

- skúškové otázky: najbližšie dnu
- ANKETA

Ak projekt vyriešite skôr, napíšte mi, prosím mail

Minule: Testy

perm. test: H_0 : rozdelenie X = rozdelenie Y (vzhľadom na štat. š.)
 H_1 : *

napr. $\bar{S}_* = \bar{x} - \bar{y}$

X : muži, Y : ženy

• $x_1, \dots, x_m, y_1, \dots, y_n$

presporiadame kvačdy $H/Z \Rightarrow S = \dots$

$\rightarrow$ * presporiadanie

- H_0 tam, ať $\bar{S}_*$ je v 5% najextrémnejších

Pr.: 3 muži, 2 ženy

H_1 : muži sú väčší ako ženy

valky $\hookrightarrow 87,95, 71 \hookrightarrow 63,79$

$\bar{S}_* = \bar{x} - \bar{y} = 13,7$

$\frac{87,95}{m_1} \frac{71}{m_2} \frac{63,79}{m_3} \frac{71}{z_1} \frac{79}{z_2} \rightarrow \bar{S} = \dots$

$m_1, m_2, z_1, z_2, m_3 \rightarrow \bar{S} = \dots$

⋮

↓
 $S!$ 5-iek



$p\text{-val} = \% \text{ extrémnejších } S$

- časťnosť: iba K náhodných permutácií: aproximácia

Big Data

- také, ať. mriem spracovať na 1 počítači
 $\rightarrow$ Twitter, FB, Google...
- rýchlosť - statické: fixné
 - "historické": pomaly pribúdajú - denné výnosy
 - streamovanie: real-time pribúdajú
- nemáme vždy kapacitu uložiť
 $\rightarrow$ súhrnne kamier, Twitter...
- kvalita - rozkierať za účelom (chceme zistiť účinnosť lieku)

Big data

- výsledok posledného zberu: všetko, čo máme } Big data
sa dostanem

→ nemusia byť reprezentatívne...

Hlavný problém: nemôžeme použiť klasické postupy (parametrické)

Riešenia: I. vyhnúť sa práci s veľkými dátami

→ 1.) Vzorovanie (sampling, subsampling) - iba časť dát

→ 2.) Sumarizácia (sufficient statistics) - pracovať iba s tým, čo potrebujeme

II. Špec. postupy pre Big data

1.) Vzorovanie:

- modely typicky majú rozumný # parametrov → môže stačiť menej dát

a) Náhodný výber: reproducibilita?

: môže byť ťažké získať viac náh. vzoriek, kt. sa nepretýkajú
(train, val, test)

Nenáhodné!

b) Príjich k záznamom → môže byť veľmi nerepresentatívne

c) "Každé 1-té": 1, 1001, 2001, 3001, ...

: take $i \equiv 1 \pmod{1000}$

alebo $\equiv 34 \dots$

: takto nepretýkajúce sa výbery: $i \equiv 1 \pmod{1000} \rightarrow 1,$
 $i \equiv 2 \pmod{1000} \rightarrow 2,$

: môže byť ok, ak dáta majú špec. štruktúru:

napr. zdarný - chový - zdarný - chový ...

Všed.: pozor na reprezentatívnosť

- nedáme štatisticky relevantných outlierov

→ "malé firmy" (napr. nezbankrotované firmy)

→ stratifikované vzorkovanie:

- najprv rozdelenie na mužov a ženy → potom z každej skupiny: výber

→ váhované vzorkovanie: rôzne pr. rôznym záznamom

(napr. vysoke' pp. pre objekty zo vzajnej triedy)

2.) Sumarizacia

- uložiť si len to, čo nutne potrebujeme pre model
 - sufficient statistics (postacujúce štat.)
- napr. kontingencná tabuľka namiesto veľá záznamov

II. Keď to inak nejde : špecializované postupy

- relačné databázy (SQL)
 - "scalelessly" paralelizovateľné výpočty
 - ak sa dá rozložiť na menšie výpočty, kt. sú spustiteľné samostatne
 - každý výpočet na 1 počítači / jadre...
 - paralelné výpočty
 - každá jednotka svoje, ale musí sa to priebežne spájať / aktualizovať
 - treba minimalizovať množstvo dát, kt. si počítače medzi sebou prešlávajú
- 3 "jazyky" : frameworky - MapReduce, Hadoop
- treba používať overené nástroje, neimplementovať vlastné riešenia
- pri veľá dátach sa nič nepočíta, frameworky to majú ošitné