

Distrib. syst. $\equiv$ systém propojení množiny nezáv. uzlů,
který poskytuje uživateli dojem jednotného
systému

ma něco jako
procesor, paměť
a připojení na
nějakou kom. linku

80. léta - distrib. protok. a alg. - teorie

90. léta - distrib. OS, systém Amoeba, T4

$\rightarrow$ neměl být rozdíl mezi:

lok. a distrib. aplikacemi

- ale FOUT - moc komplikované

0. léta - ^{lok. OS +} middleware - distrib. mezivrstva poskytující těm aplikacím,
které chtějí, to distrib. prostředí.

$\Rightarrow$ rozdíl mezi lok. a distr. apl.

201x - cloud computing, *aaS

Software
Infrastructure

- např. Gmail, napsy.cz... běž! nad tímhle

202x - trend decentralizace (web 3.0)

Cluster computing

XX

Grid computing

- na 1 fyz. místě je ^{HPC}
víc racků (např. na HS)

$\rightarrow$ Docker, Gitlab...

- HPC - většinou 1 rack
uzel a víc těch vykonávacích

- většinou jednotný OS

- rozsáhlejší výpočty třeba v celosvět.
mřížce - např. nějaké vědecké
výpočty

- volnější vazba -

např. se nepředpokládá jednotný OS

- např. OSGA

Distrib. Informační systémy - distrib. databáze, ...

Pervazivní systémy, IoT - velké množství malých zařízení

Peer-to-peer, decentral. systémy

~> nejde o výkon, ale tu spolupráci, distribuovatelnost

Motivace vývoje distr. syst.

- Ekonomika - spojením více běžných počítačů docílíme podobného výkonu drahých superpočítačů
- Rozšiřitelnost - přidání uzlu je levnější a jednodušší než upgrade jedné centralizované komponenty
- Spolehlivost - výpadek jednoho uzlu není tak velký problém
XX výpadek centralizovaného systému
- Výkon
- Distribuovatelnost - inherentní

Transparentnost

- vyšší vrstva nemusí řešit, jak to dělá nižší vrstva
- paralelní transparentnost - tedy je otevřeno, jak udělat obecně ^{automaticky} rozdělení problému na podproblémy pro paralelizaci

Přizpůsobivost

- autonomie - uzly rez. na ostatních + samostatná funkčnost
- otevřenost - obecná rozhraní

Škálovatelnost

- chceme se vyhnout čemukoliv centralizovanému
- uzly se rozhodují na základě lok. dostupných informací
- nelze se spolehnout na existenci přesných glob. hodin

servery, služby,
tabulky, atd. ...

Komunikace

- není sdílená paměť → komunikace pomocí zasílání zpráv
 ↓
 nejde třeba udělat semafor

- z pohledu klienta je transp., zda je to lok., nebo vzdálená kom.

TCP

- moc chytrý ⇒ nevhodný pro impl. clusterového systému (malá pravd. ztráty zpráv)
 (navrženo pro celosvět. kom.)
- vyšší latence - handshake

TCP/^{ack}FO^{open}, FLIP, ...

- specializované pro spolehlivé lok. sítě

Sémantika služeb

idempotentní služby - nevadí opakované provedení

- **exactly once** sémantika (ale ne vždy jde - např. tiskárna)
 nebo je to moc náročné
- **at least once** sémantika - stačí u idempotentních služeb ☺
- **at most once** sémantika - musí se dořešit na aplik. úrovni

Spolehlivost klienta

hardwar. klienta = dá požadavek a než přijde odpověď, tak umře

→ často se vůbec neřeší - jen třeba když je ten výpočet drahý či dlouhotrvající

◀ řešení

reintenance -

výpočty jsou rozděleny na

epochy, klient po zrušení

nastavuje novou epochu, přičemž server zruší výpočty z předchozích epoch

exterminace - klient po zrušení sám zruší tu službu

• expirace - úloha má přidělen čas, po překročení si server klientovi řekne "dolez"

Vzdálené volání

Kontrola!

- server - parametry, syntaxe, práva, ...
- klient - return value - jestli se to povedlo

↳ nízkorovňové !!

RPC - vypadá to jako lok. volání

→ Klient zavolá funkci ~> clientský stub ~> zpráva

clientský stub - rozbalení
↳ návrat z funkce

server skeleton (server-side stub)

skeleton - zabalení výsledku

rozbalení zprávy

zavolání výkonné funkce

mezi kompilací a vyvoláním

musí být ještě fáze

spojení

⇒ na dobře def. místo se
zanesou údaje toho serveru
(kde běží, ...)

řeší i třeba
load-balancing
directory server

→ klient se spojí s tím
directory serverem

rozhraní specifikováno v interface definition file

↓
to se předtím kompiluje

client stub, server stub

client
code

include

include

server
code

problémy RPC

- není zcela transparentní
- není tam společná paměť - někdo se dělá práce s rozsáhlejšími dat. strukturami
 - reprezentace dat - např. endianita, floaty, ...
 - ↳ musí být specifikováno
 - komunikační chyby - výpadky
 - ⇒ musí se řešit odlišně
 - bezpečnost

příklady:

- gRPC
- Google Protobufs

Skupinová komunikace

- vlastnosti:
- atomicita - zpráva je buď doručena všem, nebo nikomu
 - ↳ mechanismus pozdrčování
 - Synchronizace - aby uzly dostaly zprávy v nějakém dobře def. pořadí

uzavřená skupina - jen mezi členy
- vhodné pro keep. alg.

otevřená skupina - zasílat zprávy může kdokoliv
- spolehlivé, distribuované služby

překrývající se skupiny - může vadit, když není synchronizované
⇒ doručovací protokoly

Synchronizační algoritmy

- Synchronizace hodin mezi sebou - mezi všemi uzly
- Synchronizace s nějakým centrálním časem

C := vnitřní buh hodiny, fyzický čas t (UTC)

hodiny C ^{na uzlu p} čas t : $C_p(t)$

→ kdyby přesně, tak $\forall t: C_p(t) = t, \frac{dC}{dt} = 1$

↓

ale nejsou: $1 - \rho \leq \frac{dC}{dt} \leq 1 + \rho$

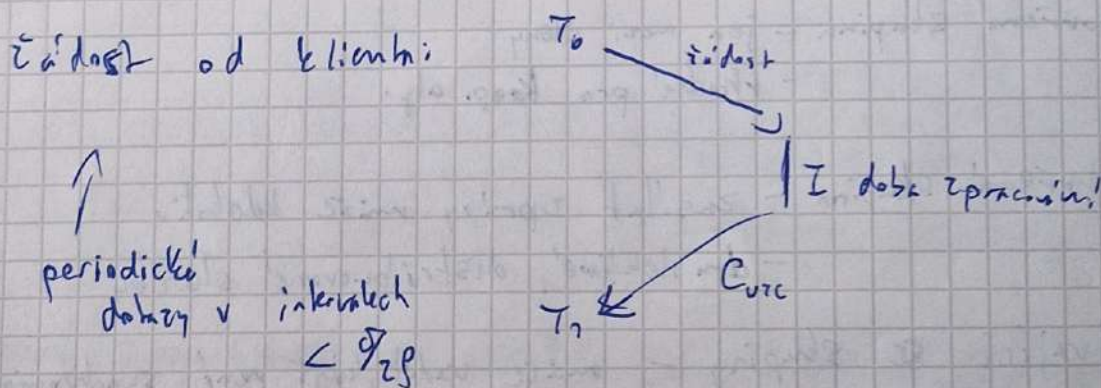
pro dvojce hodin: max. odchylka $2\rho \Delta t$

a my požadujeme odchylku δ : $\frac{\delta}{2\rho}$ Synchronizační intervaly

Fyzické hodiny

Cristianův alg.

