

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. Talwar. 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 статистики — орієнтація на алгоритми, відсутність наперед визначеної 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, розшифровка, переклад, 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.

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

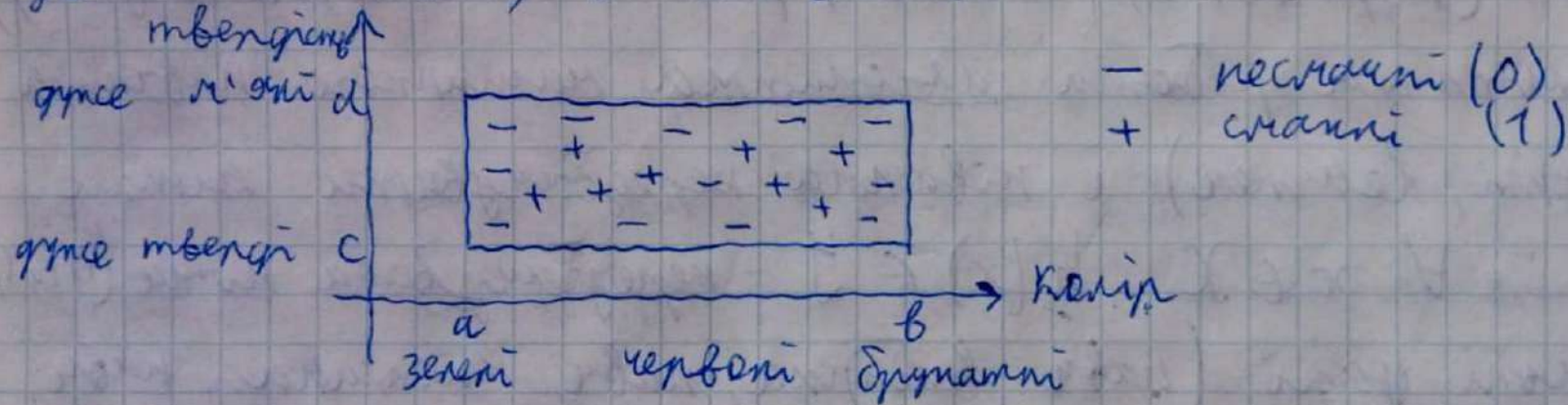
- Множина міток (область значень) Y (label set) - тут розглянемо задачу бінарної класифікації, тобто на 2 класи, тоді $Y = \{0, 1\}$ (часто беруть $Y = \{-1, +1\}$), де 0 і 1 означають належність $x \in X$ до одного з класів.

- 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 отримуює 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 некерованого

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 ?

Оскільки 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 $\underline{x}$ - сарау, де ростуть орхидея).
Ще означає, що $(X, \mathcal{F}, \mathcal{D})$ - ймовірнісний простір, тобто
- задана σ -алгебра $\mathcal{F}$ підмножин X (випірних) -
де яка 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. ($\leq$) злічених 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
(тобто такою, що містить усі відкр. замкнені інтервали,
усі їхні $\leq$ зліч. перетини та об'єднання). Далі ми,

Як правило, не будемо задованта $\mathcal{F}$ явно, 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 $\forall i \neq j$

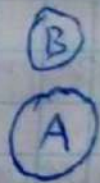
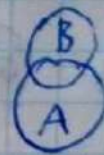
$$\mathcal{D}\left(\bigcup_{i=1}^{\infty} A_i\right) = \sum_{i=1}^{\infty} \mathcal{D}(A_i)$$

Тут цьому підмн. $A \in \mathcal{F}$ зовуть подіями, а $\mathcal{D}(A) \in [0, 1]$ —
ймовірністю події A . Вона виігрує, з якою ймовірністю
випадково обране $x \in X$ належить 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 підмн., де замість ∞ знову може бути n .



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

$\mathcal{D}(X|A) = 1 - \mathcal{D}(A)$ (чому?). Ек. Для $X \subset \mathbb{R}^d$ приймаючи ϵ орієнтований розподіл: $\mathcal{D}(A) := \frac{V(A)}{V(X)}$, де V - об'єм (міра Лебєга) на $\mathbb{R}^d$, та нормований. Отже, обернено $x_1, \dots, x_m \in X$ 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) функцію h . Тоді $\forall i$ визначимо $y_i := f(x_i)$ і таким чином побудуємо $S = \{(x_i, y_i = f(x_i))\}_{i=1}^m$ (Ек.: смачність одно-значно визначена керованим і мвергентно).

Тл.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.):

концеп з них обертається з X незалежно від інших
 згідно з $\mathcal{D}$. Тобто на X^m розкладаємо добуток
 імовірнісного мір $\mathcal{D}^m := \underbrace{\mathcal{D} \times \dots \times \mathcal{D}}_m$, ця 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 канючі,
 $S \in (X \times Y)^m$, 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 ?

Путем ми маємо "правильну" $f: X \rightarrow Y$, з якого потім
 по порівнянню h : чим вона "далі", тим гірше.

деф. Кошикової гіпотези h (справжньою помилкою

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

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

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. Помилкою навчання гіпотези h на навчальних даних $S = \{(x_i, y_i)\}_{i=1}^m$ (емпіричною помилкою, емпіричним ризиком) (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:

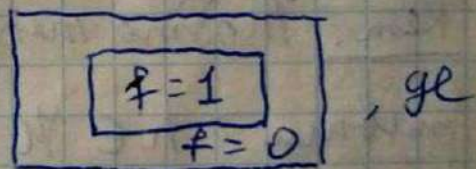
def. Творять, що алгоритм навч. А здійснює
мінімізацію емпіричного ризику (empirical
risk minimization, ERM), якщо за навчальними даними
 S він повертає гіпотезу h з найменшим значенням $L_S(h)$:
 $A(S) = h \in \arg \min_h L_S(h) := \{h: X \rightarrow Y \mid L_S(h) \leq L_S(h') \forall h': X \rightarrow Y\}$
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 $L_{D, f}(h)$. Ця традиційна
для ML проблема зветься перенавчанням (overfitting).

Ек. 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 D такий, що скінч. множини мають міру 0, то $L_{D,f}(h_S) = D(\{x \mid f(x) = 1\})$, то h_S приймає 0 майже усіх точках (міри 1). Це може бути дуже високе значення. Наприклад, у E_X з пунктами, однаковими D і f , що задається графом



площа внутр. пр-кутника у 2 рази менша за площу зовнішнього, $L_{D,f}(h_S) = \frac{1}{2} \quad \forall S$, тобто h_S "нічого не знає" про 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 підмножину H у множині 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 $h \in \mathcal{H}$:

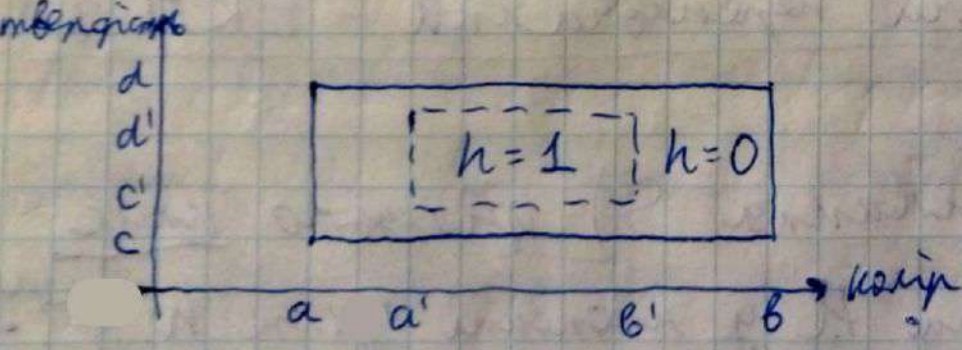
$$A(S) = \underset{h \in \mathcal{H}}{\text{argmin}} L_S(h) := \{h \in \mathcal{H} \mid L_S(h) \leq L_S(h') \forall h' \in \mathcal{H}\}$$

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

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

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

Ех.1. В задачі про смачність фруктів природно оцінювати,



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

визначені прямокутниками злі 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 гігантський "гарний" клас (ув. нарис).

Ел.2. Бачимо, не усі класи є користними. Нехай $\mathcal{K}$ - клас м.зв. порогових (threshold) 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 відповізна йому h_S - та h , що y (*).

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

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 типів дійсних чисел у нашій мові програмування) etc. Крім, $|\mathcal{H}|$ може бути дуже великою.

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

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

Rem. Це $\bar{\pi}$ 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 копієвої" $x \in X$: це так зі $\bar{\pi}$ ймовірністю 1. Зокрема, для "найменш копієї" вибірки S , тобто зі $\bar{\pi}$ 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 $\exists \text{ERM}_\mathcal{H}$ на цій вибірці, $L_S(h_S) = 0$
 (де необхід'язково $h_S = h^*$: банально, що найкраща помилка = 0).
 Але для нас банальна справсна помилка $L_{\mathcal{D}, f}(h_S)$. Кожній
 $\varepsilon \in (0, 1)$. Яка "кількість" вибірок $S \in \mathcal{S}$ потрібно, щоб
 $L_{\mathcal{D}, f}(h_S) > \varepsilon$? (ε тут зветься параметром точності (accuracy))
 Також є зв'язки (наприклад, може просто m раз повторитися
 одна й та сама пара (x, y)). Відкинувши підмножину S
 розміру 0, можемо вважати, що для такої S $L_S(h_S) = 0$,
 але h_S погана, тобто належить до підмножини

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

Пл.ч. (крім підмножини розміру 0, банально)

$$\{S \mid \exists h_S \in \text{ERM}_\mathcal{H}(S) : L_{\mathcal{D}, f}(h_S) > \varepsilon\} \subset \{S \mid \exists h \in \mathcal{H}_B : L_S(h) > 0\}.$$

Почу в силі монотонності

$$\mathcal{D}^m(\{S \mid \exists h_S \in \text{ERM}_\mathcal{H}(S) : L_{\mathcal{D}, f}(h_S) > \varepsilon\}) \leq \mathcal{D}^m(\{S \mid \exists h \in \mathcal{H}_B : L_S(h) > 0\})$$

Представимо цю групу множин як об'єднання

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

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

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

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 $1 - \varepsilon$:

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

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

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

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 зветься 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 вибіркою S з m ел-тів 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 мішкою $Y = \{0, 1\}$ 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 навчального вибіркою S 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:

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

Вип. Показати, що клас VC імпонет, що визначені
прямокутниками (див. Ex. 1. вище) $\in \text{PAC}$ навчання
для задачі "класифікації орієнтов", 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 білим числом як найменше натуральне число, для якого виконана умова з def. PAC навчальності, саме тому вище з'явився нерівності $\leq$ та округлення $\lceil \cdot \rceil$ (перехід до найближчого справа натурального числа).

Розглянемо цілі моделі:

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

Буде змінами і означення помилки.

2. Існують "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 смаканими і несмаканими 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 (заминюючи вимогу
статистичності та і.і.д.) та формулу обчислення
справжньої помилки: