

Machine Learning: Theory & Algorithms I

Mathematics of Machine Learning

Математика машинного навчання

Література

1. S. Shalev-Shwartz, S. Ben-David. Understanding Machine Learning.
2. M. Mohri, A. Rostamizadeh, A. Talwalkar. Foundations of Machine Learning.
3. D.J. MacKay. Information Theory, Inference, and Learning Algorithms.

Машинне навчання ^(ML) — область інформатики (більш точно, н.зв. штучного інтелекту), що досліджує статистичні алгоритми, які здатні модифікувати себе з використанням відомих даних ("навчатися") і робити змістовні висновки з раніше невідомих даних ("застосовувати

набуті знання").

Застосування: у задачі, де побудова явного алгоритма на противагу статистичному, ("традиційний" підхід штучного інтелекту - імітація людської мислення) дуже складна (як правило, практично неможлива).

Зокрема,

— задачі, з якими добре вправляються люди або тварини: розпізнавання образів, мови, інші задачі класифікації (виявлення стема), переклад, спілкування природними мовами, водіння транспортних засобів;

— задачі, навпаки, за межами людських можливостей: аналіз великих об'ємів даних у науці, медицині, бізнесі, знаходження прихованих закономірностей.

З іншого боку, відмінність від статистика — орієнтація на алгоритми, відсутність попереднього визначення гіпотези про розподіл даних.

Намі задачі у цьому курсі:

1. Побудувати загальні математичні моделі матричного навчання та дослідити їхні властивості.

2. Застосувати цей загальний підхід до конкретних алгоритмів ML.

Більш точно, нас цікавить перше за все кероване статистичне пакетне пасивне навчання. Значення цих слів:

Кероване (з учителем supervised) - зразки даних, на яких алгоритм навчається, ^(training data) мають суцільну інформацію, що відсутня у даних, до яких він застосовується і яку потрібно відповісти (мітки класифікації, розшифровки, переклади, продовження фраз...) на відміну від некерованого (без учителя, unsupervised), де навчальні та тестові (робочі) дані не відрізняються (наприклад, задачі кластеризації)

Статистичне у цьому контексті означає, що навчальні дані згенеровано деяким випадковим процесом, на відміну від ситуації "вчителя", що допомагає чи заважає.

Пакетне (batch learning) — алгоритми застосовуються до тестових даних лише після обробки суттєвої кількості навчальних, на відміну від онлайн-навчання, де ці процеси можуть відбуватися одночасно.

Пасивне — алгоритм лише сприймає інформацію, що йому надають, на відміну від активного, коли він може робити запити до оточення під час навчання.

Початкова математична модель ML

Елементи моделі:

— Область визначення X (domain set) — довільна множина об'єктів, що потрібно класифікувати. Часто вписують $X \subset \mathbb{R}^d$; у цьому випадку координати $x \in X$ зв.

іною ознаками (features), тобто x - вектор ознак.

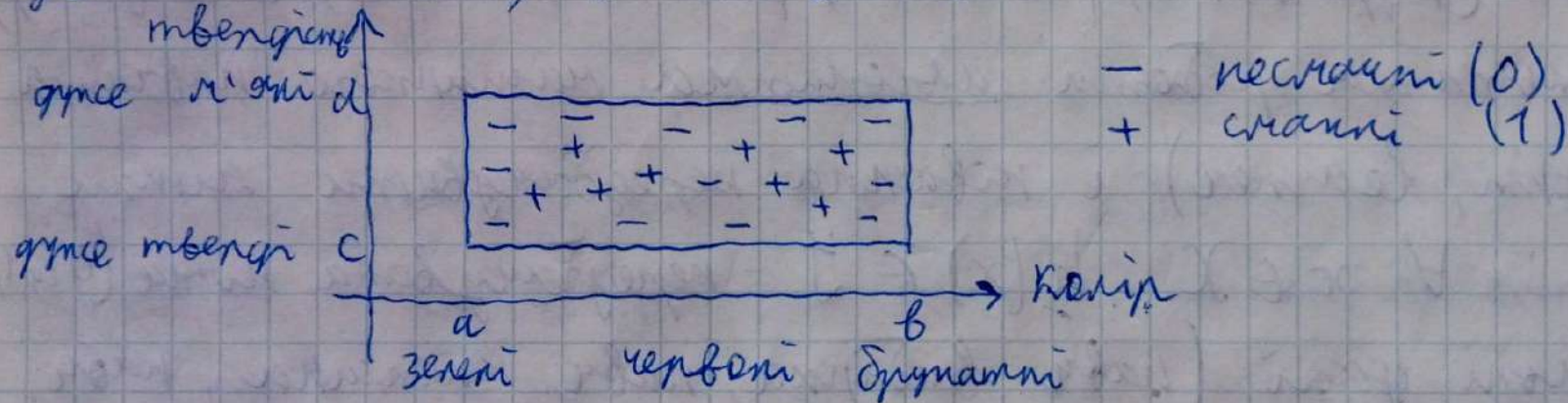
- Множина міток (область значень) Y (label set) - тут розглянемо задачу бінарної класифікації, тобто на 2 класи, тоді $Y = \{0, 1\}$ (часто беруть $Y = \{-1, +1\}$), де 0 і 1 означають належність $x \in X$ до одного з класів.

- Гіпотеза (предиктор, правило передбачення, класифікатор)
 $h: X \rightarrow Y$ (hypothesis, predictor, prediction rule, classifier),
що є результатом роботи алгоритма машинного навчання
 A (algorithm, learner) і повинна передбачувати мітки
ел-тів X : $\forall x \in X \quad h(x) \in Y$ - передбачувана мітка (0 чи 1)

- Навчальні дані (навч. вибірка, навч. множина, навч. приклади) S , що отримувє алгоритм (training data, set, sample).
Оск. навчання кероване, S тут є послідовністю $S = \{(x_1, y_1), \dots, (x_m, y_m)\}$ m пар з точок $x_i \in X$ та їхніх міток $y_i \in Y$, що ми вважаємо правильними. Для належного

пасивного навчання іпотега повинна визначатися алгоритмом за S : $h = A(S)$.

Ех. Задача визначення смачності певного фрукта за 2 ознаками: $X = [a, b] \times [c, d] \subset \mathbb{R}^2$ - прямокутник, де координати фрукта $x \in X$ означають його колір та твердість (у деякому числовому вираженні), $Y = \{0, 1\}$, де 0 - не смачний, 1 - смачний:



Навчальні дані $\{(x_i, y_i)\}_{i=1}^m$ - результат збірання m випадкових фруктів, визначення їх ознак та того, чи є вони смачними. h - функція, що за ознаками, визначає чи фрукт смачний. Ключовими є 2 наступних питання:

Що ми пропускаємо про побудову S ?

Оск. навчання статистичне, будемо пропускати, що
у $S = \{(x_i, y_i)\}_{i=1}^m$, точки $x_i \in X$ обираються випадковим
чином згідно з ймовірнісним розподілом $\mathcal{D}$ на X , що
моделює властивості оточення (у $\underline{e}_x$ - сарту, де ростуть фрукти).
Це означає, що $(X, \mathcal{F}, \mathcal{D})$ - ймовірнісний простір, тобто
- задана σ -алгебра $\mathcal{F}$ підмножин X (випірних) -
дедна сукупність підм. $A \subset X$, що

- містить X : $X \in \mathcal{F}$;
 - замкнена відп. доповнень: $A \in \mathcal{F} \Rightarrow X \setminus A \in \mathcal{F}$;
 - замкнена відп. ($\leq$) зліченних об'єднань: $\{A_i\}_{i=1}^{\infty} \subset \mathcal{F} \Rightarrow \bigcup_{i=1}^{\infty} A_i \in \mathcal{F}$ (де замість ∞ може бути скін. n).
- (звідси вона $\mathcal{F}$ замкнена й відп. до перетинів за правилом де Моргана).
Зокрема, для $X \subset \mathbb{R}$ $\mathcal{F}$ можна вважати σ -алгеброю Бoreльовою
(тобто такою, що містить усі відкр., замкнені інтервали,
усі їхні $\leq$ зліч. перетини та об'єднання). Далі ми,

Як правило, не будемо вважати $\mathcal{F}$ явн., просто зов-
рачн про $\mathcal{D}(A)$ для $A \in \mathcal{X}$.

