

По-перше, замість пари з розподілу $\mathcal{D}$ на X та $f: X \rightarrow Y$ розглянемо
т. зв. об'єднаний розподіл $\mathcal{D}$ - ймовірнісну міру на $X \times Y$
(joint distribution)
(де X та Y мають σ -алгебри, що на $X \times Y$ $\exists$ стр. ймовірнісна міра).

Подібно $\forall A \subset X \times Y$ (більш точно, A з борн. σ -алгебри) $\mathcal{D}(A)$ -
це ймовірність того, що випадково обрана пара з точки даних
 $x \in X$ та мітки $y \in Y$ належить до A : $\mathcal{D}(A) = P((x, y) \in A)$.
 $(x, y) \sim \mathcal{D}$
Зокрема, $\mathcal{D}$ визначає:

- маргінальний (marginal) розподіл $\mathcal{D}_x$ на X : для $A \subset X$
це ім. того, що випадкова $x \in X$ належить до A незалежно від мітки:

$$\mathcal{D}_x(A) = P_{x \sim \mathcal{D}_x}(x \in A) = P_{(x, y) \sim \mathcal{D}}((x, y) \in A \times Y) = \mathcal{D}(A \times Y)$$

$\exists \underline{E}_x$ з фактами це, як раніше, ім. потрапити у певну
область за координат і твердістю.

— умовна ймовірність (conditional probability) того, що $y \in B$, для деякої $B \subset Y$ за умови $x \in A$:

$$P_{(x,y) \sim \mathcal{D}}(y \in B | x \in A) = \frac{\mathcal{D}(A \times B)}{\mathcal{D}_x(A)} = \frac{\mathcal{D}(A \times B)}{\mathcal{D}(A \times Y)}.$$

Зокрема, для одноточкової $B = \{y\}$ це ім. отримати мітку y для всіх точок даних $x \in A$: $P_{(x,y) \sim \mathcal{D}}(y | x \in A) = \frac{\mathcal{D}(A \times \{y\})}{\mathcal{D}_x(A)}.$

Якщо $x \in A = \{x\}$ одноточкова, це ім. отримати y для цієї x ($y \in \mathcal{Y}$ — тобто, що функція з котрою і твердістю x буде ставити (при $y=1$) чи ні ($y=0$)). Формально це $P_{(x,y) \sim \mathcal{D}}(y | x) = \frac{\mathcal{D}(\{(x,y)\})}{\mathcal{D}_x(\{x\})}$. Для неперервного $(X, \mathcal{D}_x)$ (скажімо, $X \subset \mathbb{R}^d$ з абсолютним чи поря. $\mathcal{D}_x$) це може бути невизначення ($\frac{0}{0}$), тоді будемо (не дуже формально) розуміти це як границю $\lim_{\substack{A \rightarrow \{x\}, \\ x \in A}} \frac{\mathcal{D}(A \times \{y\})}{\mathcal{D}_x(A)}$. Дати вважасямо $P_{(x,y) \sim \mathcal{D}}(y | x)$ визначенням $\left(\begin{smallmatrix} \text{всюди, де це буде потрібно} \\ \text{(можливо, куди)} \end{smallmatrix} \right)$.

Знову ж, на практиці всі процеси дискретні, тому можна вважати, що точки мають певні міри.

Ех. Для полунепрерывной модели $z \in (\mathbb{Q}, +]$ $\mathbb{Q}_x$ - це $i \in (\text{старий})\mathbb{Q}$,
а $P(y \in B | x) = \begin{cases} 1, & f(x) \in B \\ 0, & f(x) \notin B \end{cases}$ зотн. $P(y | x) = \begin{cases} 1, & f(x) = y \\ 0, & f(x) \neq y \end{cases}$

$$\forall A \subset X \times Y$$

$$\underset{\substack{\uparrow \\ \text{у новому} \\ \text{сенсі}}}{\mathbb{Q}}(A) = \underset{\substack{\uparrow \\ \text{у старому} \\ \text{сенсі}}}{\mathbb{Q}}(\{x \in X \mid (x, f(x)) \in A\}).$$

Вар. Чи визначений $\mathbb{Q}$ мартингал. розр. $\mathbb{Q}_x$ та/або умовними йм.
у задовільному випадку?

П.ч., тепер лінійки y , взагалі кансуки, не визначені одно-
значно за x (лише з деякою йм.), що неадекватно.

По-друге, уточнено деф. помилки. Для задач класифікації,
тобто коли Y скінченна, це можна зробити так:

def. (Справністью) помилкою гіпотези $h: X \rightarrow Y$ для задачі класифікації з розподілом $\mathcal{D}$ на $X \times Y$ будемо називати

$$L_{\mathcal{D}}(h) := \mathcal{D}(\{(x, y) \in X \times Y \mid h(x) \neq y\}) \in [0, 1].$$

Рем. Для моделі з $(\mathcal{D}, f)$ це не те саме, що $L_{\mathcal{D}, f}(h)$ (чому?), а припустимість

а реалізованості означає δ , що $L_{\mathcal{D}}(h^*) = 0$ для деякої h^* з обраного класу гіпотез $\mathcal{H}$. Втім, тут ми не будемо відножувати до $\mathcal{H}$ загального випадку $L_{\mathcal{D}}(h)$ для $h \in \mathcal{H}$ повинні бути об'єктивні знання:

Ex. Байєсівським оптимальним предиктором (гіпотезою) для задачі класифікації з розп. $\mathcal{D}$ на $X \times Y$, де $Y = \{0, 1\}$, є $f_{\mathcal{D}}: X \rightarrow Y$: $f_{\mathcal{D}}(x) := \begin{cases} 1 & \text{if } P(y=1|x) > \frac{1}{2} \\ 0 & \text{иначе} \end{cases}$

Рм. Б.о.п. має найменшу справність помилку серед усіх гіпотез : $L_{\mathcal{D}}(f_{\mathcal{D}}) \leq L_{\mathcal{D}}(h) \quad \forall h: X \rightarrow Y$

$\Rightarrow \forall h \quad L_{\mathcal{D}}(h)$ - міра мисл. помилки (x, y) , де $h(x) \neq y$. У нас:

$$L_{\mathcal{D}}(h) = \mathcal{D}(\{(x, 0) \mid h(x) = 1\} \cup \{(x, 1) \mid h(x) = 0\}) = [\text{адам.}] =$$

$$= \mathbb{Q}(\{x, 0\} | h(x) = 1\}) + \mathbb{Q}(\{x, 1\} | h(x) = 0\}) \in \left(\begin{array}{l} \text{Плюс і гарні} \\ \int_A f(x) d\mu(x) - \text{інт.} \\ \text{Левора} \end{array} \right)$$

Використаємо визначення нонел. Випр. для $\mathcal{Y} = \{0, 1\}$:

$$\mathbb{Q} \int_{\{x | h(x)=1\}} P(0|x) d\mathbb{Q}_x(x) + \int_{\{x | h(x)=0\}} P(1|x) d\mathbb{Q}_x(x) \stackrel{(\circledast)}{=} [P(0|x) + P(1|x) = 1 \forall x]$$

$$\geq \int_{\{x | h(x)=1\}} \min\{P(1|x), 1-P(1|x)\} d\mathbb{Q}_x(x) + \int_{\{x | h(x)=0\}} \min\{P(1|x), 1-P(1|x)\} d\mathbb{Q}_x(x) =$$

$$= [\text{адаптивність інтеграла}] = \int \min\{P(1|x), 1-P(1|x)\} d\mathbb{Q}_x(x).$$

Якщо не $h = f \circledast$, то $h(x) = 1 \Leftrightarrow P(1|x) > \frac{1}{2}$ і $h(x) = 0 \Leftrightarrow$

$$\Leftrightarrow P(1|x) < \frac{1}{2}. \text{ Оскільки } \int_A f d\mu = \int_X 1_A \cdot f d\mu, \text{ де, нагадаємо, } 1_A \text{ —}$$

ктеристична функція $1_A(x) = \begin{cases} 1, & x \in A \\ 0, & x \notin A \end{cases}$; звідси пер-ми

(7) вище можна записати:

$$L_{\mathbb{Q}}(f \circledast) = \dots \stackrel{(\circledast)}{=} \int_X 1_{\{x | P(1|x) > \frac{1}{2}\}}^{(x)} (1 - P(1|x)) d\mathbb{Q}_x(x) + \int_X 1_{\{x | P(1|x) < \frac{1}{2}\}}^{(x)} P(1|x) d\mathbb{Q}_x(x).$$

$$\cdot P(1|x) d\mathbb{Q}_x(x) = [\text{лінійність інтеграла}] = \int_X \left(1_{\{x | P(1|x) > \frac{1}{2}\}}^{(x)} (1 - P(1|x)) + 1_{\{x | P(1|x) < \frac{1}{2}\}}^{(x)} P(1|x) \right) d\mathbb{Q}_x(x).$$

$d\mathbb{Q}_x(x) \stackrel{(\circledast)}{=}$ Показано, що $\mathbb{Q}$ -ція μ інтегровна — це

$$1 - P(1|x) \leq P(1|x) \text{ при } P(1|x) > \frac{1}{2} \text{ і } P(1|x) < 1 - P(1|x) \text{ при}$$

$$P(1|x) < \frac{1}{2}, \text{ тоді } \stackrel{(\circledast)}{=} \int_X \min\{P(1|x), 1-P(1|x)\} d\mathbb{Q}_x(x) \leq L_{\mathbb{Q}}(h)$$

за все доведення, що й потрібно. $\triangle$

Реш. Згідно з умов P_H та його доведенням, якщо, наприклад, $P(1(x) = \frac{1}{4} \forall x$ (тобто чверті точка приносена мітка 1 незалежно від точки), то $\forall h$

$$L_D(h) \geq L_D(f_D) = \int_X \frac{1}{4} dD_x = \frac{1}{4} D_x(x) = \frac{1}{4},$$

тобто припущення реалізованості невірне $\forall H$.

Нарадаємо, що наша мета — мінімізувати $L_D(h)$ для $h \in H$. При цьому просто $h = f_D$ покласти не можна, бо алгоритм не "знає" D , лише навчальні дані $S = \{(x_i, y_i)\}_{i=1}^m \in (X \times Y)^m$. Як і раніше, застосуємо парадигму ERM_H : для $h \in H$ мінімізуємо $L_S(h)$, що визначена для задачі класифікації як емпіричну помилку раніше:

$$L_S(h) = \frac{1}{m} |\{i = \overline{1, m} \mid h(x_i) \neq y_i\}| \in [0, 1],$$

Але тепер ми не можемо скористатися оцніткою $L_D(h)$

лише $\min_{h' \in H} L_D(h')$ (взагалі кажучи, ів), що, як ми покажемо,

у зад. випадку > 0 . Терм π_{loss} дає нам певне означення навантаження $\mathcal{K}$, узгодженого L_D і L_S на довільній $\mathcal{Y}$.

дев. Кожна задача ML моделюється ймовірнісним простором $\mathcal{Z}$, а $\mathcal{K}$ - клас гіпотез ($\forall$ множини). Поді (узгодженою)

функцією втрат (loss function) зв. $\forall$ відображення $\ell: \mathcal{K} \times \mathcal{Z} \rightarrow \mathbb{R}_+$ таке, що $\forall h \in \mathcal{K} \quad \ell(h, \cdot)$ (тобто відобр. $z \in \mathcal{Z} \mapsto \ell(h, z) \in \mathbb{R}_+$) є вимірною ф-цією на $\mathcal{Z}$.

Рем. Нагадаємо, що $\mathbb{R}_+ = [0, +\infty)$ і що вимірна ф-ція $\rightarrow \mathbb{R}_+$ - та, для якої прообрази вимірних множин вимірні ($\Leftrightarrow \ell(h, \cdot)^{-1}([0, a])$ вимірні у $\mathcal{Z} \quad \forall a$).

Іншими словами, $\ell(h, \cdot)$ - випадкова величина (невід-сна) на $\mathcal{Z}$. Для наших задач керуваного ML

$\mathcal{Z} = \mathcal{X} \times \mathcal{Y}$ з розподілом $\mathcal{D}$, як вище, але у випадку некерованого ML може бути і $\mathcal{Z} = \mathcal{X}$ (немає міток).

Ек. 1. Для задач класифікації розглядають завжди

м.зв. 0-1 втрата (0-1 loss). Тим Y скінченна,

$$h: X \rightarrow Y, \quad \ell_{0-1}(h, (x, y)) := \begin{cases} 1, & h(x) \neq y \\ 0, & h(x) = y \end{cases} \quad \left(\begin{array}{l} \text{Тим жову } \ell_{0-1}(h, \cdot) = \\ = 1_{\{(x, y) \mid h(x) \neq y\}} \end{array} \right)$$

Воп. Перевіримо умову валирності.

2. Якщо $Y = \mathbb{R}$, то зазвичай говорять про задачу регресії (regression). Для такої задачі ма $h: X \rightarrow Y$ типовий вибором є квадратична втрата (square loss)

$$\ell_{sq}(h, (x, y)) := (h(x) - y)^2$$

Це можна узагальнити й на $Y \subset \mathbb{R}^d$: тоді це квадратичн. втрата $\sum_{i=1}^n (h(x)^i - y^i)^2$ (метод найменших квадратів).

деф. (Справжнього) пачилко (пузак) іпачеж $h \in \mathcal{H}$ ма функції втрат ℓ зветься намерцатичне оціування ℓ і розподілу $\mathbb{P}$ на $\mathcal{Z}$.

$$L_{\mathbb{P}}(h) := \mathbb{E}_{z \sim \mathbb{P}} \ell(h, z).$$

Емпіричного пачилко h на $S = \{z_i\}_{i=1}^m \in \mathcal{Z}^m$ зветься середнє значення (пузак) навірцатичнє даннє

$$L_S(h) := \frac{1}{m} \sum_{i=1}^m \ell(h, z_i).$$

Рем. Можливо, що мат. сподівання (expectation) буд. величини (у даному випадку $l(h, \cdot)$) - це інтеграл $E_{z \sim \mathcal{D}} l(h, z) = \int_{\mathcal{Z}} l(h, z) d\mathcal{D}(z)$ (що існує в силу вимірності).

Зокрема, середнє значення - це мат. оч. для скінченного простору S і рівномірного розподілу (коли кожна z_i має ім. $\frac{1}{m}$): $L_S(h) = L_{\mathcal{D}_{\text{рівн.}}(h)}$ (на S).

Ех. 1. Для 0-1 втради:

$$L_{\mathcal{D}}(h) = E_{(x,y) \sim \mathcal{D}} l_{0,1}(h, \underbrace{(x,y)}_{\mathcal{X} \times \mathcal{Y}}) = \int 1_{\{x,y \mid h(x) \neq y\}} d\mathcal{D} = \mathcal{D}(\{(x,y) \mid h(x) \neq y\}) -$$

та, що факт, і так само для $L_S(h)$ (чому?).

2. Для квадратичної втради:

$$L_{\mathcal{D}}(h) = E_{(x,y) \sim \mathcal{D}} l_{sq}(h, (x,y)) = E_{(x,y) \sim \mathcal{D}} (h(x) - y)^2,$$

$$\text{і } L_S(h) = \frac{1}{m} \sum_{i=1}^m (h(x_i) - y_i)^2 = \frac{1}{m} \sum_{i=1}^m (l_{sq}(h, (x_i, y_i)))$$

