\text{ul}}^{\mathcal{H}}(\varepsilon, \delta) = \frac{1}{2\varepsilon^2} \log\left(\frac{2|\mathcal{H}|}{\delta}\right)$



Example:

$$\mathcal{H} = \{h(x; \theta) = \text{sign}(x - \theta) \mid \theta \in \{140, 141, \dots, 200\}\}$$

$$\varepsilon = 0.1, \delta = 0.05, |\mathcal{H}| = 60$$

$$m_{\text{ul}}^{\mathcal{H}}(0.1, 0.05) = 389.2$$

Uniform Law of Large Numbers

- ◆ Assume a finite hypothesis class $\mathcal{H} = \{h^i: \mathcal{X} \rightarrow \mathcal{Y} \mid i \in \{1, 2, \dots, H\}\}$.

$$\begin{aligned}
 \mathbb{P}\left(\underbrace{\max_{h \in \mathcal{H}} |R(p, h) - \hat{R}(T_m, h)| \geq \varepsilon}_{\text{generalization error of some } h \in \mathcal{H} \text{ is high}}\right) &\stackrel{(1)}{=} \mathbb{P}\left(\begin{array}{c} |R(p, h^1) - \hat{R}(T_m, h^1)| \geq \varepsilon \quad \text{or} \\ |R(p, h^2) - \hat{R}(T_m, h^2)| \geq \varepsilon \quad \text{or} \\ \vdots \\ |R(p, h^H) - \hat{R}(T_m, h^H)| \geq \varepsilon \end{array}\right) \\
 &\stackrel{(2)}{\leq} \sum_{h \in \mathcal{H}} \mathbb{P}\left(|R(p, h) - \hat{R}(T_m, h)| \geq \varepsilon\right) \\
 &\stackrel{(3)}{\leq} 2 |\mathcal{H}| e^{-2m\varepsilon^2}
 \end{aligned}$$

1. $a \geq \varepsilon \text{ or } b \geq \varepsilon \iff \max\{a, b\} \geq \varepsilon$

2. Union bound: $\mathbb{P}(A_1 \text{ or } A_2 \text{ or } \dots \text{ or } A_n) \leq \sum_{i=1}^n \mathbb{P}(A_i)$

3. Hoeffding inequality: $\mathbb{P}(|R(h) - R_{\mathcal{T}^m}(h)| \geq \varepsilon) \leq 2 e^{-2m\varepsilon^2}$

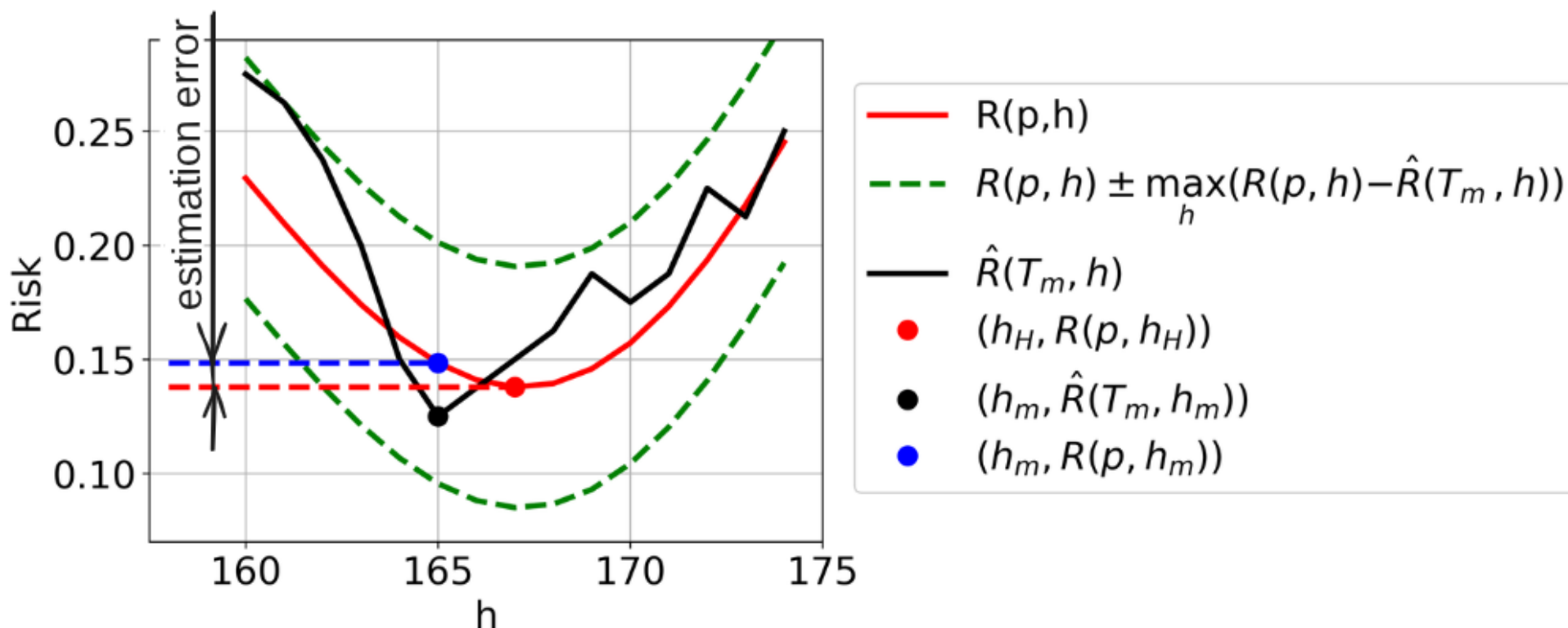
- ◆ Setting $2 |\mathcal{H}| e^{-2m\varepsilon^2} = \delta$ and solving for m , we get

$$m_{\text{ul}}^{\mathcal{H}}(\varepsilon, \delta) = \frac{1}{2\varepsilon^2} \log\left(\frac{2|\mathcal{H}|}{\delta}\right) \Rightarrow \mathbb{P}\left(\max_{h \in \mathcal{H}} |R(h) - R_{\mathcal{T}^m}(h)| \geq \varepsilon\right) \leq \delta$$

Bound on Estimation Error

- ◆ **Bound on estimation error of $h_m = \arg \min_{h \in \mathcal{H}} \hat{R}(T_m, h)$:**

$$\begin{aligned}
 \underbrace{R(p, h_m) - R(p, h_{\mathcal{H}})}_{\text{estimation error}} &= \left(R(p, h_m) - \hat{R}(T_m, h_m) \right) + \left(\hat{R}(T_m, h_m) - R(p, h_{\mathcal{H}}) \right) \\
 &\leq \left(R(p, h_m) - \hat{R}(T_m, h_m) \right) + \left(\hat{R}(T_m, h_{\mathcal{H}}) - R(h_{\mathcal{H}}) \right) \\
 &\leq \underbrace{2 \max_{h \in \mathcal{H}} |R(p, h) - \hat{R}(T_m, h)|}_{\text{Maximal generalization error}}
 \end{aligned}$$



ULLN implies ERM is PAC learner

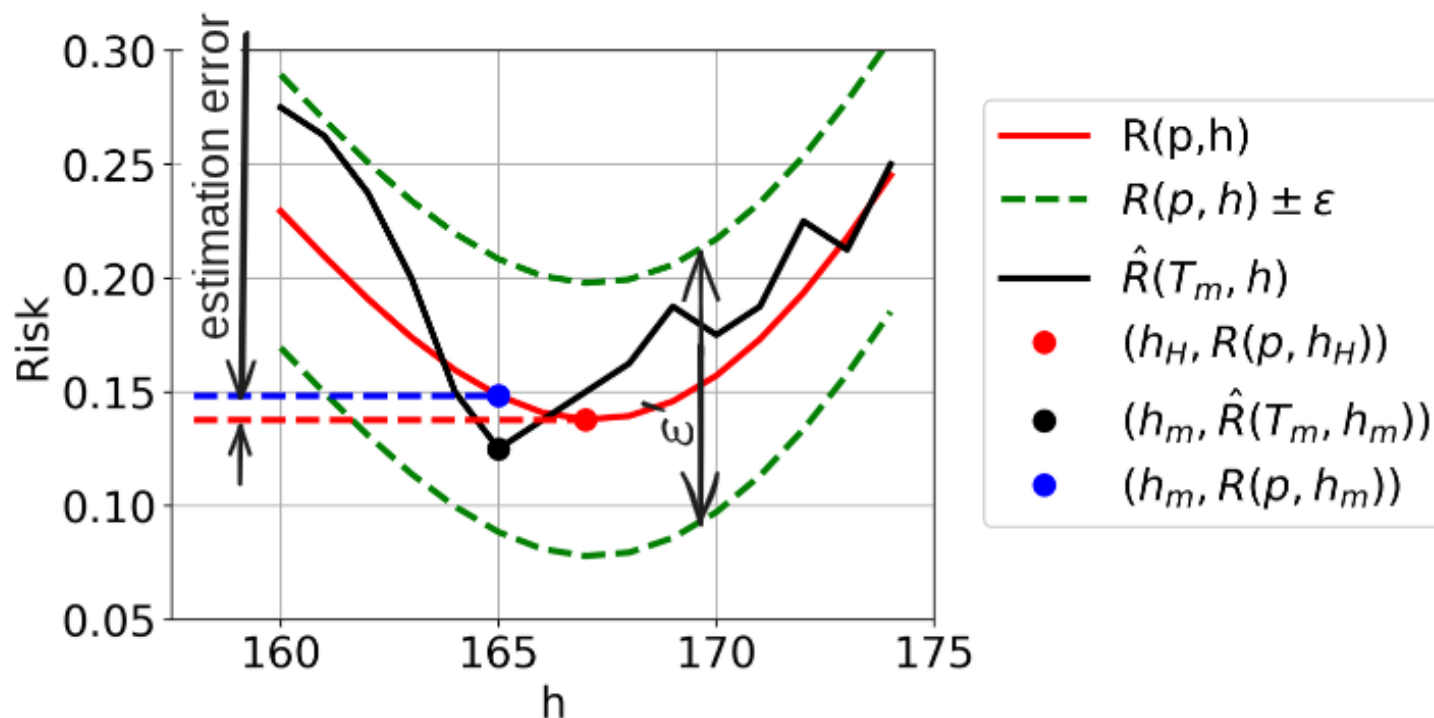
◆ **ULLN:** $m \geq m_{\text{ul}}^{\mathcal{H}}(\varepsilon, \delta) \Rightarrow \mathbb{P}\left(\max_{h \in \mathcal{H}} |R(p, h) - \hat{R}(T_m, h)| > \varepsilon\right) \leq \delta$

◆ **Bound on estimation error of** $h_m = \arg \min_{h \in \mathcal{H}} \hat{R}(T_m, h)$:

$$R(p, h_m) - R(p, h_{\mathcal{H}}) \leq 2 \max_{h \in \mathcal{H}} |R(p, h) - \hat{R}(T_m, h)|$$

◆ **ULLN + Bound on estimation error = ERM is successful PAC learner:**

$$m \geq m_{\text{pac}}^{\mathcal{H}}(\varepsilon', \delta) = m_{\text{ul}}^{\mathcal{H}}(\varepsilon'/2, \delta) \Rightarrow \mathbb{P}\left(R(p, h_m) - R(p, h_{\mathcal{H}}) \leq \varepsilon'\right) \geq 1 - \delta$$



Example:

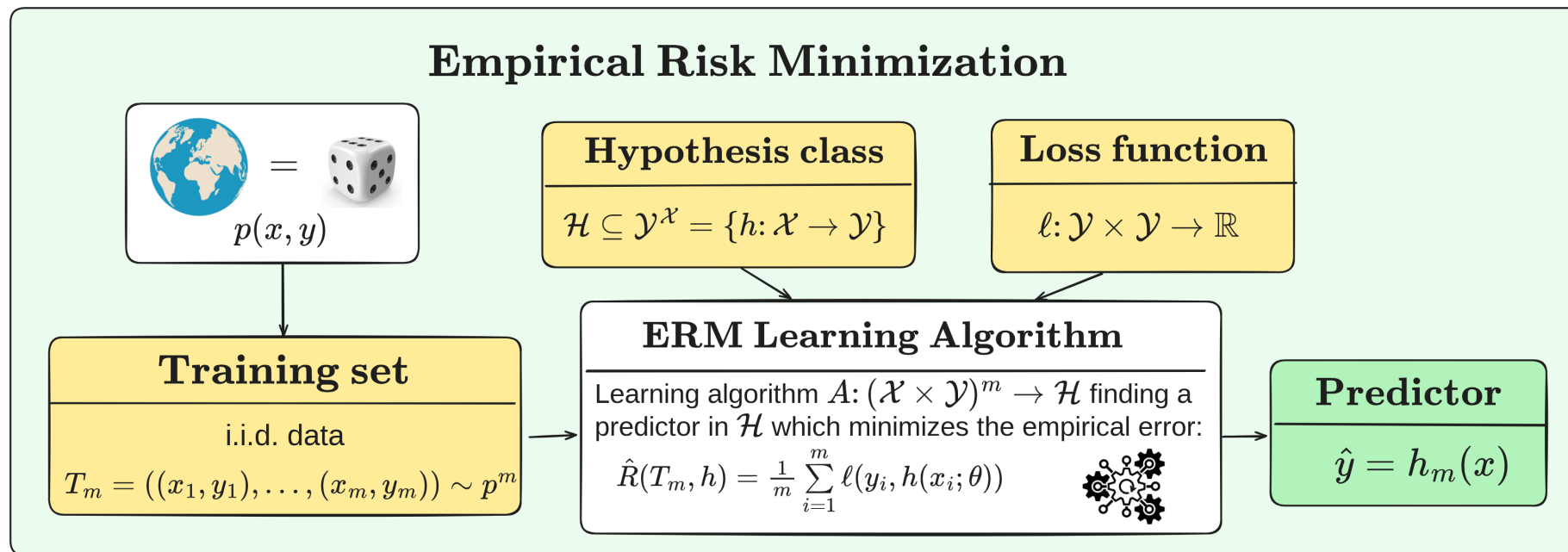
For finite $\mathcal{H}$, we have:

$$m_{\text{ul}}^{\mathcal{H}}(\varepsilon, \delta) = \frac{1}{2\varepsilon^2} \log \left(\frac{2|\mathcal{H}|}{\delta} \right)$$

Hence, the sample complexity is:

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

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

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

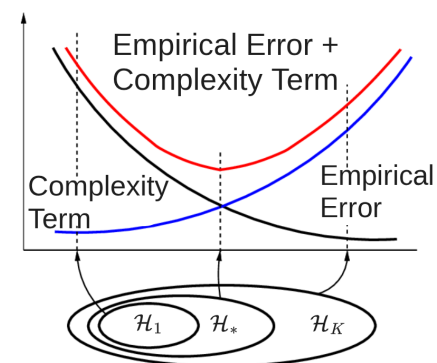
Fundamental Theorem:

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

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

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

Structured Risk Minimization



Summary of Key Concepts

◆ Error Decomposition

- Error of the learned predictor = Estimation Error + Approximation Error + Bayes Risk

◆ Probably Approximately Correct (PAC) Learning

- A successful PAC learner, with high probability, finds a close approximation of the best predictor in the class, given enough examples.
- Sample complexity: number of examples needed for PAC guarantees.

◆ Empirical Risk Minimization:

- ERM over a finite hypothesis space is a successful PAC learner.

◆ Uniform Law of Large Numbers

- Guarantees uniform convergence of empirical risk to the expected risk over the hypothesis space.