— задана ймовірнісна міра $\mathcal{D} : \mathcal{F} \rightarrow [0, 1]$ з
властивостями:

- $\mathcal{D}(X) = 1$

- зліченної адитивності: $\forall (\leq)$ зліч. набору під-
множин $\{A_i\}_{i=1}^{\infty} \subset \mathcal{F}$ (тут і далі замість ∞ може бути
скінч. n), що попарно не перетинаються: $A_i \cap A_j = \emptyset$
для $\forall i \neq j$

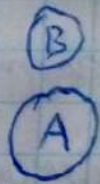
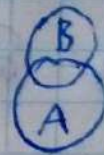
$$\mathcal{D}\left(\bigcup_{i=1}^{\infty} A_i\right) = \sum_{i=1}^{\infty} \mathcal{D}(A_i)$$

Тут зовн. підм. $A \in \mathcal{F}$ зветься подією, а $\mathcal{D}(A) \in [0, 1]$ —
ймовірність події A . Вона виігрує, з якою ймовірністю
випадково обране $x \in X$ належатиме до A . Інші власт. $\mathcal{D}$:

- монотонність: $A \subset B \Rightarrow \mathcal{D}(A) \leq \mathcal{D}(B)$

(бо $\mathcal{D}(B) = \mathcal{D}(A \cup (B \setminus A)) = \mathcal{D}(A) + \mathcal{D}(B \setminus A) \geq \mathcal{D}(A)$).

• субадитивність: $\mathcal{D}(\bigcup_{i=1}^{\infty} A_i) \leq \sum_{i=1}^{\infty} \mathcal{D}(A_i)$ для $\forall$ набору підмнож. де замість ∞ знову може бути n .



(це зветься обмеженістю об'єднання - union bound).

$\mathcal{D}(X|A) = 1 - \mathcal{D}(A)$ (чому?). Еск. Для $X \subset \mathbb{R}^d$ прамагам є острівний розподіл: $\mathcal{D}(A) := \frac{V(A)}{V(X)}$, де V - об'єм (міра Лебєга) на $\mathbb{R}^d$, та нормований. Отже, обираючи $x_1, \dots, x_m \in X$ згідно з $\mathcal{D}$ випадковим чином.

Також тут будемо припускати існування "правильної" функції класифікації (labeling function) $f: X \rightarrow Y$, що визначає за ознаками $x \in X$ мітку $y = f(x) \in Y$, і до якої ми повинні "наблизити" (дів. максимізувати) значення h . Тоді $\forall i$ визначимо $y_i := f(x_i)$ і таким чином побудуємо $S = \{(x_i, y_i = f(x_i))\}_{i=1}^m$ (Еск.: смачність одно-значно визначена кельсором і твердістю).

П.ч., S озназн. визначено набором $\{x_i\}_{i=1}^m \in X^m$, де $x_i \in X$ незалежно та ідентично розподілені (одинково) (independently and identically distributed, i.i.d.);

консен з них обертається з X незалежно від інших
 згідно з $\mathcal{D}$. Подібно на X^m розглядаємо добуток
 імовірнісних мір $\mathcal{D}^m := \underbrace{\mathcal{D} \times \dots \times \mathcal{D}}_m$, яка міра за-
 дана на $\mathcal{F}^m$ (найменший σ -алг. на X^m , що містить
 усі $A_1 \times \dots \times A_m$, де $A_i \in \mathcal{F} \forall i$) і однозначно визначена
 умовою $\mathcal{D}^m(A_1 \times \dots \times A_m) = \mathcal{D}(A_1) \cdot \dots \cdot \mathcal{D}(A_m)$. Далі
 ми будемо за її допомогою вимірювати множини,
 що складаються з вибірок S (строю канєди,
 $S \in (X \times Y)^m$, але оскільки в $S = \{(x_i, y_i)\}_{i=1}^m$ $y_i = f(x_i)$,
 тому бієктивно відповідає $\{x_i\}_{i=1}^m \in X^m$; множини
 цих наборів ми й вимірюємо).

Як ми вимірюємо успіх, тобто "гарність" гіпотези h ?

Потім ми маємо "правильну" $f: X \rightarrow Y$, з якою порів-
 няно порівняти h : чим вони "далі", тим гірше.

деф. Кошикової гіпотези h (справжньою помилкою

помилково узагальнення, (справданий) ризиком) ((true) error, generalization error, (true) risk) зветься

$$L_{D, f}(h) := \mathbb{P}(\{x \in \mathcal{X} \mid h(x) \neq f(x)\}) \in [0, 1]$$

тобто ймовірність того, що випадково обрана точка x буде мати неправильну мітку $h(x)$ від h (на відміну від правильної $f(x)$).

Рем. $L_{D, f}(h)$ дозволяє оцінити "якість" h , але не може бути використана при її знаходженні, бо залежить від D і f , а алгоритм A їх не "знає", лише S . Але він може обчислити наступне:

def. Помилково навчання гіпотези h на навчальних даних $S = \{(x_i, y_i)\}_{i=1}^m$ (емпіричного помилково, емпіричний ризик) ((training error, empirical error, emp. risk) зветься

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

(де $|\cdot|$ означає кільк. ел-тів скінченної множини), тобто

частина точок $\underbrace{(x_i, y_i)}_Z$ навч. даних S , для яких h повертає не-
правильну (у порівнянні з відомою y_i) гіпотезу $h(x_i)$.
Тепер ми можемо визначити основний елемент (навчальної
парадигми) алгоритма:

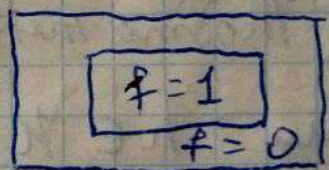
def. Творять, що алгоритм навч. А здійснює
мінімізацію емпіричного ризику (empirical
risk minimization, ERM), якщо за навчальними даними
 S він повертає гіпотезу h з найменшим значенням $L_S(h)$:
 $A(S) = h \in \arg \min_h L_S(h) := \{h: X \rightarrow Y \mid L_S(h) \leq L_S(h') \forall h': X \rightarrow Y\}$
Підто ми обираємо h , що найкраще описує S .

Рем. Цей (дуже природний) підхід крім складностей у
реалізації (треба перебрати усі відображення $h: X \rightarrow Y$)
має суттєвий принциповий недолік: він може
давати h з дуже високим $L_{D, Y}(h)$. Ця традиційна
для ML проблема зветься перенавчанням (overfitting).

Ex. Для $Y = \{0, 1\}$ розглянемо гіпотезу

$$h_S(x) := \begin{cases} y_i & , \exists i = \overline{1, m} : x = x_i; \\ 0 & \text{у іншому випадку.} \end{cases} \quad (*)$$

За побудовою $L_S(h_S) = 0 \quad \forall S = \{(x_i, y_i)\}_{i=1}^m$. Але якщо при цьому X нескінченна, а $\mathcal{D}$ такий, що скінч. класи мають міру 0, то $L_{\mathcal{D}, f}(h_S) = \mathcal{D}(\{x \mid f(x) = 1\})$, бо h_S нульове майже всюди (міра 1). Це може бути дуже високе значення. Наприклад, у Ex. 7 приклади, однаковим $\mathcal{D}$ і f , що задається зовнішньою



площа внутр. пр-кутника у 2 рази менша за площу зовнішнього, $L_{\mathcal{D}, f}(h_S) = \frac{1}{2} \quad \forall S$, тобто h_S "нічого не знає" про $\mathcal{D}$ і f - бо нагто добре підходить до S .

