

Покрі коректно визначена $h: X \rightarrow \{0, 1\} : x \mapsto \begin{cases} h_n(x), & x \in X_n, \\ 0 & \text{у інш. випад.} \end{cases}$
 $n = 1, \infty$

При цьому $h \in \mathcal{H}$ і $\forall n \ h \notin \mathcal{H}_n$, бо $h|_{X_n} = h_n$ не $\in$ обм.
ф-цій з $\mathcal{H}_n$ за побудовою. $\downarrow$ Δ
мінімізація

Приклад SRM: формальний опис

Покі що у наших прикладах вагових ф-цій w значення $w(n)$ не були жовно пов'язані з класами $\mathcal{H}_n$ (крім прикладу з поліномами, де, скажімо, для $w(n) = \frac{6}{(n+1)^2}$ це значення спадає з $n = \text{Vcdim}(\mathcal{H}_n)$ - степенем полінома + 1). Побудуємо приклад, де w визначена класами.

Класи $\mathcal{H}$ - (не білим ліне) зліченний. Покі що можна предст. як зліченне об'єднання одноточкових класів:

$\mathcal{H} = \bigcup_n \underbrace{\{h_n\}}_{\mathcal{H}_n}$. За РЧ. про рівн. збіжність скінч. класів,
усі $\{h_n\}$ рівн. збіжні з $m_{\{h_n\}}^{VC}(\epsilon, \delta) \leq \frac{1}{2\epsilon^2} \ln \frac{2}{\delta} \sqrt{V_n}$ за умови,
що ф-цією втрат даної задачі $ML \in \ell: \mathcal{H} \times \mathbb{Z} \rightarrow [0, 1]$

(наприклад, для задачі класифікації з $l = l_{0-1}$). Тоді в сумі (*)

$$\varepsilon_n(m, \delta) \leq \sqrt{\frac{1}{2m} \ln \frac{2}{\delta}}, \text{ тоді SRM можна реалізувати як}$$
$$h_S \in \underset{h \in \mathcal{H}}{\operatorname{argmin}} \left(L_S(h) + \sqrt{\frac{1}{2m} \ln \frac{2}{w(h) \cdot \delta}} \right) = -\ln w(h) + \ln \frac{2}{\delta}.$$

де менш вважимо $w: \mathcal{H} \rightarrow [0, 1]: h_n \rightarrow w(h)$ (то менш $\mathcal{H}$ у біж вимірності з N). Визначимо менш $w(h)$ у вимірності го говорячим опису h :

обв. Нехай Σ - алфавіт, тоді деяка скінченна множина символів. Рядок (string) - це будь-яка послідовність символів з Σ ; позн. через $|s|$ говорячи рядка s , а через Σ^* - множину всіх скінченних рядків з символів Σ .

Мова опису (description language) на множині $\mathcal{H}$ (скажімо, на класі інших задачі $\mathcal{ML}$) - це будь-яке вигорансена $d: \mathcal{H} \rightarrow \Sigma^*$. $\forall h \in \mathcal{H}$ його образ $d(h)$ тоді зветься описом h , а $|h| := |d(h)|$ - говорячим опису в $\mathcal{ML}$. Мова опису d зветься префіксною (prefix-free) (визначено big пре-

фіксів, prefix-free), якщо $\forall h \neq h' \in \mathcal{K}$ жоден з рядків $d(h)$, $d(h')$ не є початком (префіксом) іншого (зокр., не $i \cap j$).

Ех. Наприклад, для двійкового коду $\Sigma = \{0, 1\}$, а $\Sigma^* = \{0, 1\}^*$ усі скінченні рядки вигляду $(0, 1, 1, 0, \dots)$. Прикладом префісного $d \in$ код Хаффмана (а для $\Sigma = \{0, \dots, 9\}$ — телефонні коди країн). Для класу мов $\mathcal{K}$ конкретною реалізацією d може бути, наприклад, запис $h \in \mathcal{K}$ як програми мовою C++ та застосування до її тексту гзп (що містить кодування за Хаффманом, а отже реалізує префісну мову). Тоді $d: \mathcal{K} \rightarrow \{0, 1\}^*$ префісна.