- 1 pasivní časový server



$$T = T_{src} + (T_1 - T_0 - I) / 2$$

- nesmí být skoky = nepřerušit najednou (dozrání! - pat by se něco dalo vidět)
- dělat postupným zrychlováním/zpomalováním (např. posílček s pizzou)

Berkeley algorithmus

- 1 aktívny time server se periodicky pta' vsetch klientu
na odchylku od času (keby mu rebe)

$\Rightarrow$ zahodi' extrémny, spočíta' priemer

$\Rightarrow$ vsem pošle, o kolik si majú

preridit hodiny

zace ne
skakaj, ale
zpmaleniť /
zrychleniť
tikajú

Intersection algorithm

- nem' potrebu centralnu time server

- uzly v daných intervaloch broadcastujú svoj čas

1. uzol: 10 ± 2
2. uzol: 12 ± 1
3. uzol: 11 ± 1
:

$\hookrightarrow$ zahozem extrémny, spočítajú priemer a odchylky

$\rightarrow$ v upravenej verzii používajú NTP

Intervalový čas - čas je interval $now \pm \epsilon$

• DCE

• Google TrueTime

- napr. pri nepresnosti 200 μ s/s a synchronizácii každých 30 s

je odchylka ± 6 ms

Logické hodiny

- fyz. hodiny nejde synchronizovat dostatečně přesně pro všechny nějakých alg. - nepr. na úrovní instrukcí
- ukazuje se, že je hlavně důležité pořadí událostí, ne přesný čas
 - + procesy, které spolu komunikují, synchronizaci nepotřebují

! Kauzální závislost

zná: $e_1 \rightarrow^P e_2$ uspořádání událostí v rámci

procesu / uzlu P

řídí: jedny hodiny

send(m)

rcv(m)

odeslat / přijet zprávu m

DF: • Pokud $\exists p: e_1 \rightarrow^P e_2$, pak $e_1 \rightarrow e_2$

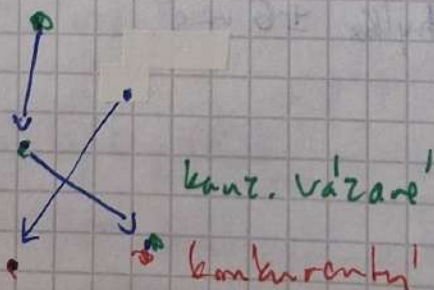
" e_1 kauzálně předchází e_2 "

• $\forall m: \text{send}(m) \rightarrow \text{rcv}(m)$

" e_2 je kauzálně vázaná událost e_1 "

• $e_1 \rightarrow e_2 \wedge e_2 \rightarrow e_3 \Rightarrow e_1 \rightarrow e_3$

DF: • Události e_1, e_2 jsou konkurenční $\Leftrightarrow e_1 \nrightarrow e_2 \wedge e_2 \nrightarrow e_1$



→ evidují kauzalitu: události $a, b \Rightarrow$ časové značky $c(a), c(b)$

ale pozor: když $c(a) < c(b)$, tak nemusí platit $a \rightarrow b$ jestliže $a \rightarrow b$, pak $c(a) < c(b)$

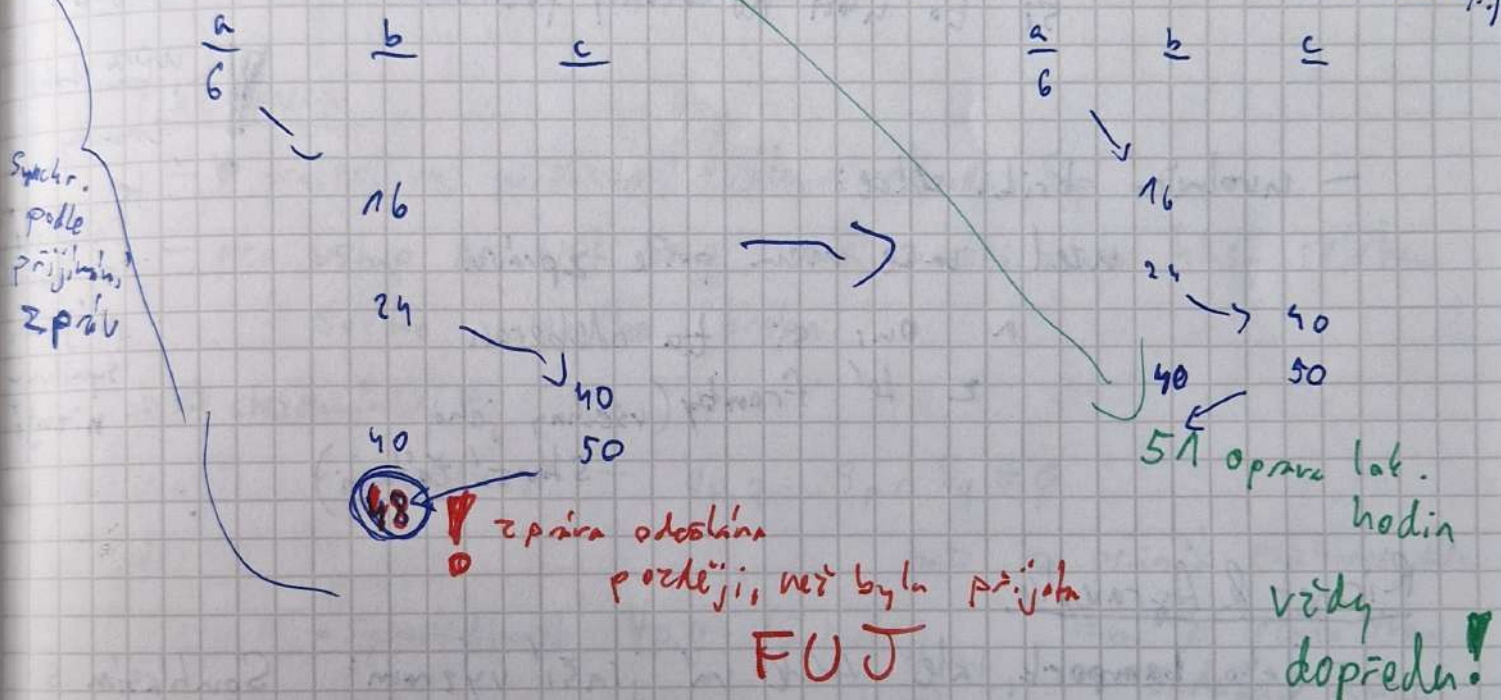
Synchronizace log. hodin

proces i vysílá v čase $C_i(a)$ zprávu $m \rightarrow T_m = C_i(a)$

proces j přijme zprávu m v čase $C_j(b)$

→ pak $C_j = \max(C_j(b), T_m + 1)$

def. náte' DS, která reprezentuje ty značky (číslo, vektor, ...)



zúplnění - aby jakékoli 2 události měly rozdílné čas. raz.

$$C(a) = C(b) \wedge P_i < P_j \Rightarrow C'(a) < C'(b)$$

↑
událost v proc. i ↑
událost v proc. j ↑
nejaké označení procesů i, j

= "byrokratické usp."

Distrib. synchron. alg. --> mutual exclusion

Centralizovaný alg.

- je jeden server s frontou, kam jsou odesílány posílají žádost a až na ně přijde řada, tak jim server řekne
- FUJ - centralizovaná komponenta

při úpadku serveru musí všichni klienti server kontaktovat znovu atd. ...

Lamportův alg.