Стандартний розв'язок проблеми перенавчання - ERM з індуктивною упередженістю

Ідея: попередньо (до отримання S) фіксовано клас гіпотез, тобто підкласи $\mathcal{H}$ у класі усіх $h: X \rightarrow Y$.

лев. Тобто, що алгоритм наш. навч. А здійснює ERM з індуктивною упередженістю (inductive bias),

в напрямку класу гіпотез $\mathcal{H}$, якщо за навчальними даними S він повертає гіпотезу h з найменшим значенням $L_S(h)$ серед $h \in \mathcal{H}$:

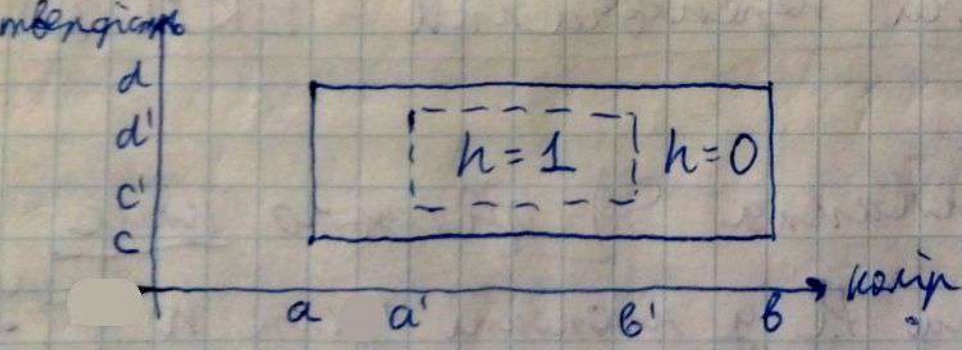
$$A(S) = \underset{h \in \mathcal{H}}{\text{argmin}} L_S(h) := \{h \in \mathcal{H} \mid L_S(h) \leq L_S(h') \forall h' \in \mathcal{H}\}$$

Рем. Подібно тепер треба найкращим чином описати S , використовувати $h \in \mathcal{H}$. Дати цей алгоритмічно позначити $ERM_{\mathcal{H}}$:

$$h_S \in ERM_{\mathcal{H}}(S) \text{ озн. } \exists A \in ERM_{\mathcal{H}} : h_S = A(S).$$

Вибір $\mathcal{H}$ залежить від конкретної задачі ML.

Ек. 1. В задачі про смачність фруктів природно очікувати,



що смачними будуть такі, твердість і колір яблук - з певних проміжків, тобто $\mathcal{H}$ - це клас ф-цій, що

визначені припокриваннями ці сторони, що паралельні осей координат:

$$h(x) = \begin{cases} 1, & x \in [a', b'] \times [c', d']; \\ 0, & x \in [a, b] \times [c, d] \setminus [a', b'] \times [c', d'] \end{cases}, \text{ де}$$

$[a', b'] \subset [a, b]$, $[c', d'] \subset [c, d]$. Це дійсно "гарний" клас (див. нижче).

Ел. 2. Вітимо, не усі класи є корисними. Нехай $\mathcal{K}$ - клас м.зв. порогових (~~thresholded~~) поліномів від координат точок $x \in \mathbb{R}^d$:

$$h(x) = \begin{cases} 1, & P(x) \geq 0; \\ 0, & P(x) < 0, \end{cases}$$

де P - поліном.

Впр. Побудувати за $S = \{(x_i, y_i)\}_{i=1}^m$ поліном P_S такий, що $y_i \text{ sign}(P_S(x_i)) = y_i$ (*).

Це означає, що вибір класу $\mathcal{K}$ призведе до переобчислення.