Рл. (нерівність Крафта) $\forall$ зліченній підмножині рядків $S \subset \{0, 1\}^*$, жоден з яких не є префіксом

іншого, задана, для $S = d(\mathcal{K})$, де $d: \mathcal{K} \rightarrow \{0, 1\}^*$ префісна,

$$\sum_{s \in S} \frac{1}{2^{|s|}} \leq 1.$$

► Один зі способів доведення: побудуємо на S ймовір-

істинний розподіл (як на дискретному ім. просторі):
 $\forall b \in S$ нехай $P(b)$ - це ймовірність отримати b ,
 підкладаючи кожною разю монетку і записуючи $\text{flip}(\text{bit})$
 (0 або 1) у рядок. В силу префіксності, цей процес може
 привести лише до одного $b \in S$ (або до нічого) зі ім.

$$P(b) = \frac{1}{2^{|b|}}. \text{ П.ч.}, \sum_{b \in S} \frac{1}{2^{|b|}} = \sum_{b \in S} P(b) \leq 1. \triangle$$

Лем. В силу лем-сти Крафта, $w: h \mapsto \frac{1}{2^{|h|}}$, де $|h| =$
 $= |d(h)|$ (визначається ф-н. до деякої префіксної мови опису d ,
 може бути обрана як вагова ф-ція на $\mathcal{H}$: довжи за
 описом іпотези мають меншу вагу. Отримана вище
 оцінка на ϵ_n тоді набуває вигляду

$$\epsilon_n(m, w(h) \cdot \delta) \leq \sqrt{\frac{1}{2m} (-\ln w(h) + \ln \frac{2}{\delta})} = \sqrt{\frac{1}{2m} (\ln 2^{|h|} + \ln \frac{2}{\delta})} < \sqrt{\frac{1}{2m} (|h| \ln 2 + \ln \frac{2}{\delta})}.$$

Тоді з одного з Рл. вище отримуюємо:

Сл. Нехай $\mathcal{H}$ - клас іпотез ML на $\mathcal{Z}$ з ф-цією втрат $\ell: \mathcal{H} \times \mathcal{Z} \rightarrow [0, 1]$,
 $\mathcal{H}$ зліченний, і $d: \mathcal{H} \rightarrow \{0, 1\}^*$ - префіксна мова його опису.

Тілоді $\forall \delta \in (0, 1)$, $m \in \mathbb{N}$ і розподілу $\mathcal{D}$ на $\mathbb{Z}$ зі ім. $\geq 1 - \delta$ має вибірку S з m ел-тів, що i.i.d. від $\mathcal{D}$,

$$\forall h \in \mathcal{H} \quad L_{\mathcal{D}}(h) \leq L_S(h) + \sqrt{\frac{1}{2m} (|h| + \ln \frac{2}{\delta})}.$$

Це, у свою чергу, мотивує наст. версію SRM:

Лем. Творять, що алгоритм A (ML для задачі на $\mathbb{Z}$ з φ -ною

втрата $0 \leq \ell \leq 1$) згінює мінімізацію довжини опису

(minimum description length, MDL) у згінюваному класі

іпозе $\mathcal{H}$ з префіксною мовою опису $d: \mathcal{H} \rightarrow \{0, 1\}^*$, якщо

$\forall m \in \mathbb{N}$ і $\delta \in (0, 1)$ за навч. вибіркою S з m ел-тів він повертає іпозе

$$h_S = A(S) \in \argmin_{h \in \mathcal{H}} \left(L_S(h) + \sqrt{\frac{1}{2m} (|h| + \ln \frac{2}{\delta})} \right).$$

Рем. Як і у загальному SRM, тут виконає упередженість на користь h з малою $|h|$: іпозе з коротким описом заслуговує на більшу довіру (принцип Бейтса Оккама).

Втім, це порівняння має сенс лише в рамках оціні

мови d і сукупово залежність від мови: вибір гра-
вальної $d \in \pi$ тут випадковою частиною розуміння
конкретної задачі ML (разом з вибором класу $\mathcal{K}$ і
ф-ції втрат ℓ).

Послідовність алгоритмів ML

У свою чергу, деб. нерівномірної навчальності можна це
покласти, розв'язавши складати вибірки залежати не
лише від ϵ, δ і $h \in \mathcal{K}$, але й від розподілу $\mathcal{D}$ на $\mathcal{Z}$:

деб. Нехай задача ML визначена областю визначення $\mathcal{Z}$,
класом гіпотез $\mathcal{K}$ і ф-цією втрат ℓ , а $\mathcal{P}$ - якась множина
ймовірнісних розподілів на $\mathcal{Z}$. Алгоритм A для цієї за-

дачі зветься послідовним (несигперечним, consistent) відносно до $\mathcal{K}$ і $\mathcal{P}$,

якщо існує ф-ція $m_{\mathcal{K}}^{\text{con}} : (0,1)^2 \times \mathcal{K} \times \mathcal{P} \rightarrow \mathbb{N}$ така, що для

будь-яких $\epsilon, \delta \in (0,1)$, $h \in \mathcal{K}$ і $\mathcal{D} \in \mathcal{P}$, якщо $m \geq m_{\mathcal{K}}^{\text{con}}(\epsilon, \delta, h, \mathcal{D})$, то зі ім. $\geq 1-\delta$ A повертає за вибіркою з m

ел-нів, що $i.i.d.$ from до $\mathcal{D}$ гіпотезу $A(S)$ з ризиком

$$L_{\mathcal{D}}(A(S)) \leq L_{\mathcal{D}}(h) + \epsilon.$$

Зокрема, якщо $\mathcal{P}$ тут - сукупність усіх ім. розподілів на $\mathcal{Z}$, то A зветься універсально послідовним (universally consistent) для $\mathcal{K}$.

Рем. Нагадаємо, що тут на $\mathcal{Z}$ визначена деяка $\mathcal{B}$ -алгебра, і всі $\mathcal{D} \in \mathcal{P}$ повинні бути визначені на ній (можі "усі розподіли" - усі, що визначені на цій $\mathcal{B}$ -алг.). З дов. ванн:

Соч. Якщо клас гіпотез $\mathcal{K}$ (у задачі ML на $\mathcal{Z}$ з $\mathcal{P}$ -цією втрат ℓ) періодично навчаний, то будь-який алгоритм A , що задов. умові період. навчальності $\mathcal{K}$, є універсально послідовним для $\mathcal{K}$.

Рем. При цьому обмеження (універсальності) послідовності функції слабше: існують алгоритми з цією властивістю для очевидно "поганих" класів, наприклад, для класу

усіх bin. класифікаторів на $X = N$ (тобто φ -цін $h: N \rightarrow \{0, 1\}$), що, як ми показали у Th. вище, не є перивн. навчанням!

Ex. Нехай $Z = X \times Y$, де $X \stackrel{(\leq)}{\sim}$ зліченна, а Y скінченна, тобто маємо задачу класифікації на X (з дов. числом класів), $l = l_0 - 1$. Розглянемо алгоритм "запам'ятовування" навчальної вибірки:

$$\text{MEMORIZE}(S) := \left(h_S : x \mapsto \begin{cases} \arg\max_{y \in Y} |\{(x, y) \in S\}|, & \text{якщо } (x, y) \in S \text{ для деякого } y; \\ y_0 & \text{інакше випадку.} \end{cases} \right)$$

де y_0 - якась фіксована мітка. Це загальнення прикладу перенавчанняї гіпотези для дов. клас., що був наведений раніше. Такі гіпотези, як завжди, можуть давати дуже високі $L_\Phi(h_S)$ для багатьох Φ , а отже цей алгоритм очевидно "поганий". Але він улюб. посл. для (менш "поганих") класу усіх $h: X \rightarrow Y$. Дійсно, покажемо це для випадку, коли Y ат. функ. χ ($\exists$ "незалежна" $f: X \rightarrow Y$).

Нехай $\mathcal{D}$ - довільний розподіл на зліч. $\mathcal{X}$, $\mathcal{X}$ - простір дискр.,
 тобто $\mathcal{D}$ визначений своїми значеннями на 1-точкових підмнож. $\mathcal{X}$,
 і $\sum_{x \in \mathcal{X}} \mathcal{D}(\{x\}) = 1$ - абсолютно збіжний ^(7,0) ряд. За власт. монотонності
 рядів, його ел-ти можна переставити, впорядковуючи за спа-
 даючим, що відп. деякому впорядкуванню $\mathcal{X}$: $\mathcal{X} = \{x_n\}_{n=1}^{\infty}$, і
 $\mathcal{D}(\{x_1\}) \geq \mathcal{D}(\{x_2\}) \geq \dots \geq \mathcal{D}(\{x_n\}) \geq \mathcal{D}(\{x_{n+1}\}) \geq \dots$. Крім того,
 у збіжного $\sum_{n=1}^{\infty} \mathcal{D}(\{x_n\}) = 1$ посл. загальної чл. сума прямує до 0:
 $\sum_{n=N}^{\infty} \mathcal{D}(\{x_n\}) \xrightarrow{N \rightarrow \infty} 0$. Оскільки вона не зростає (бо усі $\mathcal{D}(\{x_n\}) \geq 0$) це
 означає, що $\forall \varepsilon > 0 \exists N \in \mathbb{N}: \varepsilon > \sum_{n=N}^{\infty} \mathcal{D}(\{x_n\}) = [\text{адит.}] = \mathcal{D}(\{x_n\}_{n=N}^{\infty})$.

Якщо тепер вибірка $S \in \mathcal{X}^m$ містить усі x_n для $n = \overline{1, N-1}$,
 то т.ч. $\mathcal{D}(\{x | x \notin S\}) < \varepsilon$. За побудовою $A = \text{MEMORIZE}$, маємо
 і $L_{\mathcal{D}}(A(S)) < \varepsilon$, бо "печивковими" можуть бути лише точки
 за межами S . Т.ч.,

$$P_{S \sim \mathcal{D}^m} (L_{\mathcal{D}}(A(S)) > \varepsilon) \leq \mathcal{D}^m(\{S | \exists n = \overline{1, N-1} : x_n \notin S\}) \leq [\text{субад.}] \leq$$

$$\leq \sum_{n=1}^{N-1} \mathbb{D}^m(\{S \mid x_n \notin S\}) = \sum_{n=1}^{N-1} \mathbb{D}(\mathcal{X} \setminus \{x_n\})^m = \sum_{n=1}^{N-1} (1 - \mathbb{D}(\{x_n\}))^m \quad (\leq)$$

Позначимо $\eta := \min_{n=1, N-1} \mathbb{D}(\{x_n\})^{\frac{e(0,1)}{e(0,1)}}$. Як $i \in N$, це значення залежить
лише від $\mathbb{D}$ і ε . Тоді

$$(\leq) \sum_{n=1}^{N-1} (1 - \eta)^m = (N-1)(1 - \eta)^m < (N-1)e^{-m\eta}$$

Для $m > \frac{1}{\eta} \ln \frac{N-1}{\delta}$ це $\leq \delta$. Т.ч., $\forall h \in \mathcal{K}$

$$\mathbb{P}_{S \sim \mathbb{D}^m} (L_{\mathbb{D}}(A(S)) \leq \varepsilon \leq L_{\mathbb{D}}(h) + \varepsilon) \geq 1 - \delta$$

для $m > m_{\mathcal{K}}^{\text{con}}(\varepsilon, \delta, \mathbb{D})$, де $m_{\mathcal{K}}^{\text{con}}(\varepsilon, \delta, \mathbb{D}) \leq \left\lceil \frac{1}{\eta} \ln \frac{N-1}{\delta} \right\rceil$ не

залежить від $h \in \mathcal{K}$. Звідси і маємо унів. $\left\{ \begin{array}{l} \text{залежить від} \\ \mathbb{D}, \varepsilon \end{array} \right.$
послідовність.

Тим за рахунок залежності $m_{\mathcal{K}}^{\text{con}}$ від $\mathbb{D}$ вдалося забезпечити,
що δ (зі ім. $\geq 1 - \delta$) S містила усі ел-ти $\mathcal{X}$ з дост. великими $\mathbb{D}$.

Звичайно, це все ніяке не "павчак". Цей алгоритм послі-
довний лише у тому сенсі, що збільшення S гаранто-
вано $\rightarrow$ зменшує помилку.

Ех Прикладом перекладового будзе "апорт", що, Спаксно,
завжди повертає одну і ту саму n або об'єкт і впадково
визначено fig 5.