- když proces chce vejít do krit. sekce, tak pošle žádost všem - ti si to uloží do fronty žádostí a zpět mu pošlou potvrzení
- do krit. sekce může vstoupit ten proces, jehož žádost má nejstarší čas. značku a má už od všech potvrzení
(tedy ten uzel u sebe nevede starší žádost od něho jiného)
- když proces přijme potvrzení, tak si to uloží do fronty potvrzení
- uvolnění krit. sekce:
uzel zase všem pošle zprávu a oni si ~~to~~ odeberou z té fronty (všechny jeho starší žádosti)

S upravenou čas. značkou
(podle čas. značky té žádosti)
si to přijemci synchronizují

Ricart & Agrawala

- jako Lamport, ale ACK má jinší význam: Souhlasím s tím, aby s vstoupil do krit. sekce
- když chce proces vejít do krit. sekce, tak všem pošle žádost (se svou TS)
- když proces přijme žádost:
 - když ho nezajímá krit. sekce → ACK
 - když je v krit. sekci → dokonec ~~svou~~ práci tam a až pak pošle ACK všem z fronty
 - když není v krit. sekci, ale chce tam být:
 - když jeho žádost má nižší TS, tak
 - jinak ACK a čeká

Princip voleb

- proces sbírá hlasy od ostatních
 - když dostane víc hlasů než jakýkoliv jiný, tak může vstoupit do krit. sekce
- #proces má 1 hlas

Naivní volby

- prostě mu ten hlas dáme (pokud jsme ještě nehlasovali)
- deadlock! , složitost $O(n)$ ☹
- ale odohrání více případů ☹

Maekawa

- #proces má přiřazený volební okrsek S_p
- pro vstup do krit. sekce musí proces získat hlasy celého svého volebního okrsku

- podmínky:

- korektnost - $\forall p, q: S_p \cap S_q \neq \emptyset$

→ ten průnik rozhodne, kdo

- spravedlnost: $\forall p, q: |S_p| = |S_q| = k$
↑
±1

tam z těchto dvou
úspěšně může
vstoupit

- zodpovědnost: $\forall p, q: |\{S_i: p \in S_i\}| = |\{S_j: q \in S_j\}| = D$

⇒ kom. složitost $O(|S_p|)$ ☺

$N := \# \text{ okrsků}$

$k := \text{velikost okrsku}$

$M := \# \text{ procesů}$

$D := \# \text{ okrsků, ve kterých je každý proces}$

$$N \cdot k = D \cdot M$$

pro proces ≈ okrsek: $N=M \Rightarrow k=D$

max. # okrsků
je
 $(D-1)k+1$

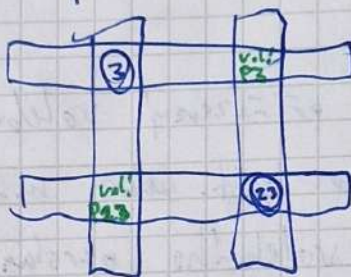
#proces q z okrsku S_p je obsazen v $D-1$ jiných okrscích
⇒ $M = k \cdot (D-1) + 1 \Rightarrow O(\sqrt{M}) = k$

Rozdělování do okresek

- opt. alg. FUJ s přizměně členství
vytvoření restruktu-
ralizací
- subopt. pro $k = O(n)$

⇒ dáme si ty procesy do úmístky
a procesu přidáme okresek
z jeho řádku a sloupce

→ ale pozor na deadlock!



Soubor P_3 a P_{23}

⇒ nikdo z nich nedostane
všechny hlasy ze
svoje okresek

rešen!
pomocí fog. hodin

— když už proces hlasoval, tak
si porovná TS_q s TS_r
tadažle: ↑
tentamín dal hlas

$TS_r > TS_q \rightarrow$ dá si číslu do fronty

$TS_r < TS_q \Rightarrow$ pošle q
zpět REJECT

q už
je v krit.
sekc. $\Rightarrow$ dodá a
pak odpoví

q ještě nemá
všechny hlasy
 $\Rightarrow$ vrátí hlas proces p
a kn ho dá
procesu r

race condition!
ale nevadí - neřešíme
výpočet,
ale obranu
proti deadlocku

Stromové alg.

Agrawal & El Abbadi - $O(\log n)$

- úplný Bin. strom - proces musí získať hlasy všetkých na lib. ceste, ktorá ide od korene do listu přes ten proces
- Fault tolerance - havarovaný uzel lze náhodit cestou v jeho levém a pravém podstromu
- Asymetrický - ten v koreni musí furt dělat nějaké hláskování!

$\Rightarrow$ řešení: nějak se to robuje

$\rightarrow$ ale tím se to zesložitěji a vyjde to horší než ten ad-hocový protokol

Token-based alg.

Raymond

- 1 token v krit. sekci
- uzly organizovaný do or. stromu, na vrcholu je vždy uzel s tokenem
- uzel chce do krit. sekce: ve směru or. pošle žádost
 - $\rightarrow$ ten, kdo má ten token a už není v krit. sekci pošle proti směru ten token (a má se orientace)

Suzuki & Kamei

- hodí se pro koop. distr. výpočty + máme broadcast ^{přímá konektivita všech se všemi}
- žádost o vstup se pošle broadcastem; ten, který drží token, ho přidá do fronty

Token ring

- má jen lokální znalost 😊

- uzel si token nechá, když chce vstoupit do krit. sekce,
a pak ho pošle dál

- různé varianty - když nikdo nechce do té krit. sekce,
tak se zbytečně vyčerpá síť
→ takže např. u síťové a
když uzlu přijde ten token
má rychle odminout,
tak ho trochu pozdrží

všechny ty distr. alg.

mají velké problémy
při výpadcích (|)

→ snazší se řešit centralizovaný problém (vstup do
jedné krit. sekce)
distributivním způsobem

takže
pro menší
síť se
vyplácí ten
centralizovaný alg.

a u větších se chceme potřebu

vzájemného
vyloučení
vyhnout

→ FVT

Volba koordinátora (leader election)

→ aby měl nějakou roli

nejdůležitější je, aby nemohli zároveň $\exists$ 2 koordinátoři

přístupy:

- detekce extrému

- předp.: každý uzel je identifikován nějakým řetězcem id

- Bully, Invitation, Kraken alg.

- race condition (+ randomizace)

- necháme ty procesy něco dělat a který náhodou vyhraje, tak se stane koordinátorem

- ale ta podmínka výhry nemusí být jednoznačná

$\Rightarrow$ randomizace

↑
s operací
↓

Detekce extrému

Bully alg.

- předp.: homogenní uzly, hodně spolehlivá síť

$\Rightarrow$ spolehlivá detekce havárie

(omezená doba přenosu zprávy T_m

+ omezená doba zpracování zprávy

a odpovědi)