Ел. 3 Нехай клас $\mathcal{K}$ скінченний: $|\mathcal{K}| < \infty$. На практиці усі вони є такими, бо алгоритм, що генерує h , повинен

закінчити роботу за скін. час, є дискретизація (наприклад, $\mathcal{H}$ з Ex. 1 перетворюється на скін., якщо брати a', b', c', d' зі стандартних мнів дійсних чисел у нашій мові програмування) etc. Крім, $|\mathcal{H}|$ може бути дуже великою.

Показано, що будь-який скін. $\mathcal{H}$ - "гарний" у певному сенсі. Для цього, ^{заздалегідь} знадобиться ще одне припущення:

def. Говорять, що клас гіпотез $\mathcal{H}$ задовольняє припущенню realizability (realizability assumption), якщо існує $h^* \in \mathcal{H}$ така, що $L_{\mathcal{D}, f}(h^*) = 0$.

Rem. Це $\bar{\mu}$ означає, що цей $\mathcal{H}$ добре підходить для даної задачі. Отже, $f(x) = h^*(x)$ для x з підмножини X міри 1, тобто для "найменш конічно" $x \in X$: це так зі $\bar{\mu}$ ймовірністю 1. Зокрема, для "найменш конічної" вибірки S , тобто зі $\bar{\mu}$ ймовірністю 1, $L_S(h^*) = 0$, бо $h^*(x_i) = f(x_i) = y_i \quad \forall (x_i, y_i) \in S$. Це означає, що для $\forall h_S \in \text{ERM}_{\mathcal{H}}(S)$, тобто для $\forall$ реалізованих

роботи алгоритма $\exists ERM_{\mathcal{H}}$ на цій вибірці, $L_S(h_S) = 0$
 (де необхід'язково $h_S = h^*$: банально, що найкраща гіпотеза = 0).

Але для нас банальна справність гіпотези $L_{\mathcal{D}, f}(h_S)$. Нехай
 $\epsilon \in (0, 1)$. Яка "кількість" вибірок $S \in \mathcal{S}$ потрібно, щоб

$L_{\mathcal{D}, f}(h_S) > \epsilon$? (ϵ тут зветься параметром точності (accuracy))

Після ϵ залежи (наприклад, може просто m раз повторитися
 одна й та сама пара (x, y)). Відхилюючи гіпотезу S

рівні 0, можемо вважати, що для такої S $L_S(h_S) = 0$,
 але h_S погана, тобто наслідок до відхилення

$$\mathcal{H}_B := \{h \in \mathcal{H} \mid L_{\mathcal{D}, f}(h) > \epsilon\}.$$

Пл. 1. (крім відхилення рівні 0, банально)

$$\{S \mid \exists h_S \in ERM_{\mathcal{H}}(S) : L_{\mathcal{D}, f}(h_S) > \epsilon\} \subseteq \{S \mid \exists h \in \mathcal{H}_B : L_S(h) = 0\}.$$

Потім в силу монотонності

$$\mathcal{D}^m(\{S \mid \exists h_S \in ERM_{\mathcal{H}}(S) : L_{\mathcal{D}, f}(h_S) > \epsilon\}) \leq \mathcal{D}^m(\{S \mid \exists h \in \mathcal{H}_B : L_S(h) = 0\})$$

Представимо цю групу множин як об'єднання

$\bigcup_{h \in \mathcal{H}_B} \{S \mid L_S(h) = 0\}$ за фіксації елемента $\mathcal{H}_B$:

$$\textcircled{=} \mathcal{D}^m \left(\bigcup_{h \in \mathcal{H}_B} \{S \mid L_S(h) = 0\} \right) \leq [\text{субадитивність } \mathcal{D}^m] \leq \sum_{h \in \mathcal{H}_B}$$

$$\mathcal{D}^m(\{S \mid L_S(h) = 0\}) \textcircled{=}$$

Нагадаємо, що суми тут скінченні. $L_S(h) = 0$ для $S = \{(x_i, y_i)\}_{i=1}^m$

означає, що $h(x_i) = y_i \quad \forall i = \overline{1, m}$, тобто це декартовий

добуток класів: $\{S \mid L_S(h) = 0\} = \{x_1 \in X \mid h(x_1) = f(x_1)\} \times \dots \times \{x_m \mid h(x_m) = f(x_m)\}$

$$\textcircled{=} \sum_{h \in \mathcal{H}_B} \mathcal{D}^m \left(\prod_{i=1}^m \{x_i \in X \mid h(x_i) = f(x_i)\} \right) = [\text{власн. } \mathcal{D}^m] = \sum_{h \in \mathcal{H}_B} \prod_{i=1}^m \mathcal{D}$$

$$(\{x_i \in X \mid h(x_i) = f(x_i)\}) \textcircled{\leq}$$

Але за означенням $\mathcal{H}_B$ кожна з цих ймовірностей не перевищує $1 - \varepsilon$:

$$\textcircled{\leq} \sum_{h \in \mathcal{H}_B} \prod_{i=1}^m (1 - \varepsilon) = |\mathcal{H}_B| (1 - \varepsilon)^m \leq |\mathcal{H}| (1 - \varepsilon)^m \textcircled{<}$$

$$\text{Оск. для } \varepsilon \in (0, 1) \quad e^{-\varepsilon} = 1 - \varepsilon + \underbrace{\frac{1}{2}\varepsilon^2 - \frac{1}{6}\varepsilon^3 + \dots}_{>0} > 1 - \varepsilon,$$

$$\textcircled{<} |K| e^{-\varepsilon m} \textcircled{\leq}$$

Керай тепер m достатньо велике (тобто навч. дані S містять дост. велику кількість точок), а саме $m \gg \frac{1}{\varepsilon} \ln \frac{|K|}{\delta}$ для деякого $\delta \in (0, 1)$. Тоді

$$\textcircled{\leq} |K| e^{-\ln \frac{|K|}{\delta}} = \delta,$$

тобто ще раз:

$$\mathbb{P}^m \{ S \mid \exists h_S \in \text{ERM}_K(S) : L_{\mathcal{D}, f}(h_S) > \varepsilon \} < \delta -$$

імовірність того, що за випадково обраною навчальною вибіркою S алгоритм з ERM_K видасть поганий результат h_S , менша за δ , аніж зі імовірністю $\geq (1 - \delta)$ (що зв'язано параметром впевненості (довіри, confidence parameter) цей результат буде добрим, тобто $L_{\mathcal{D}, f}(h_S) \leq \varepsilon$. При цьому оцінка на m не залежить від $\mathcal{D}, f$, лише від K, ε та δ (та від припущення про реалізованість) П.ч., ми маємо загальний результат:

Сон. Нехай $\mathcal{H}$ - скінченний клас гіпотез у задачі ML з довільними $X \text{ і } Y = \{0, 1\}$, $\epsilon, \delta \in (0, 1)$ і $m \in \mathbb{N}$ такє, що $m \geq \frac{1}{\epsilon} \ln \frac{|\mathcal{H}|}{\delta}$. Тоді для будь-якої розподілу імовірностей $\mathcal{D}$ на X і $f: X \rightarrow Y$ такє, що $\mathcal{H}$ задовольняє припущенню реалізованості ($L_{\mathcal{D}, f}(h^*) = 0$ для деякої $h^* \in \mathcal{H}$) зі імовірністю не меншою за $1 - \delta$ алгоритм $ERM_{\mathcal{H}}$ за видіркого S з m ек-мів X , що $\in$ і.і.д. за $\mathcal{D}$, поверне гіпотезу h_S , для якої $L_{\mathcal{D}, f}(h_S) \leq \epsilon$:

$$\mathcal{D}^m(\{S \mid \forall h_S \in ERM_{\mathcal{H}}(S) L_{\mathcal{D}, f}(h_S) \leq \epsilon\}) \geq 1 - \delta.$$

