

Vyhodnocování výkonosti počítačových systémů ~~WTS~~ VUPS

pozn. studijní povinnosti - na konci zkušebního (přesného)
neexistuje literatura ani skript! přednáška je kompilát.

<http://d3s.mff.cuni.cz/teaching/peva>

- slajdy v pdf, login performance: pe13slides
mailing list

př: ping benchmark

- server/client na jednom počítači
- RPC přes CORBA
- opakované požadavky a měření odezvy

- CORBA - na řadu věcí řeší problém, ale funguje a používá se

avšak různé implementace mutexu (mutex, old mutex)
- zamýšlení a odemykání

- některé služby se chytí ať už při opravě nějaké množství podmínek - dočasné řešení systému
- výsledky často těžko interpretovat jen pomocí průměru
- někdy se výsledky liší mezi dvěma spuštěními (úhly) - postupná stabilizace
- sledování i při změně v systému
- příčinou může být např. memory layout → využívání cache apod.

- cíl vzhledem k výkonosti

- system s maximálním výkonem
- system s požadovanou výkoností za minimální cenou
 - otrávená šifrovatelností (skří) jen u.c. řešení?
- realtime system - reakce v pevně daném čase (ne můžeme max výkon)
- system, kde chceme snižovat analyzovat worst execution time

hlubší implementace a testování

- modelování - matematický / statistický model pomocí formalismů
 - abstrakce známých charakteristik systému a prostředí
 - systémový model abstrakce, petriho síť, ...
- simulace
 - analýza je často nemožná, než se řeší numericky

hlubší implementace a testování

- měření (co, jak, proč) - benchmarking
 - workload relevantní pro daný systém
- monitoring - stará data
- profiling - kde hledat v kódu a kódu nejvíce času?

- výhody modelování
 - odhalení problémů během design time → úspora
 - např. součerení o prostředky (centrální databáze)
 - opakovatelnost
 - analýza extrémních / krajních příkladů - přitáhnutí tisíců uživatelů najednou, apd.
 - hledání vzájemných souvislostí různých elementů - zjištění, na co se zaměřit
 - relativně nejrychlejší a nelevnější

- nevýhody modelování
 - snadno se vytvoří příliš komplexní model, který nelze analyzovat s dostupným výpočtovým výkonem - je nutné hodně abstrahovat
 - příliš mnoho abstrakce může vést k zanedbání důležitých faktorů
 - nižší důvěryhodnost výsledků modelování (hlavně numerických)
 - spíše se soustředí o zachycení vztahů mezi faktory

- simulace
 - doplnění dohadů neimplementovaných částí systému

- výhody
 - snížení nákladů
 - více detailů než modelování
 - (slabší)

- nevýhody
 - problém udržet délkou simulace v rozumných mezích
 - zjednodušení → výsledky nemusí být reprezentativní
 - pokud je realističtější, nemusí být opakovatelná
 - většinou pomalejší a dražší než modelování
 - náchylnost na (nezřetelné) chyby (např. undog serverů mih. 200)

- kombinace modelování a simulace
 - potvrzení (odhalení chyb v analýze)
 - kompletnost - modelování neodhalí vliv hw implementace
 - simulace nemusí odhalit chování v extrémních případech

Měření výkonu

- metiky - použitelné i na simulaci
- cíle
 - porovnání alternativ
 - vliv nové firmy
 - tuning systému (např. pro konkrétní hw)
 - měření relativního výkonu (metri vs. reálné apl. zruš)
 - metod. performance debugging
 - odhad chování v nejhorších případech
- porovnávání alternativ
 - nejspektativnější měření podle reálných apl. zruš
 - nelze ale #aplikace x #HW kombinace
 - měření výkonosti jednotlivých komponent
 - proveditelnější, ale méně spolehlivé
 - jednoduchý model → středně přesné výsledky
 - ps. Open (veřej. Benchmarking project)
 - platform benchmark - základní operace
 - odhad výkonosti na dosud neexistujících platformách - snaha o odhadnutí závislosti komponent
- vliv nové firmy
 - dvě verze systému odlišné jen v jedné věci ("před a po")
 - nastavení kompilátoru apod. - také může mít velký vliv
- ladění systému
 - hledání optimálního nastavení systému pro konkrétní prostředí
 - nastavení ds, vm, ... velikosti bufferů, cache, nastavení plánovače, počet vláken
 - možnosti se mohou vzájemně ovlivňovat, někdy velmi neintuitivně
 - typicky nelze testovat všechny kombinace (mnohdy hodně parametrů...)
 - experiment design theory - odhady vzájemných závislostí komponent, ...
 - ps. Shell Project

- Shall
 - počítačové testy, experiment design theory
 - odhad závislosti možnosti na závislosti experimentu
 - na závislosti toho identifikovat nastavení kritické pro výkonost
 - validovat a aktualizovat úsudky dalšími experimenty
 - ostatní nastavení - randomizovat, zkusit různé další úlohy
 - prohledávání prostoru nastavení (lokalně) - viz optimalizační metody

- měření relativního výkonu
 - změna mezi dvěma verzemi systému

- př. regression testing - zda nová verze dělá to samé (co se výkon týče)
 - je dobré dělat to pravidelně, (k automatickosti (ale dost často))
 - experiment se objeví poměrně brzy po vzniku
 - př. projekt Bttn

- Indikátor výkonu

- vyhledávání hotových programů
- identifikace pomalých nastavení
- menší úskalí, odhalují dříve problém
- profilování - které části jsou nejpomalejší? → prioritizace
- př. cíl by neměl být zjevný
 - experiment → zachytit a hledat menších změn výkonu
 - hledat korelace s událostmi v systému
 - cache misses, page faults, GE, context switches,
- př. performance explorer

- profilování
- trace alignment - během určitého času měřit různé události (jako je to v případě)
- trace visualization

- odhad výkonových případů (důležité pro realtime systémy)

- certifikace - "manuální" prohledání kódu a odhad maximální doby provádění - může poskytnout znalost všech podmínek, hw, ...
- často nutno napsat program, aby byl testovatelný
- optimalizace pro určitý case - pokud změna tento neoptimalizuje, změnu neprovádět - jinak se nevíde až dost často na to, aby se to dal měřit
- př. Myronome Real-time GC

- měření výkonu - jak rychle jednáme?

- typický - počet událostí
 - dosažení / meti událostmi
 - podle výsledku

- systém dostane požadavek na službu a může

→ splnit požadavek korektně
→ splnit požadavek nekorektně → nekvalitní spolehlivost: pro nějaký druh chyby
→ odvrátit splnění - psst chyby (dle velikosti, např. požadavek, pakety)
čas mezi chybami

↓ jak dlouho trvá stav, kdy se tohle děje?

↓ čas mezi takovými událostmi (někdy se hodně zkrátí systém
resetovat, pokud neodpoví do určité doby)

- analýza spolehlivosti

- hw a sw, dost čisté

- regrese: testy neodhalí vše, chyby jsou často zcela nahodné

- měření může odhalit chyby

- např. „průběh“ výkonu →

př. Mono Regression Benchmarking

- neaktivní!

chyba, obsluha, významy a množství
normálního provozu

- občas může systém padat, to by nikdy nemělo nastat v testování

- testy spolehlivosti jsou náročné ⇒ dají hlášet poruchy, pokud open source

- nemůžeme nahradit regresní testování, prověřování je
podstatně složitější (i: nároky na stabilní prostředí)

- analýza výkonu

- čas odezvy (ten mezi přijetím požadavku a odpovědí)
an - měří „humanatelné“ výsledky

- produktivita

- pro produktivitu systému - jaké jednotky dat / čas se podařilo zpracovat
- může se odhadnout, rozumíme, může hodnotit

- využít zdroje - užití

- měří, kolik % času je prostředek užitý nad určitou mez

→ odhalování úzkých míst a systémů nebo předimenzovaných
komponent

Metriky

AM

- měření výkonu - v jakých jednotkách?

- typicky - počet událostí
 - dosažení / mezi událostmi
 - podle výsledku

- systém dostane požadavek na službu a může

- ✓ splnit požadavek korektně
 - ✓ splnit požadavek nekorektně → metrika spolehlivosti pro nějaký druh chyb
 - ✓ odmítnout splnění
- počet chyb (dle velikosti, např. počet rozestupů, paketů)

čas mezi chybami

↓ jak dlouho trvá stav, kdy se takhle děje?

čas mezi takovými událostmi (někdy se hodně zhorší systém, pokud neodpoví o něco dříve)

- analýza spolehlivosti

- hw a sw, dost čisté

- regrese: testy neodhalí vše, chyby jsou často zcela nahadné

- měření může odhalit chyby

- např. „průběh“ výkonu → ~~průběh~~

chyba, obsluha, významy a množství normálního provozu

př. Mono Regression Benchmarking

- neaktivní!

- občas může systém padat, to by nikdy nemělo nastat v testování

- testy spolehlivosti jsou náročné ⇒ dají hroší patiči, pokud open source

- nemůžou nahradit regresní testování, prověřování je podstatně složitější (i: nároky na stabilní prostředí)

- analýza výkonu

- čas odezvy (ten mezi přijetím požadavku a odpovědí)

an - měří „humanatelné“ výsledky

- produktivita

- pro produktivitu systému - jaké jednotky dat / čas se podařilo zpracovat

- může se odhadnout, rozumíme mezi hodnot

- využití zdrojů - utilization

- měří, kolik % času je prostředek vytížený nad určitou mez

→ odhalování úzkých hrdel a systémů nebo předimenzovaných komponent

- konzistence - jednotky měřící by měly mít stejný význam na různých systémech a jejich konfiguracích
- př. MIPS, MFLOPS nesplňují (RISC/CISC grad.)
↓
watstone, dnystone benchmarky

- nezávislost - neměla by být optimalizovaná na konkrétní prostředí ani naopak prostředí na benchmark (což výrobci někdy dělají) → SPEC
→ zhruba výpočetní hodnoty

- př. metricky pro výkon procesoru a systému
 - frekvence → ne spolehlivá, nelineární (různé architektury, různé podmínky, ...)
 - MIPS, MFLOPS - o něco lepší než MIPS, FLOPS je dobře definovaný, vhodný jen pro systémy, kde na FLOPS záleží
 - SPEC - sada standardizovaných benchmarků → vyšší výpočetní hodnoty
 - QWIPS
 - doby běhu

- SPEC

- sada benchmarků, z jejich výsledků geom. průměr (aritm. průměrem dostanete extrém, medián je tak více robustní)
ovšem - není lineární vůči době běhu programu
- není intuitivní

- spolehlivost je dobrá pro jednotlivé benchmarky, ale méně využitelná pro porovnání

- o něco lepší pro application benchmarks - př. SPECint2000

- nezávislost: není ideální, řada způsobů optimalizace pro SPEC, aniž by se tím výrazně zrychlily reálné aplikace

- zlepšuje se na vyšších úrovních a př. vyšší počty benchmarků

- base a peak profil (u base jen základní příprava)

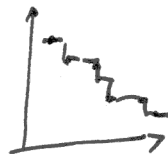
- občas z benchmarků obměňují

- využití na základě sady benchmarků není nutně podstatné, jen je nutno se vyhnout zjevně nevhodným optimalizacím

- Virtualizční benchmarky
 - měření schopnosti HW + hypervisoru spouštět zároveň VM s různým využitím
 - standardní ~~na~~ průmyslové benchmarky nejsou lineární a nepoužívají dobré metricky
- VMmark
 - sada benchmarků nad VM a set virtualizace („reference“)
 - postupně přidáváme VM
 - měří se propustnost
- TMT
 - reference a VM
 - měří se $\left\{ \begin{array}{l} \text{celková propustnost} \\ \text{průměrná response time} \end{array} \right.$

- Pareto curves

- ~~break~~ graf odezvy / throughput
- pareto-optimalní servery
 - je jich už možná,
 - záleží na ~~term~~ prioritách



- Execution time

- intuitivní, jednoduchá
- pozor na background load
- čas CPU?
 - neprovidají o celkovém (hovrn)
 - někdy se měří interrupt based accounting
 - ignoruje nečistě číselnou
 - take' out of the background load (TLB, cache, ...)

- Kompromis: CPU & wall clock time

- splňuje většinu požadavků
 - problém s jednoduchostí měření - přesnost
 - opakovatelnost - pouze statisticky
 - už podléhají předpoklady

- QUIPS - Quality improvement per second
 - matematický problém řešený iterativní metodou
 - kvalita = jak blízko je současný stav k výsledku (v každém kroku se zlepšuje)

- HINT benchmark: $\int_0^1 \frac{(1-x)}{(1+x)} dx$

- postupný přesňování integrálu

- kvalita = $\frac{1}{\text{horní met} - \text{dolní met}}$

- vyhovující hodnota je omezená - specifický numerický problém

- jak vytvořit analogický benchmark pro I/O, apod.?

- někdy se snadno rozlišet nežádoucími vlivy

- měří se výkon, ne činnost - doplňuje FLOPS

- práce - při velkých zátěžích

- vytvořit si benchmark na míru

- zveřejnit výsledkem, nechat je optimalizovat systém na benchmarku

- smluvně požadovat opakovatelnost

metry $\left\{ \begin{array}{l} \text{úsili' na dosažení cíle (MIPS, FLOPS, ...)} \\ \text{cíl (QUIPS apod.)} \end{array} \right.$ - inherentně nespojitlivé

- nedeterminismus v době těch
- warmup - několik prvních sémí typicky trvá (i u známé) déle
⇒ vynechat ze statistik

- předpoklady (které ale někdy nuplatí)
 - doby běhu jsou náhodné (skryté závislosti)
 - nezávislé
 - rovnoměrně distribuované

⇒ náhoda je nepochopitelná souvislost

→ běžné statistické metody

- často požadavek na jedno číslo

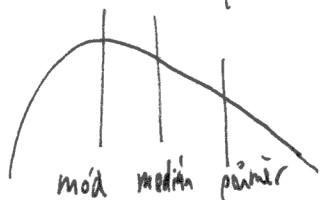
- u jednoduchých benchmarků se tato někdy sblíží dají

- různé statistické metody

- distribuce nemusí být unimodální, může být multimodální (uíc / ot. maxim), ...
- měla by mít intuitivní vzhled, něco opravdu znamenat
- medián, průměr, ...

- doby běhu → náhodné proměnné $X_1, \dots$ - normálně jejich distribuce
- $x_1, x_2, \dots$ - realizace n-h. prom. (naměřené hodnoty)

- přehledovité: distribuce, frekvence, pstr, dist. a spoj. NV, kvantily, střádlosty



→ co vybrat?

kategorická data → móda

počet dat daná (tzv.) smysl? →

vykložené distribuce → aritmetický průměr nereprezentativní, → medián

jindy → průměr

- problémy metrik

- průměr - velmi citlivý na extrémní výsledky

→ ořezaný průměr (kromě několika % krajních výsledků) - spolehlivější

- medián - nepohodlně ignoruje hodnoty informace - extrémní informace

- vícemodální data se zhruba stejně velkým: cluster ⇒ fragilní

- pozor na korelace → nelze násobit průměry

- jiné průměry - nutno pečlivě vybrat

- agregativní metricky

- linearity, spolehlivost!

- geometrický průměr: $\sqrt[n]{\prod_{i=1}^n x_i}$

- použít pokud má smysl součin hodnot (mnohočetná)

- vícemodální data → cluster

→ průměry zrychlení ve vícečetném systému

- nepoužívat na bare → průměry (stále je o křivce, aby geom. průměr dával a jmenovatel měl význam!)

- harmonický průměr
 - "hybná" usnadní: data obsluhy požadována - pozor na obecné trendy
 - ...
- poměry
 - podíl měřené hodnot dané syst $\Rightarrow$ podíl sum
 - konstanta v čitateli $\rightarrow$ harmonický průměr
 - "skoro konstantní" poměry $\rightarrow$ geom. průměr
 - ne při různých slovních, apod.
 - vhodné při srovnání o odhadnutí nějaké konstanty: $a_i = C \cdot b_i$
- př. ~~průměr~~ průměrné změněné velikosti zdroj. údaje (před optimalizací a po ní)
 - program / kB před / kB po / poměr $\frac{po}{před}$
 - součet čísel i součet jmenovatelů dané syst, násobí se posčítat velikosti ~~každá~~ a spočítat jejich podíl = podíl průměrů
 - $\rightarrow$ průměrný se programy různě velikosti, každý kB má různý vliv (dominují velké programy)
 - daný vliv programů syst? optimalizace máli korekci s velikostí programu, což x bezohledně
 - ideální by bylo provázet velikosti programů s frekvencí jejich vstupu - neznáme
 - raději průměrnat přes programy
 - geom. průměr - výsledky v blízkém okolí
 - $\rightarrow$ vhodnější
 - stále bezohledně velikost programu $\sim$ změně
- př. average speeds:
 - benchmark / ref. time / measured time / speedup
 - součty a průměry času opět dávají smysl
 - benchmarky mají ale velmi odlišné hodnoty
 - geom. průměr nelze použít - velký rozskok hodnot
 - model $a_i = C \cdot b_i$ resení - velká řada různých funkcí
 - SPCC je použitelný, ale v principu to není dobré
 - $\rightarrow$ ideálně přidat velký benchmarkům
 - jinak, můžeme-li přiblížit, že naměřené hodnoty jsou reprezentativní, použít harmonický průměr zrychlen
 - není dobré průměrnat podíly odlišná čísla

- vážení průměrů
 - součet vah vždy = 1
- m44 může být ve zhoršivé roli

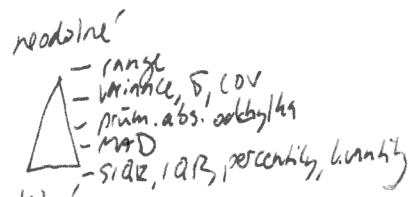
- variabilita u nedeterministických měření
 - také (spolu s průměrem) vypovídá o kvalitě systému
 - histogram - potřeba vhodné zvolit jemnost (z důl. histogramu)
 - pozor na warm-up - warmout

- se statistikou nelze pracovat moc dobře algoritmicky
- každé rozhodnutí velmi přesně zdůvodnit

- metricky variability dat
 - průměrná absolutní odchylka
 - mediánová absolutní odchylka (MAD)
 - mezi kvantilový interval (IQR)
 - semi-interquartile range
 - 10-90. percentil (očerání extrémů)
 - variance, směrodatná odchylka, COV

- výše metricky závisí na

- pořadovost robustnosti máli okrajovým hodnotám:
 - ↓
 - víc není vždy lepší, outliers mohou mít velký vliv
- použít metricky benchmark
- konkrétních dat



- rozsah
 - polovina proměnná má smysluplné hranice
 - polovina měříme hard - realtime systém

- rozptyl

- outliers vzácné a nepříliš velké
- polovina je distribuce spíše unimodální symetrická
- polovina je jaká metrika ante. tendence použít průměr
- upřednostňování COV

- průměrná absolutní hodnota

- std. odchylka není příliš robustní máli okrajovým hodnotám

- robustní metricky
 - pokud se ~~po~~ měří medián,
 - pokud je distribuce multimodální, má hodně odvržených hodnot, good.
 - kvantily a IQE jsou dobré, kvantily mají smysl pro soft-realtime systems
- obecně při vyhodnocení závisí na aplikaci a cíli měření
 - také na dalším vyhodnocování, s nevhodnými metrikami se špatně pracuje
 - např. chceme odhadnout distribuci

- Jak interpretovat a reportovat variabilitu ve výsledcích?
 - systém má vlastnosti, které neznáme, ať můžeme být deterministické
 - měřicí nástroje, prostředky, ve kterých se měří → chyby, šum
 - termíny
 - v češtině se termíny značímají potv.
 - accuracy - jak se měření blíží reálné hodnotě - "správnost"
 - precision - množství šumu v naměřených hodnot (jejich stabilita) - "věrnost"
 - resolution - nejmenší hodnota, kterou jsme schopni změřit - "rozlišení"
- chyby v měření
 - systematické
 - ovlivňují správnost
 - měly by být eliminovány
 - např. overhead time, grad.
 - velmi opatrně vše nastavit, kontrolovat
 - minimalizace rušení (pro služby uživatele, ...)
 - stabilní prostředí (SW/HW update)
 - randomizace naměřených prostředků
 - převedení na naměřené chyby
 - náhodné
 - ovlivňují přesnost
 - nedeterministické, nepředvídatelné - často nelze kontrolovat
 - mají stejnou pst. zhoršení / zlepšení měření!
 - reportovat vše dohromady, většinou ruše izolovat
- model náh. chyb
 - předpoklady v praxi mnohdy neplatí (nebo zdaleka ne všechny), přesto jsou velmi užitečné
 - normalní rozdělení
 - střední hodnota - μ : přidání a náh. systematická chyba
 - náhodné chyby ovlivňují rozptyl σ^2
 - vždy reportovat alespoň naměřenou hodnotu $\pm$ odchylku
 - nezmiňovat nic zjevně špatného nebo nepotřebného
 - NE uvádět čísla ~~max~~ chyby maximální
 - "sigma rules"
 - $\sim 68\%$ hodnot v hranicích chyby měření od stř. h.
 - $\sim 95\%$ hodnot $\in [\mu - 2\sigma, \mu + 2\sigma]$
 - $\sim 99,7\%$ hodnot $\in [\mu - 3\sigma, \mu + 3\sigma]$

- jak jsou naše odhady spolehlivé?
- * intervaly spolehlivosti: oprávněná střední hodnota je ve zvoleném intervalu se zvolenou pdf
- nepřímá úměra mezi přesností odhadu a spolehlivostí
- o numerických ~~průměrech~~ průměru je nutno uvážet jeho o NV

$$E(X_i) = \mu \quad \text{Xi i.i.d.} \rightarrow E(\bar{X}_i) = \mu$$

$$var(X_i) = \sigma^2 \rightarrow var(\bar{X}_i) = \frac{\sigma^2}{n}$$

$$sd(X_i) = \sigma \rightarrow sd(\bar{X}_i) = \frac{\sigma}{\sqrt{n}}$$

(standard deviation)

standardní odchylka
užitelná průměr
- "standardní chyba"
- lze uvést do výsledku
(spolu s n)

- interval spolehlivosti je počítán vždy pro ~~některé~~ konkrétní měření
- $pdf(\alpha)$... pravděpodobnost, že průměr neleží v intervalu, je α
- pro jiné měření jiná spolehlivost!

... viz slides...

Porovnání alternativ

- dva benchmarky, jejichž výsledky se „přibližují“
- nepárování
 - dvě sady dat z různých systémů, ale různé vlastnosti
 - chceme porovnat různé průměry
 - předp: variabilita = nřh. chyby
 - jiné metody: confidence interval pro rozdíl, dva CI, ...
- pokud se CI obou měření ~~nepřekrývají~~ nepřekrývají, lze s „obvyklou“ úrovní spolehlivosti předpokládat, že rozdíl mezi průměry ~~je~~ systémů
- přibližující CI, průměry nejsou v CI Ankeho měření
 - statistický test, pokud nelze, neuskotat chodit
- přibližující CI, skřh. v CI Ankeho měření
 - nelze dokázat, že nejsou stejné
- přibližující interval spolehlivosti
 - směrnost metody - implementace, vnitřní, ale intuitivita
 - nezávislost na distribuci
 - není zřejmá spolehlivost celkového výsledku (neodvážlivé, někdy nelze)
 - 95% přísl. konstant (takto nedostatečně podstatné rozdíly), stl chybí interval 68% přísl. konstant
- kvalita porovnání
 - porovnání hodnoty nezáleží - jsou nebo nejsou stejné, nelze mluvit o psti
 - rozhodnutí na nřměřených skřhodnotě
 - rozdíl > threshold => ...
- nulová hyp. H₀ (H₀) - default „pokřdli hodnoty sledované veličiny“
 - H₀: průměry měření jsou stejné
 - dostatečně silná důvěra, že rozděl v břecech nejsou
 - pokud na rozděl > odhodnotit H₀ > průměry jsou různé
- dva typy chyb
 - typ 1: false positive (zaměnit pravdivé H₀)
 - typ 2: false negative (přijít nepravdivé H₀)
- míra ~~spolehlivosti~~ významnosti
 - vícje chyb 1. typu, tradeoff s chybou 2. typu
- síla testu je čísla významné
 - závisí na míře významnosti, vřl.losti vzorku a
 - rozměru rozděl průměry

α ... míra ~~spolehlivosti~~ významnosti
 β ... síla testu

	H ₀ = true	H ₀ = false
zaměnit H ₀	typ 1 P = α	true positive P = 1 - β
zaměnit H ₀	true negative P = 1 - α	typ 2 P = β

- jaké zvolit míru významnosti: testu?

- zkrátka, co by měl státní činovník a co říkám, když rozhodnu činovník?

- předpoklady

- normální distribuce náh. (110) chyb

- ~~procenty~~

- stá. odchylky se mohou v měřeních lišit

- jakékoliv náh. přesným CI

- pro ~~průměr~~ měření st. míry významnosti s použitím CI,

je potřeba vzít míru spolehlivosti pod 95%. Pro měření s podobnými
rozhledy Aho fungují hodnoty zhruba 84%.

- reportovaný výsledky

- odhadnutí H0 s 5% významnosti ...

- nelze na odhadnutí H0 na hladině 5% významnosti $\Rightarrow$ nerostí významnost

- říkám více dat $\rightarrow$ lze říci CI

- porovnání proporcí

- např. % čísel státních voleb

- power - měry spolehlivosti

- statistické testy

- testování statistické potřeby z Aho (např. vzhled příměsí)

- kritická oblast

- pokud se výsledky testu nacházejí zde, zamítáme H0

- zpráva na CI

- p-hodnota: významní hladina významnosti množství zamítání H0

- čím menší, tím více důvěry máme pro zamítání H0

- zpráva odpovídá, že dostaneme tuto/extremnější data, pokud H0 platí

- srovnání: průměry, st. hodnoty, odchylky, distribuce, ...

průměr: dva měření pro měření polohy (na stejných úsecích)

- o něco silnější

ne už dva měření (populace)

- hierarchické modely

- ANOVA

- porovnání porovnání

- převod na nepárový (rozdíly měření mezi konstanty)

- CI: konstanty, zda, 2. konstanty 0

- testy: předpoklady - zda prům. rozdíl je 0

- nepárový porovnání

- CI (pro $\neq$ metodu, pro kterou máme důstojnost (I))

- statistické testy

- T-Test (welchův t-test)

- k-test - začíná předpoklady o distribuci, parametricky

- kolmogorov - Smirnov

- neparametricky

- porovnání distribuce (dva vzorky a referenční)
 $\rightarrow$ lze už jako test normality (ale z jiných)

- ANOVA

- znalosť o rozličných významoch (výsledky) výslední měření
- odhad jejich stat. významnosti
- podrobnosti výsledky

- v praxi nie je to nikdy jednoduché a pred príkladom splniť → myslieť!
- i tak máme štatistické metódy porovnávať Anna Anta
- stat. metódy nemôžeme nahradiť analýzou systému
- naučiť sa štatistiky ktorú je možné si sám napísať

- veľkosť efektu

- ~~trivializácia~~

- Cohenovo D a ďalšie odhady
- vhodné tvoriť min. ktorú budeme ešte zanedbávať

- Exploratory data analysis

- Anscombeov čtvereček
 - vit slajdy, úplne stejne charakteristiky (priemer, medián, rozptyl, ...)
 - ale ve skutečnosti se hodnoty liší → dost věcí lze zjistit z grafu
 - vždy se dívat na data!

- analyzovať data pred použitím stat. metód
 - pre odhalenie štruktúry, outliers

- závislosť na grafickom zobrazení

- pr. SciMark2 (vit slajdy)
 - FFT v CH Mono

- Run-sequence plot

- index x value
- posun v hodnotách / variáci, šum, outliers, multimodalita
- interferencia s jiným procesem (podobnosť outlierů)

- histogram - odhad hustoty, predstavuje o distribúci, streda, rozptyl
 - spoľahlivejšie než run-sequence plot, ale nevhodný, časovo závislý
 - vytváranie množiny (použití multimodality model na odhad množiny)
 - nemôže byť se môže po ďalších metódach dobre zmeniť

- nezavislost meren: lag plot
- oty: x : namerene hodnoty
 y : namerene hodnoty bez prvich h L } nebo
 } nepat, nep. mist
- namernost dat, schvacheni korelace, outliers, odhad chybneho modelu pri
 casove zavislosti
- jedna hodnota $\Rightarrow$ na zaklade jedne hodnoty nelze rict, jak
 bude nasledujici
- data v clusteru $\Rightarrow$ ~~prakticky v nem mohou byt~~
 prechazeni mezi clustery?
- presamplovani (zmen priklad vysledni mereni) - pomoci ziskanou clusteru
 podobne, rejde o casovou zavislost
- zkusnout lag $\rightarrow$ lze vysledovat periodu (cluster se dopynaji)

- bootstrap plot
 - náhodné převzorkování (když užší 1x užší vícekrát)
 - pro každý subsample znovu vypočítat statistiku
 - > stabilita výsledků? tvar celkové distribuce?

t

- Quantile - quantile plot
 - pro porovnání distribucí dvou množin dat
 - $x = y$
 - stejné $\Rightarrow$ na úhlopříčce
 - porovnání s normální distribucí

boxplot

- violinplot

- histogram

- Prezentace dat

- Jain
- neřizovat, nemaniplikovat
- jednoduše, pochopitelně, popisně