ne
třeba
TCP

$2T_m + T_p + c$

nelze použít
v async.
systému

1. kolo: osloví se
všichni
možní
kandidáti

2. kolo: vítěz všem oznámí,
že je novým
koordinátorem

- proces, který chce volit, pošle všem uzlům s vyšším id zprávu

přišla odpověď: hotovo

neřišla: stane se koordinátorem

BLE

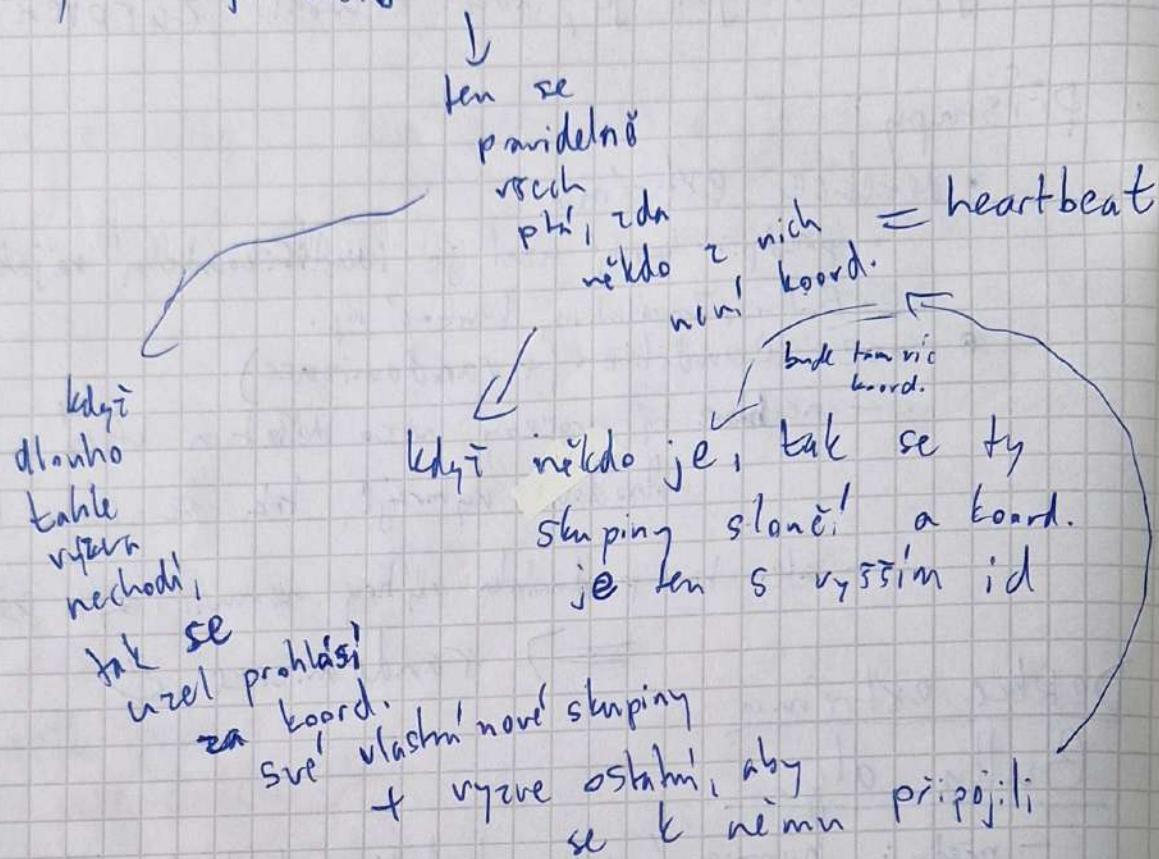
víc koord.

byly by
falešné
detekce
havárií

- proces dostane zprávu o volbě: odpoví a poše zprávu s vyšším id procesem

Invitation alg.

- už nepředpokládá, že absence odpovědi = havárie ☺
- Skupiny - mají svého koord.

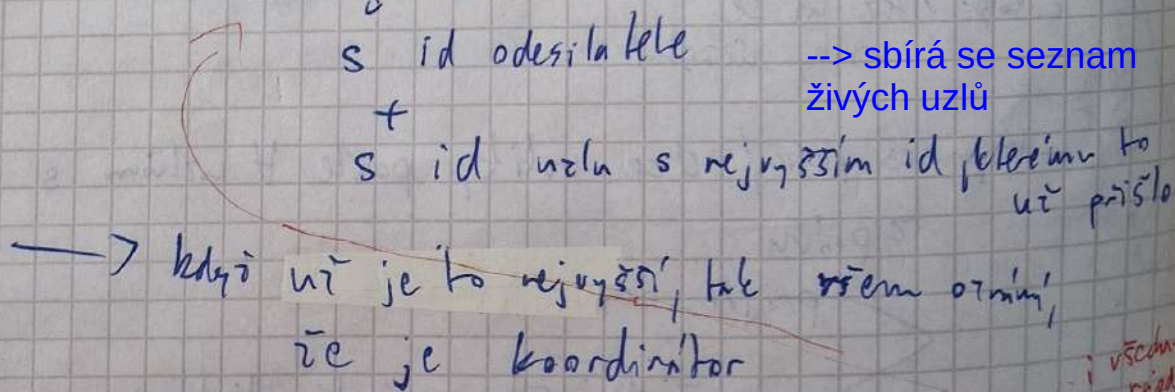


Kruhové alg.

- pro rozsáhlejší prostředí, jednotlivé uzly nemají ^{kompletní} znalost

Le Lann, Chang, Roberts

- proces pošle zprávu podél orient. hrany



- složitost $O(n^2)$ - může znát volit lib. proces

HS alg. - optimalizace

- orientace hran je obousměrná

- v několika kolech,

postupně se združují sobě
obklopující, ve kterém

se volí a vítězí se
účastní jen ten

$$O(n \log n)$$

vítěz z
předchozího
kola

Randomizované protokoly

- úplně za začátku jsou všichni **followeri**; každý po rand. timeoutu přejde do role **kandidáta** a pošle ostatním žádost o hlas pro tuto **epochu**
 - => příjemce má novější epochu => nedá mu hlas
 - => příjemce má aktuální epochu a ještě nehlasoval (ani sám není kandidát) => dá mu hlas
 - => příjemce je kandidát v $\geq$ epoše => nedá mu hlas
 - => příjemce má starší epochu => dá mu hlas a pokud je kandidátem, tak se kandidatury vzdá
- když se dlouho nedří získat nadpoloviční většinu hlasů, tak každý uzel počká náhodný timeout a zkusí se to celé znovu
- funguje dobře, když $T \gg$ doba přenosu zprávy
- v clusterech, cloudech
- když followerovi dlouho nechodí heartbeat, tak se stane kandidátem a zahájí novou epochu

např. v **RAFTu** (viz později)

2-fázový commit

- jiný přístup: uzly náhodně vygenerují číslo -> pošlou ostatním

příjem zprávy
- v bufferu

doručení zprávy
- přijetí + splnění podmínek

když má někdo
lepšího kandidáta,
tak to
nepotvrdí

Doručovací protokoly

Globální usp.

- > zprávy doručování
v pořadí odeslání
- NEJDE - nemáme dost
přesné glob.
hodiny

Sekv. (totální) usp.

- > Všechny mají
čísla a podle
nich jsou usp.
- bez ohledu na
řez. čas odeslání
leže podle těch čísel
určit jejich pořadí

server - centralizovaná
komponenta udržuje
ka čísla
dvofázový distrib. alg.

Atomický multikast

- sekv. + zahrnuje atomický
bude doručeno
všem, nebo
nikomu

Kanáliz. usp.

- > kanál. vázané zprávy v pořadí kanáliz.
kombinováno v lib. pořadí
- když přijde zpráva, ale by k
ní kanáliz. vázané, ještě ne
všechny, tak se ještě nesmí
doručit

kanaliz. -> obsah vázané
zprávy
může být
ovlivněn

Jak příjemce zjistí, že už může tu zprávu zpracovat

$dest(m) :=$ množina uzlů, kterým je zaslána zpráva m

$deliver_p(m) :=$ událost doručení zprávy m uzlu p

$m_1 \rightarrow m_2 \Rightarrow \forall p \in (dest(m_1) \cap dest(m_2)):$

chceme?

$deliver_p(m_1) \rightarrow deliver_p(m_2)$

Vektorové hodiny

$n :=$ # procesů ve skupině

$VT(p), VT(m)$ - čas. značky procesů/zpráv
(vektory!)



na začátku $VT(p_i) = 0$

(svou složku)

při odeslání: zrušíme si VT a to bude $VT(m)$

při doručení: upravíme si podle zprávy
(ne přijeli!) své $VT(p)$

- po složkách na max.

Kanální doruč. protokol

odeslání zprávy m uzlem p_i :

• $VT(p_i)[i]++$

• $VT(m) := VT(p_i)$

$\Rightarrow$ pošle $VT(m)$ a tu m

přijetí uzlem p_j : doručení se pozdrží, dokud:

• $VT(m)[k] = VT(p_j)[k] + 1$ pro $k = i$

• $VT(m)[k] \leq VT(p_j)[k]$ jinak

$\rightarrow$ = musí být doručeny všechny by
kan. včasně předchozí zprávy

doručení:

p_j si upraví VT : $\forall k = 1, \dots, n: VT(p_j)[k] = \max(VT(p_i)[k], VT(m)[k])$

Kauzální doruč. prot. pro překrývající se skupiny

→ maticové hodiny - Vuzel si udržuje ty skupiny,

odesílání m uzlem p_i do skup. g_a :

$$\bullet VT_a(p_i)[i]++$$

$$\bullet VT(m) := \bigcup_{b: p_i \in g_b} VT_b(p_i)$$

jejich
je součástí

příjem zprávy m uzlem p_j (od p_i , s $VT(m)$, do g_a)

→ doručení se pozdraví, dohled nepřít:

jako
vekt.

$$\begin{cases} \bullet VT_a(m)[i] = VT_a(p_j)[i] + 1 \\ \bullet \forall k: (p_k \in g_a \wedge k \neq i): VT_a(m)[k] \leq VT_a(p_j)[k] \end{cases}$$

kauzální
na

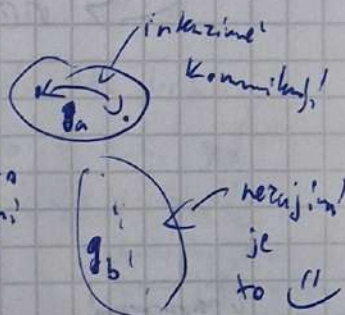
$$\bullet \forall b: (p_j \in g_b): VT_b(m) \leq VT_b(p_j)$$

primární
všech skupin
primární
odesílatele & příjemce

~~vektorové hodiny přes všechny skupiny~~

→ mat. hodiny umožňují izolaci

lokálního
chození
uzlu



→ u vekt. bychom vynutili

že by se musely všechny
zprávy posílat všem (1)
(protože na
ně musí čekat)

neřeší
se fyz. usp.
ale dodržuje
jednu
posl. událost
na všech příjemcích

Distribuce / total-order prot. => sek. doručování

=> zprávy uzlům doručeny ve stejném pořadí -> stejn. sledování
hodiny

příjem zprávy: Synchronizace lok. hodin s $TS(m)$

+ pošle se potvrzení s $TS(p_j)$

při obdržení
uzlu
zprávy
doručí
(je jediné
uspořádané
těch událostí)

→ odesílatel pak pošle finální zprávu s $TSF := \max(TS(p_k))$

Centralizovaný total-order prot.

11 - zase je to centralizovaný problém
tak nemůžeme
sáhnout
k řešení
bez
vstupu
alg.

Sequencer = proces, který synchronizuje pořadí událostí

→ jemu se posílají ty zprávy a on je pak přeposílá
(příjemci ještě musí zajistit, aby se ty zprávy nepředcházely - pomocí nějakého čísla)

a příjemci mu pošlou potvrzení

Spolehlivé doručování

→ členům skupiny bude doručena zpráva

→ spolehlivý kanál? → ALE: • výpady příjemce

• výpady odesílatele

→ řešení: • transakce
- třífázové
- příliš striktní "

ještě
HORSÍ

• záplavový alg.

- při příjmu dosud nepřijaté zprávy ji uzel přepošle všem členům skupiny

→ spolehlivé, ale neefektivní ($O(n^2)$)

• alg. s potvrzováním

- stabilní zpráva = byla doručena všem uzlům

- jako záplavový, ale nepřeposílá se všem, ale jen potřebným

zhasnutý uzel
= po nějakou dobu nebyl schopný na nic odpovídat

→ příjemci si ji taky po chvíli dala načíst (a pošle ALE)

- pak už ty zprávy přeposílá všem, kdo ji nepotvrdili, jak zjistit: ...

→ když uzel odešle zprávu, tak si ji příjemci načítají, dokud nedostanou od všech ostatních potvrzení, teprve ji přijali (nebo do lineárního)

Trans alg.

transitive Acks

→ ack(m) se posílá, když přijmeme m a v kauzální řadě
↳ nemusí se potvrdit ty předchozí předchozí zprávy

- Uzly si pomáhají graf kauzality slumping (DAO)

- m' potvrzuje m → (m', m) ∈ G

- 3 komponenty: příjem ack a vypořádání s ack

Gp := { zprávy, které p přijal a ještě nejsou stabilní }

Jak p zjistí, které uzly zprávu přijaly:

zpráva obsahuje i seznam všech zpráv, které už byly kauz. potvrzeny

Uzel má: ack-list = seznam zpráv pro potvrzení

• nak-list = nepřijaté zprávy

• undelivered-list = přijaté, ale nedonosené

pošle s každou zprávou

když přijme zprávu:

1) podívá se na její nak-list

a pokud něco z toho má, tak to VŠECH pošle

(šlo by jenom tím jedním uzlem)

5) tu zprávu si dá do undelivered-listu

6) prohrabe se undelivered-listem, jestli nějaká zpráva nemá splněné ty podmínky (jsou doručeny její kauzální předchůdci)

7) když už ji dostal (m ∈ Gp), tak konec

8) když je v jeho vlastním nak-listu, tak ji z toho vyndá

9) podívá se, jestli v ack-listu má zprávy nemějící, kterou nemá

3) z Gp si odebere { stabilní zprávy }

4) dá ji do ack-listu

a když tak ji doručí

když si ji dá do nak-listu

protože když je problém na jednom uzlu, tak je typický i na dalších

optimalizace - snížení latence řešení problému

když ale nějaký uzel bude zhasínat, tak alg. do nekonečna čeká!!

Virtuální synchronie

--> koncept řešení problému konzistence při selhání
- při havárii nastane změna
pohledu na nový ($L^x \rightarrow L^{x+1}$)

Group view - množina uzlů ve skupině
značí: L, L_p, L^x, L_p^x ← verze pohledu

↑
jak
proces
P vidí
kumulativně

↓ procesy

P, q současně L^x a L^{x+1}

↑ verze
značí pohledu

$$\Rightarrow \text{install } L_p^x < \text{deliver}_p(m) < \text{install } L_p^{x+1} \Rightarrow \text{install } L_q^x$$

↑
 $\text{deliver}_p(m)$
↑
 $\text{install } L_p^x$

↑
pohledově-synchronní komunikace

$$\Rightarrow p \in L_q \Rightarrow q \in L_p \quad (= \text{podmínka vzájemné konzistence})$$

implementuje

- když je zpráva doručena nějakému uzlu a odesílatel zhasavne, tak musí protokol zajistit, aby byla doručena i ostatním
- nesmí se doručovat do jiných pohledů

Transis alg. - rozšíření Trans

- Paranoia = stáčí, aby jeden uzel měl podezření, že jiný uzel je zhasavovaný, a už ho ten uzel označí za zhasavovaného a všechny uzly z toho pohledu to musí respektovat

- jednosměrnost = když je někdo označen za zhasavovaného, tak je ze skupiny vyloučen (a nic proti tomu nemůže dělat)
burdčí, protože v distr. prostředí nelze při nesplnitelnosti dosáhnout společného konsenzu
→ jenom pak znovu počítat o členství

$Last[i] :=$ uzly označové proc. i za zhasavované

$J :=$ nově přidávané uzly

$JLast[i] :=$ uzly naposledy označové proc. i

konkrétně vázané
způsob si poradit
(blocked)
za
přidávané

v konkrétním uzlu se nastaví hranice inkluze nového pohledu a individuálně konzistentní inkluze

ISIS/VSync protokol

- používá maticové hodiny (MT)

$MT_{P_i}[j][k]$ - co uzel P_i ví o doručení zpráv uzlu P_j od P_k

- příjem zprávy od P_j :

$$MT_{P_i}[i][j] = VT_m[j]$$

$$MT_{P_i}[j][*] = VT_m[*]$$

- proces udržuje zprávy, dokud nejsou stabilní

- při instalaci nového pohledu:

- přepošle všechny nestabilní zprávy
- pošle mu pošlou flush zpráv - potvrzení instalace
- když jich má dost, tak nainstaluje ten nový pohled (od všech, kteří nejsou v množině zhařadovaných)

Globalní stav a Konsensus

Diffusing Computation

uzly tvoří orientovaný graf, ^{dynamický} kostka, v opačném směru kam. kanálu jsou signální kanály

Ukončení distr. procesu

stav uzlu $\left\{ \begin{array}{l} \text{aktivní} \\ \text{pasivní} \end{array} \right.$

- může přijímat zprávy a přejít při tom do akt. stavu

stejně je ten, od koho přišla první zpráva v tom výpočtu

konec výpočtu $\rightarrow$ všechny uzly pasivní

Dijkstra-Scholten (DS) algoritmus

Strom

- když se list dostane do pas. stavu, pošle signál ohe
- u aktiv. uzlu když dostane signály od všech synů, tak pošle signál směrem ohe

DAS

- v k hraně je číselník = deficit - rozdíl mezi apl. zprávami poslanými tam a signálními zprávami poslanými zpět
- končí proces počítá na nulový deficit dat. kanálu
- a pak každým signálním kanálem pošle tolik signálů, aby v nich vynuloval deficit

Obecný graf

- dynamicky se vytváří kostka
- proces s signálními kanály končí tak, že pošle signál, čísla na zpětné signály a pak pošle signál ohe

Huanghir alg.

V proces^u ^{zpráva} na váhu

W = váha iniciátorů, ostatní mají váhu 0

- při odeslání zprávy uzel dá část své váhy té zprávě
- při příjmu si proces váhu zprávy vezme
- při ukončení proces pošle zprávu s celou svou vahou
 - konec výpočtu: váha iniciátorů = W

problémy:

- dělitelnost vah
- havárie procesů
- ztráty vah

Značkové alg. (TM)

- značka = speciální zpráva
- když iniciátor uzel má pocit, že by měl výpočet už být ukončen, tak všem výst. kanálům pošle značku
- při příjmu značky:
 - uzel aktivní - nade
 - uzel pasivní - propaguje značku výst. kanálům

→ definice glob. stavu

- množina událostí E

- řez $C \equiv$ ^{disj.} rozdělení E na P_C a F_C

↳ je kauzálně konz. $\equiv a \rightarrow b \wedge a \in F_C \Rightarrow b \in F_C$

(minulost nelze měnit v budoucnosti)

- stav $\equiv$ množina událostí, co už proběhlo

↳ konzistentní $\equiv S = P_C$ pro C kauz. konz. řez

- stav uzlu = množina přijatých a odeslaných zpráv + ^{dodatečné informace} ^{specifické pro to, co řešíme}
- stav kanálu = množina zpráv, které byly kanálem odeslány, ale ještě nebyly doručeny

→ iniciátor všem výst. kanálům pošle značku

- při příjmu první značky uzlem: = okamžik lokálního řezu

- zapamätuje si svůj stav,
využije stav všech vstupních kanálů

- všem výst. kanálům pošle značku

- příjem zpráv kanálů, ze kterých ještě nepřišla značka

→ uzel si zapamätuje jejich čísla

- příjem dalších značek → stav kanálu = zprávy dostě uzel přijímá první a tento stav

→ konec po přijetí všech značek

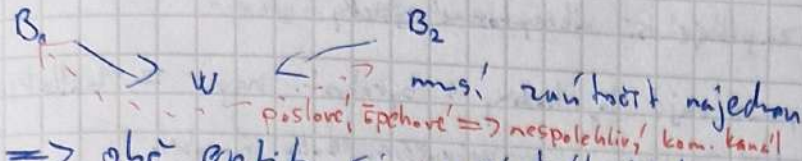
konzistentní
stav systému
je zároveň
uzel a kanál

Distribuvovaný konsensus

Byzantský = 1 → 1
 konsensus = n → 1
 Interaktivní konzistence = n → n

Problém dvou armád - koordinace entit - spolehlivý uzly, nespolehlivý kanály

- partit. např.
- replikace serveru
- replikace logů
- synchronizace
- baricerní orchestrace
- distrib. key-val. databáze
- obecné cloudové aplikace



⇒ obě entity si musí být jisty, že ví, co mají dělat

Praktická řešení

- agresivní strategie - vysle se dostatečně množství zpráv
- pravděpod. strategie - když se ty charakteristicky kanál mění - pošle se víc zpráv a k tomu i jejich počet
 → ta druhá strana si porovná, kolik toho přišlo, s tím počtem

Problém Byzantských generálů

- uzly mohou lhát a chovat se zkeřně
- spolehlivá komunikace

BÚNO 1 generál, ostatní důstojníci, rozkaz bude vydán na základě většiny
 chová: • Uzly se shodnou na hodnotě
 • Pokud hodnota navrhl nekooperativní uzal, tak se uzly na ni shodnou

⇒ pro f zrádců řešení pro $n \leq 3f$ uzlů, je potřeba $n \geq 3f+1$ uzlů

Praktická řešení

- PBFT
- kryptografie, podpisy

4 fáze: (od klienta)
 • request zvěřý stran - řekne nějakému nodu → ten ho přepíše všem + zakašuje
 • pre-prepare - nody ověří, že zpráva s hashem a všem přepošlou

Paxos protokoly

- nespolehlivá síť, nespolehlivé uzly, ale nejsou zkeřné (když nefungují, tak neškodí)
- Fault-tolerant replikační automat
 → Uzly provádí stejný zvěřý na stejném automatu

fault-tolerance - není to výjimečná událost, ale partit. se tím
 + timeouty, checkpoints
 • prepare - po obdržení $2f+1$ prepare zpráv všechny svoji → ověříme a pošleme COMMIT
 • commit - po obdržení $2f+1$ commit zpráv → ověříme a provedeme

polyomální, ale náročnost a citlivost na Sybil attack
 → útok si vygeneruje víc trůh identit a tím přehluší ty korektní uzly

Part-time parliament

- poslanci > nespolehliví, asynchronní systém (bez timeoutů) + není potřeba exactly-once semantics
- posíláči

→ každý poslanec má zálohu a nesmazatelný interval (disk) info. tam zůstane zachráněná, i když si poslanci odkaš!

- role:

- Client
- Proposer - navrhuje nový stav
- Acceptor v Quorum - schvaluje návrhy
- Learner - zapisuje si akceptované výsledky

v každém uzlu - moduly

- pro log. usp. se používají log. hodiny s byrokrat. usp. = ballot numbers

→ akceptují akceptují jen zprávy s nejvyšším Ballot number

→ používají si:

- Last Ballot
- Accept Ballot
- Accept Value

- protokol:

(zabýváme)

1. Fáze: volba leadera, ještě se vůbec neví, že hodnota

2. Fáze: akceptace hodnoty

leader navrhuje

cílem je potvrdit
proposerovi,
že má dostatek
hlasů na
navrhování

kyž umře
starý leader, ale
ještě se nestihla
dosáhnout
konsenzu,
tak aby to
nový leader
ještě dodělal

ale nepoužívá
se - neefektivní

Korektnost není
deterministicky
zaručena

→ v praxi se

často při konfliktech
lávují randomizované
timeouty

• Multi-Paxos

→ 1. fáze jen jednou
a pak už jen 2.
(dokud nedojde ke
konfliktu)

• Cheap Paxos

• Fast Paxos

• Byzantský Paxos

→ V zprávy mají ještě jednu položku - číslo instance

jsou tam záložní uzly pro výpadky

Client rovnou posílá acceptacím, při konfliktu to vyřeší leader

ještě verifikace, ale moc se nepoužívá

(určí pořadí a udělá to
v tom pořadí jako v
basic verzi)

- impl. a aplikace:

- Google Chubby - distrib. zamykání
- Google Borg / Omega / Kubernetes
- Derecho

- ale je to složité a těžko pochopitelné !!

RAFT

- navržen, aby byl snadno pochopitelný a implementovatelný

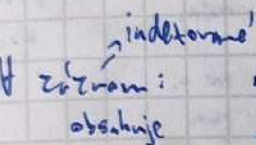
- role:

- leader - navrhuje
- Follower - akceptuje
- candidate - dočasná role při volbě leadera

- komponenty:

- volba leadera - on pravidelně posílá heartbeat a když to nedojde, tak je to důvod pro volbu nového leadera
- replikace logu - když leaderovi přijde nějaká zpráva od klienta, tak to pošle všem followerům
- udržování konzistence záznamů

- prüfend:



První replik s akt. kermem je ten committing stav (ten z nich velikos

- kdo nedostane heartbeat, tak se přihlásí za kandidáta a v poše request vote
- ⇒ když dostane většinou / # hlasů, tak se stane leaderem

- do určitého
timeoutu
od začátku
volby se
nezulidne
zvotit leader
- randomi zvráto
timeoutu

leader
men pak postupně
posílá
rozšíření
Append Entries
S data předložení zápisu,
dokud se neshodnou
→ follower si to
odsad celé
přepíše na
to od
leadera

 ZAB

- součástí Yahoo Zookeeper frameworku
- účel: backupy systémů → zálohový se nesmí pokazovat

Distribuovaná sdílená paměť

paral. arch.

multiproc.

sedlá zóna
multicomp. = kuzel
místní paměť

sběrnice arch.

přepínací arch.

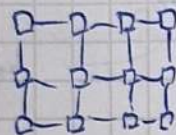


nízká škálovatelnost

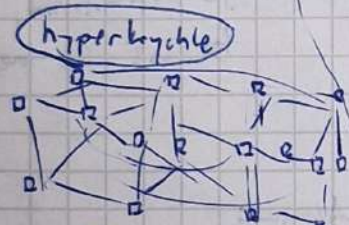
sběrnice arch.



přepínací arch.



→ např. počítač



→ superpočítáče



inžinýrská

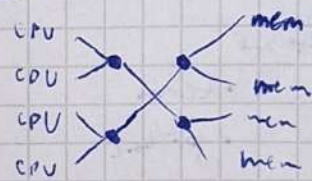
Crosspoint switch

- na různé adresy přepíná

⇒ větší škálovatelnost, ale hrušná náročnost !!
(kvalita přepínání)

je potřeba dobře
optimalizovat SW,
aby se to
rozprostřelo
přes ty
blahy
paměti

Omega

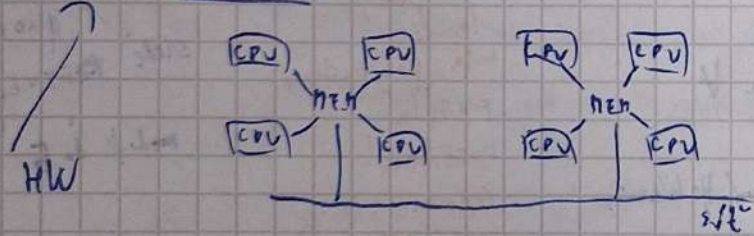


⇒ menší hrušná náročnost,
ale stejně

NUHA

GPSPU

CPU + GPU



KDMA

speciální síť. karta → přímé sdílení paměti !!
s velmi nízkou latencí !! !! !!

Konzistenční modely

pro virt. paměť - bez potřeby interakce s vnějšími procesy - transparentní

pro sdílená data - pomocí frameworku → efektivnější

• striktní / atomická konzistence

- všechny zápisy okamžitě všude viditelné
- v distrib. systému neúnatitelné - musí být glob. přesný čas

Znač.

$W(x)$ - zápis x
do x
 $R(x)$ - přečtení x
z x
 S - synchronizace
Acq - vstup do krit. sekce
Rel - výstup z krit. sekce
 P_i - proces i

Signatura
= výstupy jsou
seřazeny
podle uslo

• sekvenční konzistence

- výsledek je ekvivalentní s tím, že by z jeden rozvrh, pořadí respektuje lok. pořadí
- na vuzlu vidíme stejné změny proměnných, jen trochu v jiném čase
- není determin. - dvakrát spuštění stejného programu může dát jiné výsledky
- nejsilnější model, který je rozumně možné implementovat v distrib. prostředí
- ale pomalý $\ll$ $rtw \geq t$

→ jde optimalizovat jen čtení, nebo zápis

• externí konzistence

- sekv. + čas propagace omezen
- v distrib. databázích "citlivost lidského oka"

• kauzální konzistence

- zápisy, které se na sebe kauzálně vztahují, musí být na všech uzlech ve stejném pořadí
- odesláni, čtení $\approx$ přijetí

$$W(x)_a \rightarrow R(x) \xrightarrow{p_x} W(y)_b$$

kauz. záv.

- složitější impl. - musíme mít ten graf závislosti

• PRAM konzistence

- zápisy jedním uzlem jsou na vuzly doručeny v tomhle lok. pořadí, ale zápisy různých uzlů mohou být viděny různě
- Uzel má vlastní rozvrh (x sekvenční konz.)

• Slow memory

- synchronizují se jen zápisy do jednoho místa

na úrovni

správy vřt.

paměti - nepotřebují!

podpora od jazyka

akt. - jen propagaci

!! zápisy, transparentní !!, ale pomalější !!

je stránka

propagace až při opuštění

Konz. modely se synchr. proměnnou

- proces to přímo musí ovládat - třeba vstup do krit. sekce
- ale ztráta transparentnosti - musíme vědět, že pracujeme s distrib. pamětí, jak se ovládá ten framework, ...

Znač.: synchronizační proměnná $\rightarrow S, Acq, Rel$

• Weak consistency

- přístupy v blocích omezených přístupy k synchr. proměnným $\Rightarrow$ sekv. konzistentní propagace hodnot

- prakticky nepoužitelné - nerozlišuje se Acq a Rel - vždy se dělá synchronizace pro oba případy

- release consistency
 - Acq a Rel prakticky konzistentní
 - ↳ dokončení zápisu i čtení
 - ↳ čtení

- impl.:

- eager rel. cons. - při Rel se všem rozešlou všechny aktuální verze dat
 - ⇒ ostatní to dostanou hned ☺, ale zbytečně to zatěžuje síť ☹
- lazy rel. cons. - propadne se až při Acq procesu
 - ⇒ pošle se to jen těm, kteří to chtějí ☺, ale vyss!

problém u Acq ☹

- entry consistency
 - Acq a Rel jsou vázány na nějaké množiny (mohou se prolínat)
 - exklusivní a neexklusivní přístup k synchr. proměnné
 - ↳ zápis
 - ↳ čtení
 - ⇒ může být jeden zapisovatel / více čtenářů

Distrib. stránkování

- po dobu, kdy je stránka načtená, nejde nijak detekovat čtení/zápis do ní
 - ⇒ to my ale potřebujeme kvůli té synchronizaci ⇒ dělá se read-only napování

Načtení stránky

- centralizovaný manager - eviduje, kdo má jaké stránky
 - v nějakých malých clusterech
- replikovaný manager
 - různé skupiny stránek spravovány různými managery (trocha práce spočítání kb)
 - už i pro větší clustery ☺
- broadcast (hu) - a majitel té stránky se sám přihlásí a pošle nám
 - menší clustery

⇒ při zápisu dáváte k přeměně

Správa kopií

- broadcast - všichni, kdo máte verzi 57 a nižší, tak si to zinvalidujte
- copy set
 - zodpovědnost na to, kdo všechno má kopie té stránky, je přenesena na vlastníka stránky - a děláme to přes něj

Uvolňování stránek - kleron vyhodit

- nevlastní read-only kopie - máme jistotu, že vlastník ji má, takže o to nikdo nepříjde
- vlastněná replikovaná kopie (v copy setu) → přeneseme vlastnictví na někoho, kdo má tu kopii
- lokální heuristika

False sharing

- prevence - hinty překladači, aby to řekl linkeru ⇒ ztráta transparentnosti ☹
- hysterese - když máme stránku jenom chvíli a netřeba ji číst, tak mu ji nedáme

Sekv. konz. stránkování

- při zápisu musí být stránka jen na jednom uzlu
- když se normálně čtou/zapisují namapované stránky, tak to rejde nijak zdetekovat
- > akce až při výpadku stránky

<u>my-chce</u>	<u>mapováno u vlastnicka</u>	
R	R	✓ ani nepotřebuje
R	W	vlastník si ze stránky udělá R=0
W	W	my = vlastník
W	R	povýšime na W
W	R	všechny z copysetu získáme, až si to zinvalidují
W, mkr	R	vlastník nám předá vlastníku, invaliduje si to
W	W	

Konz. konz. stránkování

- vektorové hodiny - VTs, VTP, složky jsou stránky

↳ verze stránek v procesu p
na kterých stránkách stránka konz. závisí

ale efektivnější

Distrib. sdílené proměnné - na úrovni Frameworku - není abstrakce (X stránkování)

- distrib. objekty ← základní kámen v těch Frameworkech - viz předmět Middleware

Epidemické protokoly - ve VELKÝCH rozsáhlých systémech

→ eventální konzistence

- entropie - když má uzel nějakou informaci, tak ji pošle i na nějaký další

push - funguje dobře na začátku, když jich je ještě málo nabitých
pull - uzly si o ty informace samy vlekají → na začátku pomalé rozšiřování
push+pull

- kdy ale přestat infikovat nové uzly?

→ gossip protokol - když uzel natrefí na nějaký už infikovaný, tak se s ním jakou pravidel.

- problém mazání dat - jak zrušit nějakou informaci

• naivní - prostě smažeme, ale z jiných uzlů se ta informace zase vrátí

• certifikát smrti - nová informace o zrušení ke staré

- trénování much

- spočítání, jedna výplata, která se pak replikuje, počítají si průměry
- největší mucha
- celková váha

Správa prostředků a procesů

→ krátkice ve Wait-for grafu (WFG)

Detekce deadlocků

- korektnost $\Leftrightarrow$ 1) $\nexists$ deadlock v kon. čase detekován
2) $\nexists$ det. deadlock opravdu $\exists$

Centralizovaný alg.

- do jednoho uzlu se reportují události • změnách WFG
- je potřeba mít ^{strukturu} prot. pro doručování zpráv $\Rightarrow$ jinak falešná detekce
nebo při podezření na deadlock nejdrův kontrola

Path-pushing

- uzel si udržuje část WFG, která se ho týká
- vyměňují si externí závislosti + kandidáty na zabití
- průběžně se kontroluje, jestli není deadlock $\rightarrow$ kdyžtak se pošle informace tomuto uzlu, kde je ten

nejčastěji nejmladší proces
(s nejvyšší hodnotou log. hodin)

Edge-chasing

- podle WFG se roznese speciální značka, ve větveních
- $\rightarrow$ když se dostane zpátky k iniciátorovi, tak on usoudí, že došlo k deadlocku

po cestě sbírá kandidáty na zabití

se replikuje na větve

(viz dříve)

Diffusing computation - podobně jako alg. na detekci zastavení výpočtu

- distrib. výpočet po WFG
- v praxi moc ne

Detekce glob. stavu

- "deadlock na nás počká" - dokud nerozbijeme tu krátkici
- najde se kauz. konz. řez

Distrib. správa procesů

- dělení dle toho, kde jsou ty procesy spuštěny:

Load balancing / task scheduling

- replikovaný master (např. pomocí Paxosu)
- round robin $\rightarrow$ rozděluje ty úlohy
- na krátké / srovnatelně náročné úlohy
- jednoduché ☺

služby Klientům

• krátkodobé úlohy $\Rightarrow$ chceme min. latenci

• dlouhodobé $\Rightarrow$ max. výkon

distrib. výpočty $\rightarrow$ chceme rovnoměrné rozložení zátěže

- weighted round robin - např. $\frac{1}{2}$ serverových jsou starší stroje, zbytek novější $\Rightarrow$ vyšší váha
- dynamic round robin - měří se, jak moc jsou uzly vytíženy - pravidelně posílají - a podle toho se master rozděluje
- least connections - ^{požadavek} dynamic - ale master si jenom proská paruje, kdo má kolik spuštěných procesů
- random, threshold - možnost, že uzly sami odmlknou, když už toho mají moc

- agent-based, ..., AI - jenom aby ^{spíš} dědci měli o čem psát

Koop. systémy - uzly po celém světě, volně organizovány

- když uzel nemá co dělat, tak si řekne centrálnímu registru o job ↑ anotační hlava, co třeba potřebují a tak..

Koop. load balancing

- distrib. heuristický alg.
- det. graf. alg. - data-flow driven computing
→ toky v sítích
- bidding alg. - procesory nabízejí výp. sílu a procesy si ji kupují
- nepoutivá se
- centralizované řešení - UP-down alg.
- koordinátor si eviduje uzel trestné body ↑ procesy na jím uzlu + nespokojení p. řádky -
→ když se uvolní kap. uzlu, tak se vyberou ty s nejmenším # trestných bodů ↑ počítačů z uzlu

Migrace procesů / virt. strojů → přenos, ^{konkrétní, transparentní} procesů během výpočtu ↳ dnes kontejnery a virt. stroje

MOSIX - distrib. OS

- byl na něm implementován vekt. alg. pro řízení zátěže

T4

- 2 MFF

Migrační load balancing

Párový alg.

- náhodně se vytváří páry, konstruují se seznamy svých vybraných procesů a počítá se, o koho by se zlepšilo za zatížení, když se nějaké procesy přenesly na druhý uzel z páru

Vektorový alg.

- uzel si eviduje vektor s vytížením uzlu ↑ vekt. menší než # uzlu
↳ 0. prvek je vlastní vytížení
- v nějaké periodě si uzel vybere náhodně # uzlu odpovídajících polce vektoru a od nich si vytváří jejich horní polovinu vektoru ↑ zbytek zátěže
- v případě potřeby vyslání procesů se uzel podívá na ten vektor ↑ z ipem spojit

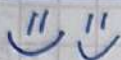
Bidding alg. a další

- FUJ

Centralizovaný / hierarchický alg.



Lok. rand. alg.



Distrib. dat. struktury

CRDT - Conflict-free Replicated Data Type

✓ chceme

- updates se asynchronně přenášejí mezi replikami

→ při konfliktech se udělá merge

→ konvergence ke stejnému výsledku ve všech replikách

State-based replication

- lokální příprava + async. posílání stavu ostatním - gossip prot.

- merge: komutativní, asoc., idempotentní, monotónní vzhledem k čas. usp.

DF: kauzální historie C replik x_i objektu x :

1) na začátku: $C(x_i) = \emptyset$

2) po op. update - F : $C(F(x_i)) = C(x_i) \cup \{F\}$

3) po op. merge stavů x_i, x_j : $C(\text{merge}(x_i, x_j)) = C(x_i) \cup C(x_j)$

Operation-based objekty

- nepřenáší se stav, ale ty operace - není potřeba žádná merge operace

- operace jsou kom. a asoc., ne idemp.

- doručování v kauz. usp. !, je potřeba spolehlivý broadcast

Konvergence (eventual convergence)

$C(x_i) = C(x_j) \Rightarrow$ stavy i a j jsou ekvivalentní

$f \in C(x_i) \Rightarrow$ někdy nastane i $f \in C(x_j)$

Last-Writer-Wins registr

- každý update má svůj timestamp a ID

Množiny

- add a remove nejsou

komutativní ☹

Řešení: pomocný (ale podle potřeby znovu vložit)
- prvky + ID