Тепер менше реплісати дів. PAC навчальності:

def. Для задані ML на просторі Z з функцією втрач ℓ клас гіпотез $\mathcal{H}$ (на якому ℓ визначена) зветься априорно PAC навчанням, якщо існує φ -гін $m_{\mathcal{H}}: (0,1)^2 \rightarrow \mathbb{N}$ і, аналогічно ML , що для $\forall \epsilon, \delta \in (0,1)$ та розподілу $\mathcal{D}$ на Z (таким, що зберігається умова вимірності на ℓ) за навчальною вибіркою S з $m \geq m_{\mathcal{H}}(\epsilon, \delta)$ елементів, що i.i.d. відносно до $\mathcal{D}$, зі ймовірністю не меншою за параметр впевненості $1-\delta$ повертає гіпотезу $h_S \in \mathcal{H}$, для якої ризик $L_{\mathcal{D}}(h_S)$ обмежений зверху значенням $\min_{h \in \mathcal{H}} L_{\mathcal{D}}(h) + \epsilon$, де ϵ - параметр точності.

$$\mathcal{D}^m(\{S \mid L_{\mathcal{D}}(h_S) \leq \min_{h \in \mathcal{H}} L_{\mathcal{D}}(h) + \epsilon\}) \geq 1 - \delta.$$

Рем. "Априорно" тут означає без "випи" у припущення реалізованості (якщо не було випи, то $\min_{h \in \mathcal{H}} L_{\mathcal{D}}(h) = 0$). Як бачимо, це означення працює для різних типів керування навч. і навіть для некерованого! Як пам'ятає, такі вважати, що $m_{\mathcal{H}}(\epsilon, \delta)$ - найменше, для якого виконана умова - класичність вибірки.

Впр. Показати, що $\forall$ скінченною PAC навчальний клас для задачі класифікації $\in$ PAC навчання (для того не алгоритма).

Достатня умова скінченною PAC навчання – рівномірна збігність

Згадаємо, що алгоритм ML, що реалізує $ERM_{\mathcal{H}}$, мінімізує емп. ризик $L_S(h)$ у класі $h \in \mathcal{H}$; наша ж задача – мінімізувати справжній ризик $L_D(h)$. Розглянемо S , для яких вони близькі:

def. Навчальна вибірка S зветься ϵ -репрезентативною для задачі ML з простором $\mathcal{Z}$, розподілом π і класом гіпотез $\mathcal{H}$ і φ -їм втраат ℓ , якщо $\forall h \in \mathcal{H}$

$$|L_S(h) - L_D(h)| \leq \epsilon.$$

Рн. Якщо (в умових поперед. дев.) $S \in \frac{\epsilon}{2}$ -репрезентативною,
то $\forall$ гіпотези $h_S \in \text{ERM}_{\mathcal{H}}(S)$

$$L_{\mathcal{D}}(h_S) \leq \min_{h \in \mathcal{H}} L_{\mathcal{D}}(h) + \epsilon.$$

Враховуючи, що умова означає $h_S \in \text{argmin}_{h \in \mathcal{H}} L_S(h)$, маємо

$$L_S(h_S) \leq L_S(h) \quad \forall h \in \mathcal{H}. \text{ Тоді}$$

$$L_{\mathcal{D}}(h_S) \leq \left[\frac{\epsilon}{2} - \text{нер.}\right] \leq L_S(h_S) + \frac{\epsilon}{2} \leq L_S(h) + \frac{\epsilon}{2} \leq \left[\frac{\epsilon}{2} - \text{нер.}\right] \leq L_{\mathcal{D}}(h) + \epsilon,$$

$\forall h \in \mathcal{H}$. Перескодав справа до $\min$, отримавши потрібне. Δ

дев. Тобто, що клас гіпотез $\mathcal{H}$ має властивість рівномірної збіжності (uniform convergence) для задані M_L з простору $\mathcal{Z}$ і φ -цією втрач ℓ , якщо $\exists$ така φ -ція $m_{\mathcal{H}}^{\text{UC}}: (0,1)^2 \rightarrow \mathbb{N}$,
що $\forall \epsilon, \delta \in (0,1)$ і розподіл $\mathcal{D}$ на $\mathcal{Z}$ (таким, що вик. умова
випірності на ℓ) навчальна вибірка з $m \geq m_{\mathcal{H}}^{\text{UC}}(\epsilon, \delta)$ елементів,
що і.і.д. відрізняє до $\mathcal{D}$, є ϵ -репрезентативною (для $\mathcal{Z}, \mathcal{D}$,
 $\mathcal{H}, \ell$) зі ймовірністю $\geq 1 - \delta$.

Рем. Тут зазвичай мене беруть $m_{\mathcal{H}}^{\text{UC}}(\epsilon, \delta)$ найменше можливе

натуральне. "Рівномірно" означає незалежність
оцінки від $\mathcal{D}$ або конкретного $\mathbf{h} \in \mathcal{H}$. З означень і
попер. Рн. очевидно випливає:

Сол. Якщо (в умових попер. дов.) клас гіпотез $\mathcal{H}$ має власт.
рівномірної збіжності з деякого $m_{\mathcal{H}}^{VC}$, то він априорно
РАС навчання, причому складність вибірки $m_{\mathcal{H}}(\epsilon, \delta) \leq m_{\mathcal{H}}^{VC}(\frac{\epsilon}{2}, \delta)$,
а алгоритмом може бути будь-який, що здійснює $ERM_{\mathcal{H}}$.

Ця умова дозволяє узагальнити (тобто посилити)
доведення раніше РАС навчання скінч. класів на априорний
випадок:

Рн. Будь-який скінченний клас гіпотез $\mathcal{H}$ має властивість
рівн. збіжності для задачі ML з довільним $\mathbb{Z} \subseteq \mathbb{C}: \mathcal{H} \rightarrow \mathbb{Z} \rightarrow$
 $\rightarrow [0, 1]$, причому
$$m_{\mathcal{H}}^{VC}(\epsilon, \delta) \leq \left\lceil \frac{1}{2\epsilon^2} \ln \frac{2|\mathcal{H}|}{\delta} \right\rceil,$$

а отже априорно РАС навчання для $m_{\mathcal{H}}(\epsilon, \delta) \leq m_{\mathcal{H}}^{VC}(\frac{\epsilon}{2}, \delta)$.

$\bar{i}$ $\forall$ алгоритма з ERM_K .

Рем. Таким чином суттєвого обмеження ρ -ції немає:

$\ell(h, z) \leq 1 \quad \forall h, z$ (це вірно і для $a \leq \ell(h, z) \leq b$, як було введено з доведення, але оцінка m_K^{uc} може змінитися).

$\Rightarrow$ В силу конв. Сок., дост. довести саме рівн. збіжності (і отримати оцінку m_K^{uc}). Початок практично той же, що у дов. Сок. про симетрич. K вище: щоб оцінити збіжність імовірності отримати ϵ -пер. S , розглянемо множини відсію, що такими не є: за дев.,

$$\{S \mid S \text{ не } \epsilon\text{-пер.}\} = \{S \mid \exists h \in K: |L_S(h) - L_\Phi(h)| > \epsilon\} = \\ = \bigcup_{h \in K} \{S \mid |L_S(h) - L_\Phi(h)| > \epsilon\}, \text{ тому за субадитивністю імовірності}$$

$$\mathbb{Q}^m(\{S \mid S \text{ не } \epsilon\text{-пер.}\}) \leq \sum_{h \in K} \mathbb{Q}^m(\{S \mid |L_S(h) - L_\Phi(h)| > \epsilon\}) \quad (\otimes)$$

Розглянемо $S \mapsto L_S(h)$ для фікс. $h \in K$ як випадкову величину (на $\mathcal{Z}^m$ з конв. ім. $\mathbb{Q}^m$). Її мат. оч.:

$$S \stackrel{\sim}{\sim} \mathbb{Q}^m(L_S(h)) = E_{S \sim \{z_i\}_{i=1}^m \sim \mathbb{Q}^m} \left(\frac{1}{m} \sum_{i=1}^m \ell(h, z_i) \right) = \left[\begin{array}{c} \text{імовірність} \\ \text{мат. оч.} \end{array} \right] =$$

$$= \frac{1}{m} \sum_{i=1}^m E \left(\ell(h, z_i) \right) = \left[(z_1, \dots, z_m) \mapsto \ell(h, z_i) - \begin{array}{l} \text{p-гіа big 1 з коор.}, \\ \text{власн. } \mathbb{D}^m \text{ і інтеграла,} \\ \mathbb{D}(z) = 1 \end{array} \right] = \frac{1}{m} \sum_{i=1}^m E \left(\ell(h, z) \right) =$$

$$= [\text{def. } L_{\mathbb{D}}] = L_{\mathbb{D}}(h) \quad (\text{пер. маює з огляду з задан практики})$$

Підмо $S \mapsto |L_S(h) - L_{\mathbb{D}}(h)|$ - відхилення вил. величина big її мат. очікування. Його можна оцінити за допомогою наступної версії закону великих чисел:

Лем. (нерівність Хoeffdinga концентрації міри). Нехай $\{\theta_i\}_{i=1}^m$ - i.i.d. випадкові величини на ймов. пр-ті $\mathbb{Z}$, i
 $E(\theta_i) = \mu \quad \forall i$ Якщо крім того $P(\theta_i \in [a, b]) = 1$ для де-яких $a < b$ і $\forall i$, то $\forall \varepsilon > 0$

$$P \left(\left| \frac{1}{m} \sum_{i=1}^m \theta_i - \mu \right| > \varepsilon \right) \leq 2e^{-\frac{2m\varepsilon^2}{(b-a)^2}}.$$

В Див. додаток В з [1] або курс мат. статистики. $\triangle$

Підмо середні значення наближаються до мат. очікування так, що різницю можна оцінити. У нас для $i=1, m$
 $\theta_i = \ell(h, z_i)$, а отже $\frac{1}{m} \sum_{i=1}^m \theta_i = L_S(h)$ і $\mu = L_{\mathbb{D}}(h)$.

Пл.ч. за пер. $\mathcal{H}$. (при $a=0, b=1$),

$$\textcircled{\leq} \sum_{h \in \mathcal{H}} 2 e^{-2m \epsilon^2} = 2|\mathcal{H}| e^{-2m \epsilon^2} \textcircled{\leq},$$

а отже при $m \geq \frac{1}{2\epsilon^2} \ln \frac{2|\mathcal{H}|}{\delta}$

$$\textcircled{\leq} 2|\mathcal{H}| e^{-\ln \frac{2|\mathcal{H}|}{\delta}} = \delta.$$

Потім $\mathbb{D}^m(\{S \mid S \text{ } \epsilon\text{-порн.}\}) > 1 - \delta$, для малих m , що й потрібно. $\triangle$

Рем. Ця оцінка працює і для (теоретично) нескінченних класів на практиці в силу дискретизації. Дійсно, нехай ϵ -мн $\mathcal{H}$ задані d дійсними параметрами (у наших прикладах з лекцій та практики: $d=2$ для прямокутників, зовільне $d \in \mathbb{N}$ для паралелепедів, $d=1$ для кругів або багатобоків, кулі або для порогових ф-цій $h(x) = \text{sign}(x-a)$ з пар. a), які задані реально, скажімо, 64 бітами кожен. Тоді $|\mathcal{H}| \leq 2^{64d}$, отже за Рн.

$\mathcal{H}$ достатньо PAC навч. для $(\epsilon \in \delta)$

$$m_{\mathcal{H}}(\epsilon, \delta) \leq m_{\mathcal{H}}^{\text{VC}}\left(\frac{\epsilon}{2}, \delta\right) \leq \left\lceil \frac{2}{\epsilon^2} \ln \frac{2|\mathcal{H}|}{\delta} \right\rceil \leq \left\lceil \frac{2}{\epsilon^2} \ln \frac{2 \cdot 2^{64d}}{\delta} \right\rceil = \left\lceil \frac{2 \ln \frac{2}{\delta} + 128d \cdot \overbrace{\ln 2}^{\frac{1}{\sqrt{2}}}}{\epsilon^2} \right\rceil.$$