PAC навчання

Попередній розділ мотивує наступне загальне означення:

Def. Для задачі ML з довільною областю визначення X і дискретною міткою $Y = \{0, 1\}$ клас гіпотез $\mathcal{H} \subset \{h: X \rightarrow Y\}$ зветься імовірно придатно коректно навчанам (probably approximately correctly learnable, PAC

learnable), якщо існує функція $m_K: (0,1)^2 \rightarrow \mathbb{N}$ і алгоритм ML, що для будь-яких $\epsilon, \delta \in (0,1)$, розподілу ймовірностей $\mathcal{D}$ на $\mathcal{X}$ та функції класифікації $f: \mathcal{X} \rightarrow \mathcal{Y}$ таких, що K задовольняє припущенню реалізованості ($L_{\mathcal{D},f}(h^*) = 0$ для деякої $h^* \in K$), за навчальною вибіркою з $m \geq m_K(\epsilon, \delta)$ елементів, що незалежно та ідентично розподілені відповідно до $\mathcal{D}$ та позначені за допомогою f , зі ймовірністю не меншою за параметр впевненості $1-\delta$ повертає гіпотезу $h_S \in K$, для якої правдива помилка $L_{\mathcal{D},f}(h_S)$ не більша за параметр точності ϵ : (approximately correct) (probably)

$$|\mathcal{D}|^m \{S \mid L_{\mathcal{D},f}(h_S) \leq \epsilon\} \geq 1-\delta.$$

Пит "ймовірно" відноситься до $\geq 1-\delta$, а "приблизно коректно" - до $\leq \epsilon$.
Тепер можна переписати попередній наслідок:

Сол. Будь-який скінченний клас гіпотез K для будь-якої задачі ML є PAC навчальним, де $m_K(\epsilon, \delta) \leq \left\lceil \frac{1}{\epsilon} \ln \frac{|K|}{\delta} \right\rceil$, а алгоритм - з ERM $_K$.

Впр. Показати, що клас $\mathcal{H}$ іпотез, що визначені
прямокутниками (див. ех. 1. вище) $\in$ PAC навчання
для задачі "класифікації ординат", тобто з $X =$
 $= [a, b] \times [c, d] \subset \mathbb{R}^2$, причому $m_{\mathcal{H}}(\epsilon, \delta) \leq \left\lceil \frac{4}{\epsilon} \ln \frac{4}{\delta} \right\rceil$.
Узагальнити це на довільну вимірність d .

Ром. Ф-ція $m_{\mathcal{H}}$ ще зветься складністю вибірки (sample complexity). Її часто визначають білом менше як найменше натуральне число, для якого виконана умова з def. PAC навчання, саме тому вище з'явився нерівності $\leq$ на окруженні $\Gamma \cdot 1$ (перехід до найближчого справа натурального числа).

Підсумки цієї моделі:

1 Обмеженість лише бінарною класифікацією. Було б добре розглядати задачі й для інших $\mathcal{Y}$, зокрема, дійсних скінченних або $\mathbb{R}$. Але тоді потрібно

Буде змінами її означення помилки.

2. Існуюча "правильного" класифікатора $f: X \rightarrow Y$:
в реальності мітка y не обов'язково однозначно визначена
точкою x , наприклад, два фрукти однакового кольору і
твердості можуть бути смачним і несмачним через
інші фактори. Але одн. помилки залежать і від f .

3. Припущення реалізованості: теж часто насправді
невірне: наприклад, у випадку прямокутників завжди
 $\in$ певна, але не міри 0 частка точок з міткою
1 за межами $\forall$ пр-ка і з 0 - всередині. Це ще
не означає, що цей клас інозем K - позитив.

Щоб придати ці підстави, потрібно змінити
модель у частині побудови S (заминувати вимогу
статистичності та і.і.д.) та формулу обчислення
справжньої помилки: