

SIMD

- více ALU $\hookrightarrow$ levně
- synchronizace $\hookrightarrow$ dražší
- samotný výpočet $\sim 10\%$ energie, ta zbyte $\sim 90\%$ na CPU (20% na GPU)
- vekt. instrukce (viz v. 1974)
- vekt. registry (256 - 512 b) většinou

Vekt. registry

- nebyly povahové - to, co tam je, se interpretuje podle té instrukce, která se na tom udělá
- lze použít i jako takovou další úroveň cache

Vekt. instrukce

- prostě normální, ale operandy jsou velké (vezmi 2 reg., něco s nimi udělej, a ulož to do jiného reg.)
- N matematických op. najednou
- $\hookrightarrow$ N lajn (lanes), (nezávislých) (ale nemusí) dokonce většina instrukcí z instr. sady
- dnes podpora pro neuronové sítě (např. 16 b floaty)
- např. reálné přesuny bitů, Kryptograf. operace

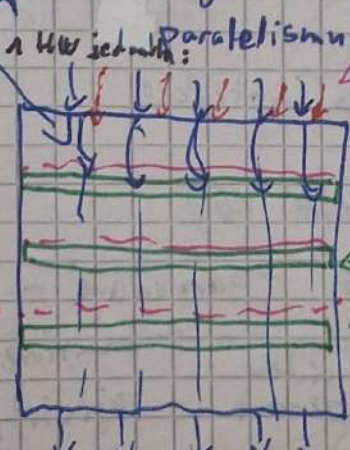
hezké instrukce:

- N HW jednotek

1 HW jednotka

K HW jednotek $\rightarrow \frac{N}{K}$ instrukcí

cesty, různé dlouhé

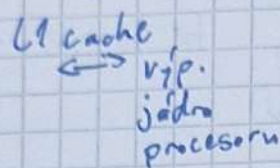


1 instr. tam nape ze sebou ty operandy $\rightarrow$ jednodušší manipulace s daty

at se to dělá na jednom patře, tak už se tam můžeme dít dále

Paměťové přesuny

- kvůli vekt. instr. se musela zvětšit tloušťka těch kanálů (jedině ty ji dokážou celou využít → další důvod je používat)

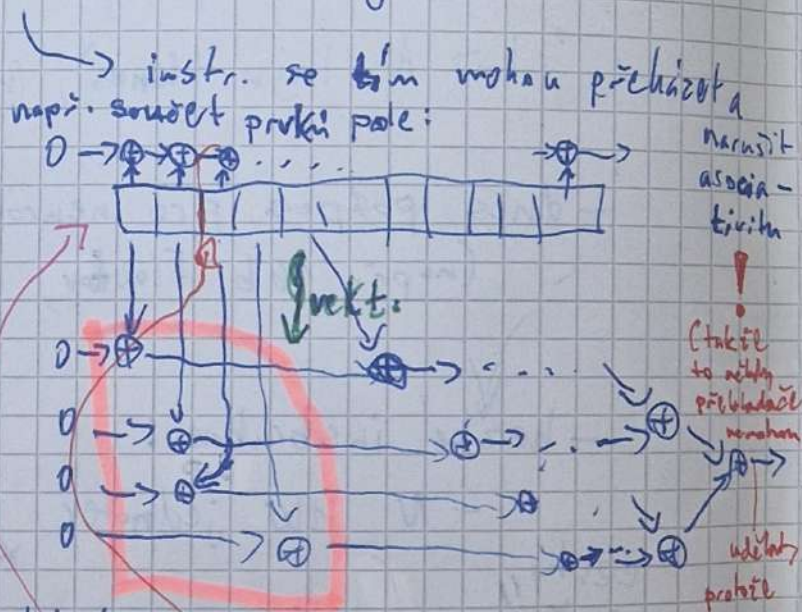


např.
memopy -
ani ne pro k
vekt. aritmetiku
ale kvůli
těm
přesunům

SV support

- vektorizace překladači
- cyklus $\rightsquigarrow$ spojí se N po sobě jdoucích iterací do jedné

- polyhedrální kompilace -
prohazování
zaměřenosti
cyklů...



- vektorizovaný knihovní kód
- explicitní vekt. dat. typy a instrukce
- zarovnání (na adresy dělitelné 16)
 - třeba dostaneme float 4 - zarovnané na 4

u každého si ten
můžeme posunout

- musíme z té adresy
vjistit, jestli i na
16

např.
u floatu

SIMD na Intel/AMD CPU

MMO - MM7

1997 - Intel využil existující FP registry (80-bit) na intové operace (8, 16, 32-bit)
- AMD i 32-bit Floaty

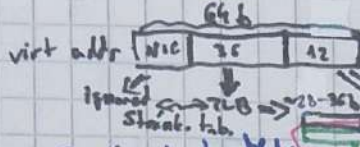
1999 - SSE - 128-bit registry XMM0-7

2001-2007 - SSE2-SSE4 - 64-bit Floaty a 8/16/32/64-bit inty

2003 - (AMD 64) XMM8-15

2011 - Intel + AMD

AVX
extensions
used



128-bit. vektory

512 = 3 * 64 / 26
FP opce

64-bit procesor

bitů ALU
- 8(16) b v 8080
adresy

Velikost adresy
- 16 b na 8080
- 24 b na 8086
- 32 b na 8088

adresovatelnost paměti
- 16 b na 8080
- 20 b na 8086
- 32 b na 8088

Velikost adresy
- 16 b na 8080
- 20 b na 8086
- 32 b na 8088

Velikost fyz. adresy
- 20 b na 8086
- 32 b na 8088
- 34-42 na AMD64

2013 - AVX2 - +intová aritmetika

gather / maskstore instrukce

načítání
i nesouvislých
úseků z paměti
zapisuje se
do souv. úseků
ale jen
nevyužívané
části

2012 - (Intel) XMM

32 registry ZMM0-31, 60 jader

- používáno jako Coprocessor

- ale FUJ

2016 - (Intel) AVX512

32 registry, 512-bit
jako hlavní procesory

Vekt. instr.

→ #include <immintrin.h>

- přístupy do paměti
- arit. instr.
- posloup.
- inter-lane aritmetika
- inter-lane načítání
- konverze (typy, skalár ↔ vektor)

sub r1, r2
latence = doba, kterou
se musí
počkat na
tu instrukci
(inverzní)
Propustnost
CPI
kolikrát
je na
odstředivém
(např. 0.5 → 2 instr. za takt)

r1 = r2
ak překladač
umí dělat iluzi
takhle: r3 = r1 - r2
(a stejný princip
je i u vektor-
ných instr.)
umí se odstraňovat
části instr.
za takt

Přístup do paměti

společná mechanisme na přístup k paměti

- load/store
 - loadu/storeu - i nezarovnané operandy
 - stream-load - obchází cache (většinou platí jen pro L1)
- + jiné vekt. instr. mohou mít jeden operand z paměti

souvislé úseky

adr. nedělitelné 16 B

Saturované operace = když by to mělo přelst, tak to uděle jeho ziskem na min/max

dues to un? všechny instr. kromě loadu, ale někdy to je stejné jako load, když to není zarovnané, tak je to pomalejší

Datové typy

- třeba mohou být FP vekt. reg. a jiné intové reg.

+ i nějaké zkratky, že to jde třeba rovnou do ALU

bez zastávky v registrech

- ≥ 3 typy loadu:
- intové
 - single prec. floaty
 - double prec. floaty

nesouvislé úseky: pomalejší než souvislé!

- gather
- scatter (jen AVX512)

ale když to jinak vejde, tak je to rychlejší, než to dělat nějak skalárně

- další vstupní vekt. reg. jsou určeny adresy

gather: for (i in 0...N-1): $v[i] = a[c * x[i]]$ 1/2/3 podle velikosti toho elementu

scatter: For (i in 0...N-1): $if(m[i]) a[c * x[i]] = v[i]$

Podmínky

For (...)

if (...)

$x[i] = f(...)$

else

$x[i] = g(...)$

a na závěr se udělá podmíněné přiřazení
musí se udělat obě

- v SSE a AUX:

$$e[i] = a[i] == b[i] \approx c[i] : d[i]$$

$\Rightarrow \text{cond} = \text{cmpeq}(a, b)$ $\begin{cases} \text{true} : \text{same} \ 1 \rightsquigarrow -1 \\ \text{false} : \text{same} \ 0 \rightsquigarrow 0 \end{cases}$

$$\text{left} = \text{and}(\text{cond}, c)$$

$$\text{right} = \text{andnot}(\text{cond}, d)$$

$$e = \text{or}(\text{left}, \text{right})$$

↓
pozor, vlastně
je to notand

$\Rightarrow \begin{matrix} \text{cond} \\ 1 \rightarrow c \\ 0 \rightarrow 0 \end{matrix}$

$$\text{left} = \text{cond} \& c$$

$\Rightarrow \begin{matrix} \text{cond} \\ 1 \rightarrow 0 \\ 0 \rightarrow d \end{matrix}$

$$\text{right} = \text{ncond} \& d$$

- v AVX512:

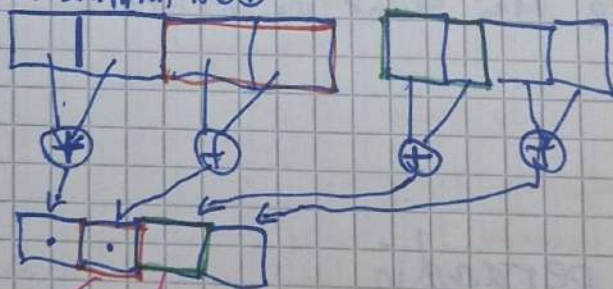
- maskovací registry $\leftarrow$

- porovnávací op. do nich dávají tv výsledky (1 bit/lajna)

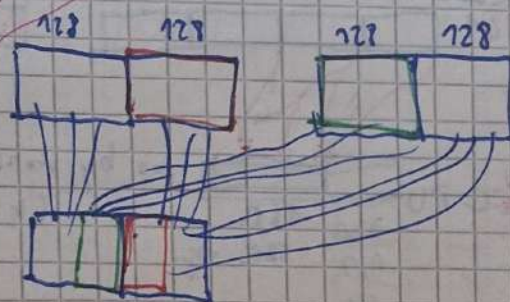
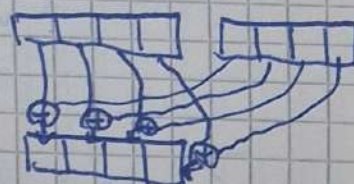
$\hookrightarrow$ pak se v další instr. použije ta výsledná maska (+ dá se jí alt. hodnota, která se tam má dát, když ta maska je 0)

Inter-lane aritmetika

Horizontální ADD:



$\times$ normální:



opět - FUJ! POZOR!

Inter-lane shuffles

unární: $r[i] = a[p[i]]$

binární: $r[i] = (p[i] \ll TOP_BITS) \mid b[p[i] \ll LOW_BITS]$

• permute, shuffle, align

bit. posun
směrem k
nižším bitům



provo je nalevo
(je to LE)

hodí se třeba
pro čtení nezarovnaných
hodnot, je-li
načteme zarovnaně
a posuneme si je
(vystřihneme si
z nich tu
užitečnou část)

- extract - vekt. reg. $\rightarrow$ skalár/menší vekt. reg.
- insert - opačný extract
- broadcast - rozkopíruje daný skalár do vektoru

- varianty:
 - set1
 - setzero
 - cast

Zarovnaní paměť. operandů

SSE: nezarovnaný $\rightarrow$ segfault
na 16B

AVX: nezarovnaný $\rightarrow$ pomalejší
na 32B



číslový
malloc
zarovnaný
na 8
a dočetl
nemohl
new předat
informaci o tom
typu té alokaci
málo, takže tam to
bylo na 8B

klíč. re.
tak můžeme
použít
alignas(16)

ještě klíč. to / tak
ještě zarovnaný
ale klíč. je
to na konci
tak je to
víc lepší

to by mohlo v těch
64B byt
hbitativ
bylo by na
to potřeba
hodit další
ještě
překladač
udělá to
za nás
tak se to ještě
neudělá

alignas

- `struct alignas(16) aligned { int32_t a[4]; };` ← vlastnost typu
- `alignas(16) int32_t v1[8];` ← vlastnost proměnné

nejde to dát třeba do vektoru jako `std::vector<alignas(16) int32_t>`

↳ tam se mus. používat ty speciální typy — `m256`.

⇒ když pak chceme používat skalární, tak se mus. použít `reinterpret_cast`

- číselná — na data nějakého typu se začneme dívat jako na data jiného typu

(`v.data()`) — to nám

dá — `m256`

→ to dáme

do `reinterpret_cast<int32_t*>()`

— např. pro zapsání do binárního formátu —

- může se tím porušit záruka toho, že když se přistupuje k různým typům, tak nemohou ležet na stejném místě:

místo: `double *p = ..`
`int *q = ..`

↗ překladač
↘ může probíhat

→ ale `reinterpret_cast`em to rozbijeme — může

nám takhle prokazovat naše skalární a mat. instr.

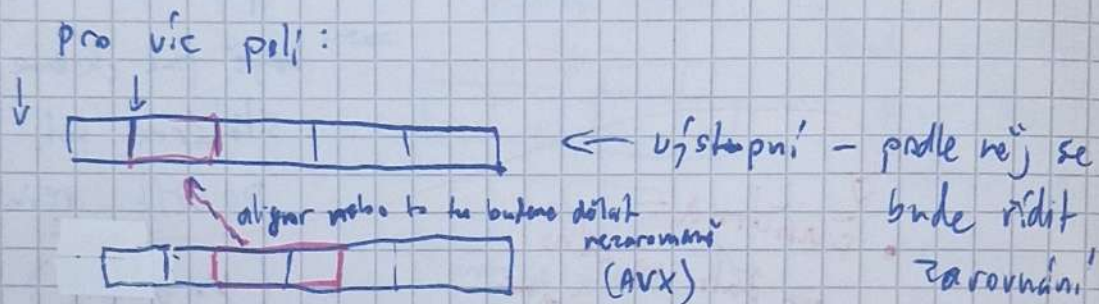
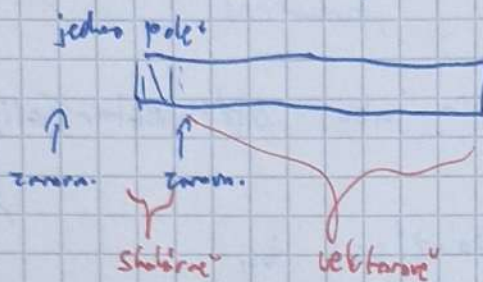
- v C++23 to řeší

`std::launder` — chová se jako fence — nedá se přes něj prokazovat !!

našleš! ale překladač to umíle rozumí a funguje to

Zarovnění za runtime

- pointer reinterpret casteme na `std::int_ptr_t` a uděláme % 16



Vektorizace kompilátorem

- nejvnitřnější cykly s přehledným chováním

```
for (int i = 0; i < N; ++i)
```

```
    a[i] = b[i] + c[i];
```

překladní dělá rovnou takhle

} loop unrolling - ekv. úprava

```
for (i = 0; i + 3 < N; i += 4) {
```

```
    a[i] = b[i] + c[i];
```

```
    a[i+1] = ...
```

```
    a[i+3] = ...
```

```
}
```

```
for (; i < N; ++i)
```

```
    a[i] = b[i] + c[i];
```

```
for① (i = 0; i + 3 < N; i += 4) {
```

```
    __mm_store_ps(a+i,
```

```
    __mm_add_ps(
```

```
        __mm_loadu_ps(b+i),
```

```
        __mm_loadu_ps(c+i)))
```

vektizace
nemusela by být ekv. úprava

```
for② (j; j < N; ++j)
```

kdyby se to a, b, c a[i] = b[i] + c[i];
vždy překrývaly! => ne překladač musí zarovnat ekvivalenci

- podmínky autom. vektorizace:

- predikovatelná podmínka

např. $i+3 < N \rightarrow i < N, i+1 < N \dots$
musí se vědět

- počítatelné smyčky

$\rightarrow$ # iterací lze za běhu určit podle té podmínky
 $\rightarrow$ např. tam nemusí být nějaký if a ta řídicí proměnná se nemusí nijak měnit v tom těle

- nemusí tam být závislosti mezi iteracemi:

$$S = S + a[i]$$

- každý další iterační
závisí na předchozím S

$$a[i+1] = a[i] + 1$$

- závislost přes
elementy toho
pole

$$a = b + 1$$

$$a[i] = b[i] + c[i]$$

- aliasing

- ale my chceme:

- vyhledávat v poli a skončit, když najdeme
 $\rightarrow$ nepředvídatelné smyčky

- akumulátor

$\rightarrow$ závislosti mezi iteracemi smyčky

- vygenerování výstupního pole

$\rightarrow$ může tam být aliasing
vstupního a výstupního pole

teď měř
všechny
účinné
smyčky budou
mít nějakou
překážku

da tam
na to testy

```
if((a <= b || a > b+3)
    || (a <= c || a > c+3))
```

for(1)

for(2)

takže by takle
mohlo vypadat

1.2.6: Floyd-Warshall

konstr.: alokace + zaroveň clear()

velikosti matic budou delitelny 64 - nemí potřeba řešit zkrácení

násobení matic: $c[i][j] = \min(c[i][j], a[i][k] + b[k][j])$

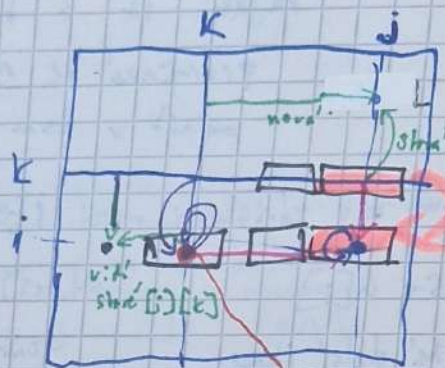
F-W

min

c

+ c

-> aliasing



$$i = k \Rightarrow c[k][j] = \min(c[k][j], c[i][k] + c[k][j])$$

pred vstupem

do j-čkového cyklu

se da' instr. broadcast na

rozumnosti $c[i][k]$

1) základní řešení:

vektorizované

ještě ten

najmenší

cyklus - ten j-čkový

$$\Rightarrow i \cdot j \cdot (\text{add} + \text{min} + 2 \cdot \text{load} + \text{store})$$

-> to přičteme

=> minimum s $c[k][j]$

=> výsledek
zapsáno

2) budeme na jednom

čísle několikrát (4):

$$\Rightarrow \frac{1}{4} \cdot j \cdot (4 \cdot \text{add} + 4 \cdot \text{min} + 5 \cdot \text{load} + 4 \cdot \text{store})$$

experimentovat (norma 2)

$c[i][k]$ - broadcast

$c[i+1][k]$ - broadcast

to může být

norma 2

ještě ten

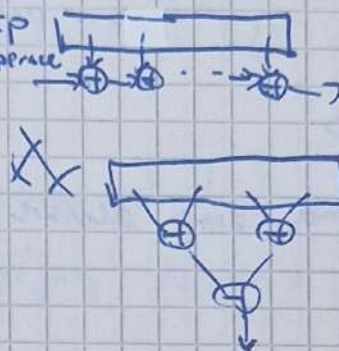
cyklus - ten j-čkový

32 elementů
na AVR512

C++ rozšíření pro optimalizace

- ekviv. úpravy, které by ale překladač rozhodl jako nezitečné
- neekviv. úpravy, které by překladač nesměl dle pravidel jazyka udělat
 - my tu správnost nenarušíme
 - narušíme

na pří: FP operace



tohle by
překladač
sám
nemohl
udělat

Inlining

- " => odstranění overheadu instrukcí call + ret
- ✓ => větší hřiště pro překladač
 - ví, co se tam děje
 - meziprocedurální optimalizace

- " => delší překlad
- => větší velikost kódu

Intel

#pragma inline/forceinline/noinline ← k volání

GCC

void f() —attribute— ((always_inline/flatten/noinline)) ← k funkci

↓
všechna
volání
se musí
inlineovat

(a kdyžtak
kompil. chyba)

↓
všechna
volání z této funkce
chtějí inlineovat
(pokud to jde)

32 bitů
AVX512

Vynucení vektorizace

Intel

#pragma vector always

#pragma simd

← před smyčkou

✓ zvektorizuje jen pokud je to bezpečné

OpenMP

#pragma omp declare simd

GCC

void f() __attribute__((simd))

```
for( ... ) {  
    a[i] = f(b[i]);  
}
```

nešlo by normálně zvektorizovat

1) inline
2) zvektorizovat

NEBO

se vytvoří další verze té funkce, která bude mít vektorové argumenty a vst. hodn.

```
for( ... i+=4 ) {  
    a[i...i+3] = f4(b[i...i+3]);  
}
```

Ignorování potenciálních závislostí

Intel

→ ignore vector dependencies

#pragma ivdep

GCC

#pragma GCC ivdep

ale FUJ

ale příklad je jen různě číselný
řekne, že má to normálně by závislosti o kterých si není jistý

```
void f(int *a, int k, int c, int m) {
```

#pragma ivdep

```
for(int i = 0; i < m; i++)
```

```
    a[i] = a[i+k] * c;
```

ale tímhle garantujeme, že je to v pořádku
kdyby k bylo záporné, tak to vadí

SolarisCC (Oracle)

#pragma noalias (pointer, pointer [, pointer]...)

#pragma may-not-point-to (pointer, variable [, variable]...)

restrict - C99

nebo někde -- restrict --

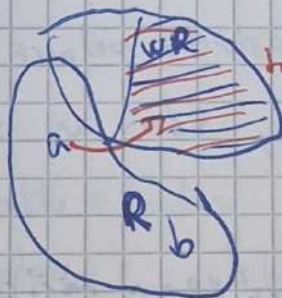
```
void copy(int * restrict a, const int * b, int m){  
    for(int i=0; i<m; i++) a[i] = b[i];  
}
```

~~pro~~ všechny paměťové přístupy se mají vystopovat k přístupu na nějaký pointer + offset

když se použije restrict pointer na a, tak se zavazuje, že ke všemu, na co přistupujeme přes ten pointer, se pak bude přistupovat už jenom přes ten pointer a

pro proměnné json škatulky:

- if restrict proměnná má vlastní
- ostatní mají jednu společnou

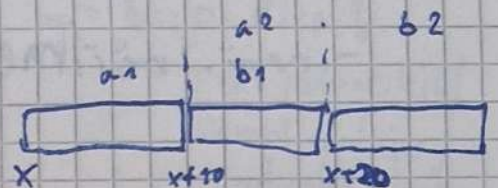


takle se zapisuje jenom přes a

přístupy jsou prokazovat, pokud jsou dotčeny proměnné v různých škatulkách

```
int m[100];  
int * restrict x = m;
```

- copy(x, x+10, 10);
- copy(x+10, x+20, 10);



ACE:

- copy(x, x+10, 20);

```
{ int * restrict y = x;
```

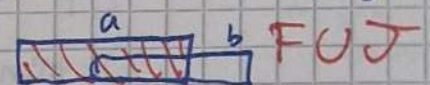
```
copy(y, x+10, 10);
```

```
copy(y+10, x+20, 10);
```

z pohledu překladače

ty copy zapisují na různé místa v paměti

můžeme ověřit, že jsme přestali zapisovat přes x a zapisujeme přes y



povolili jsme interleave, tedy voláme copy!

Ladění rychlosti programu

Problematický koncept (hotspot) - místo, kde má cenu optimalizovat

Problémy: • je těžké odhadnout to úsilí na tu optimalizaci
• je těžké odhadnout efekt ke optimalizaci

tedy: hotspot je místo, kde celkový čas vykonávání dělený velikostí toho kódu je největší

chceme zahrnout i ten čas procedur, které jsou odsud volány?

Instrumentace ($\Rightarrow$ není potřeba podpora HW ani OS " ")

$\rightarrow$ modifikace programu, aby sám u sebe něco změřil
 $\hookrightarrow$ ta ale ten chod

ovlivní $\Rightarrow$ nemírá cenu měřit čas, ale počty průchodů nějakými důležitými místy

- např. měříme, kolikrát a odkud byla volána která procedura

a z toho se pak odvodí i další místa

- optimalizace řízené profilem

1) pustí se instrumentovaný program (na reprezentativních datech)

2) ten výsledek profilování se předhodí dalšímu překladu a podle toho se optimalizuje

např. na těch hotpotech se pustí plně řešit nějakého NP

(a jinde jenom úprava)

úplného problému

Sampling

⇒ vzniká se konkrétně na normálním běh toho programu

↓ podle toho, jak někdo přišel

způsobem:

- podle IP
- podle call stacku

- v náhodných okamžicích
→ nezdr. na běh programu
- podle nějakých událostí

• softwarové samlování

- zmanipuluje se časovač, je potřeba podpora OS
(OS používá jednu z hodin, které tikají jednou za 10 ms)

↓
posílá
přeruš.

↓
řekneme, že to
bude častěji
použijeme jiné
hodiny

- pozor, takže se budou

sampleovat i ostatní procesy a kernel

⇒ bezp. díra!

• hardwarové samlování

- CPU udělá po daném počtu datových událostí přeruš.

- pro pokročilejší samlování je potřeba být root

↳ • tikání hodin - takový procesor !!

- #instrukcí (ale pozor - CPU spouští spoustu instr., ale jen některé doběhnou do konce - počítat všechny, nebo jen dokončené? → Intel)
- přístupy do paměti
- špatná predikce skoků

budeme si to nechat
překontrolovat

z toho
si chceme
trackovat
jenom
pár
nejednou

(a jinak jenom
čistě)

úplně problém

Analýza závislosti

→ zjištění, že nějaká instrukce musí být puštěna až po jiné

(+ latence - minimální doba, která mezi nimi musí být)

→ část. usp. na instr.

- HW nějaké záv. může ignorovat: - např. když jsou rejstříky

(→ a pak třeba zjistí, že není → pusť znovu)

např.

- závislosti:

• datové - Write-Read

- typický latence - read musí počkat, až to to první spočítá a zapíše

• anti dependence

- kdybychom to prohodili, tak dostaneme nějaký jiný stav

• Read-Write

• Write-Write

Samostatným instrukcím to nevadí, ale nám ano

• kontrolní

- je tam instr., která může způsobit nějaký vážný problém

- např. overflow výjimka

Kontr. záv.

while (p && p->v != x) p = p->next

WR

• započítá nějaké
kde bychom se už k
přirozené hodnotě
nedostali
• sáhnout
do paměti
která není

2.D.Ú. Bsearch

bsearch-inner - Read only dat. struktura

bsearch-outer - Voláno vlastním, RV

→ pole, ve kterém vyhledáváme, data už v konst.

konst.: osize = kolik těch hodnot bude dle zároven vyhledávat (N)

dostaneme seřazená, my si to třeba můžeme doplnit na velikost mocnin dvojky

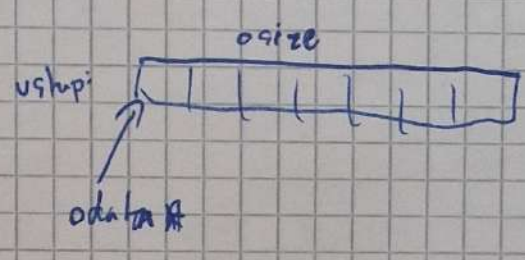
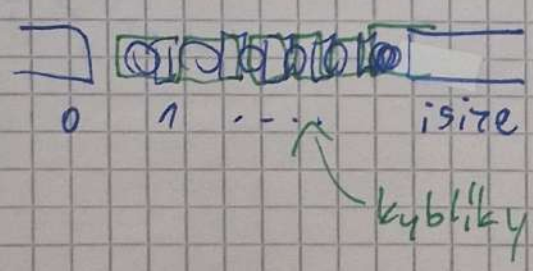
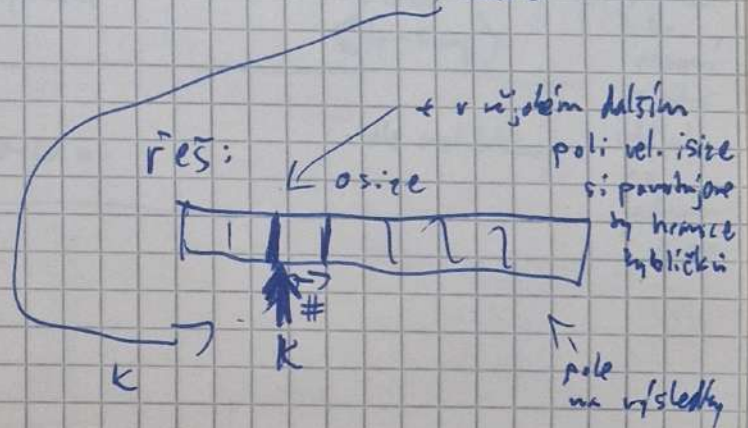
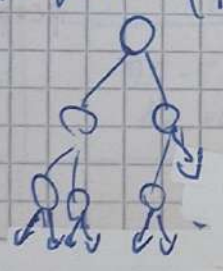
naalokujeme si buffery, nainicializujeme dat. strukt.

není měřeno

bucketize se pak bude několikrát volat s daty vel. osize

výsledek si uchováme a vracíme ve funkci bucket

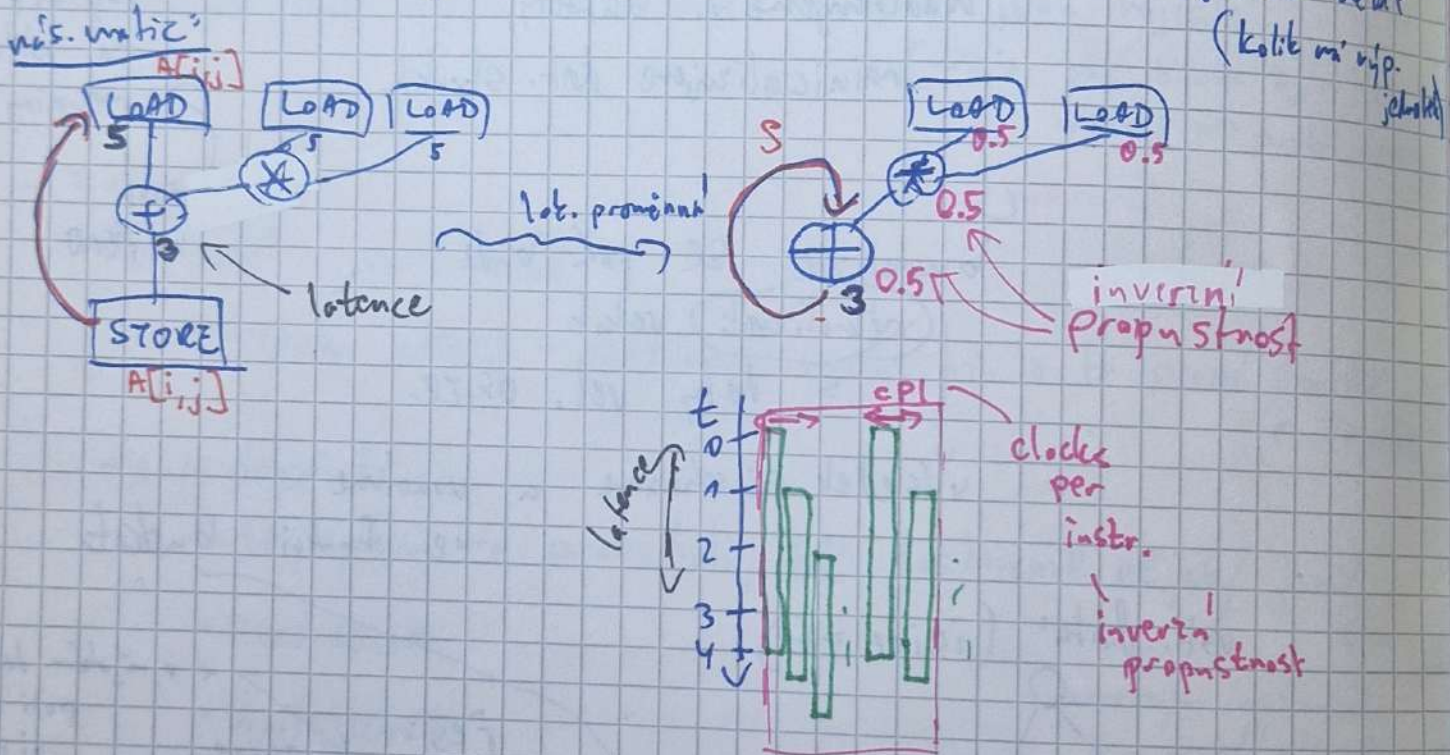
vst. data: (isize = 6)



čas: $osize * \log_2(isize)$
trošičku rektorizovat

Optimalizace vnorených smyček

- změna zanořenosti
- využití algebry (např. že XOR je kom. a asoc.)
- polyhedrální kompilace - řeší se vztahy mezi datovým a iteracním prostorem
- brzdi nás:
 - latence operací
 - # operací, které ten procesor může najednou dělat (kolik má výp. jednotek)



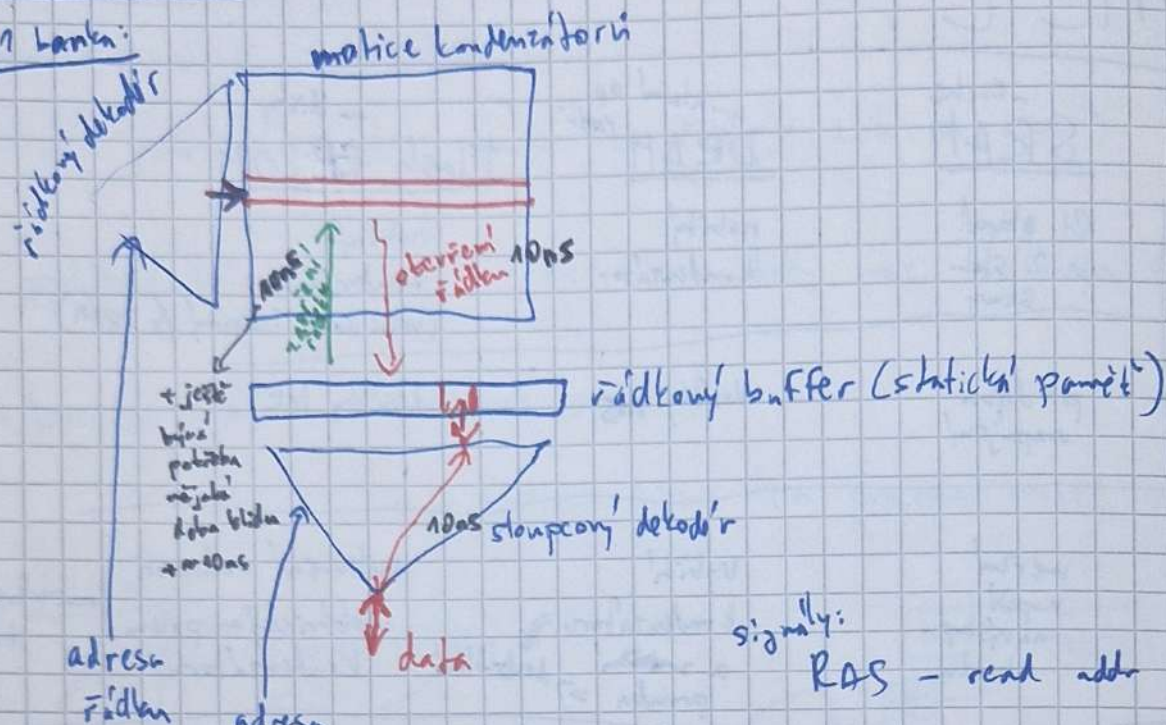
Paměti

	- cache <u>SRAM</u>	- hlavní op. paměť <u>DRAM</u>	- disky <u>Flash EPRAM</u>
princip základu	el. obvod se 2 slab. stavy	nabitý kondenzátor	nabitý kondenzátor * (velmi dobře izolovaný / X DEAN)
výdrž základu	po dobu napájení	desítky ms	desítky let
princip čtení	měření napětí na výstupu obvodu	vybití kondenzátoru a změna proudu → destruktivní	ovlivnění vodivosti oleobrickým polem kondenzátoru → bezstrukt. ture
princip zápisu	změna napětí vstupní obvodu	nabití kondenzátoru	nabití/vybití kondenzátoru - takže je na tom (nap. tunelování) to technicky náročné
transistorů na bit	6	1	1/3
výrobci	USA	Japonsko	Korea (Samsung) → nikdo to nemohl předtím udělat na jednom čipu
latence	0.5-5 ns dle velikosti	10-20 ns	dle architektury

* nos
ex e m
t i n
i d e conductor

DRAM

1 banka:



signály:

RAS - read addr select

CAS - column addr select

multiple-
xvstho
(CAS
je
potřeba
at po
RAS)

VDDR3 (PC3)
jsem 2 moduly

modul

2 ranky

rank

tabule
celé x ~ 8

a nad
tím je ještě
jedno patro

tabule
s tím
nad tím



kanály (nezávislé,
=> větší propustnost)

V DDR5:

- modul přes
dva kanály (= subkanály)
=> větší paralelismus

DDR4 - 2400

=> hodiny = 1.2 GHz

64b dat
x 16
Cpř. 16
Víc čipů se společnými
tými signály

=> SDRAM

napi.
dalšíme
adr.
řádku

po 15 tabulkách
musíme
dodat
adr. sloupce

alemní
notiček
(sest.)

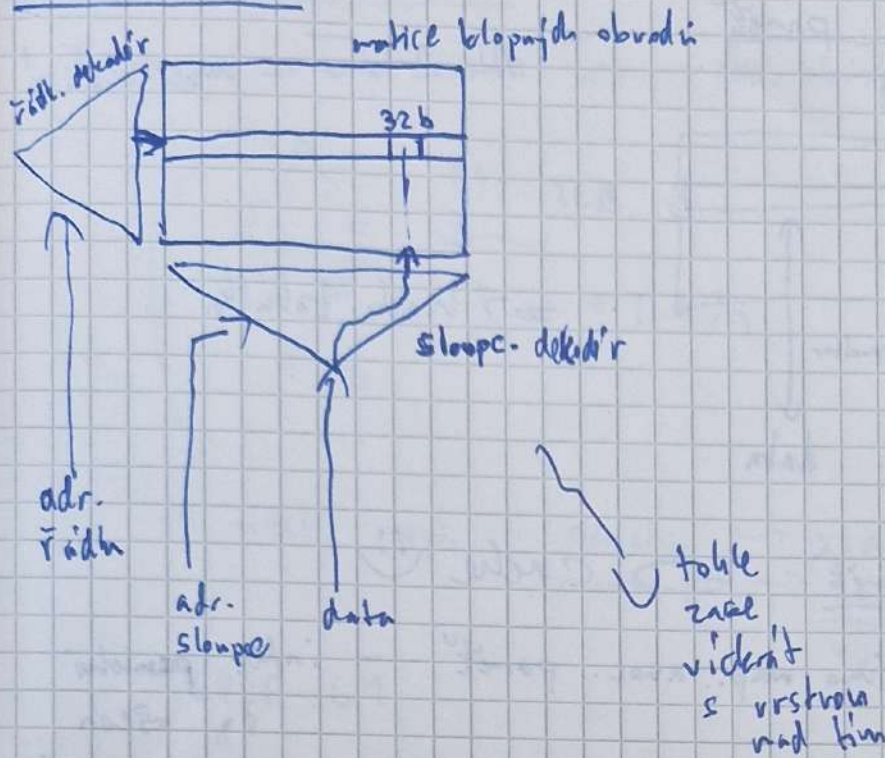
v náhlíkách
v baticích
po 4/8

Po dalších
15 tabulkách
dostaneme ta data

(v DDR5 po 16)

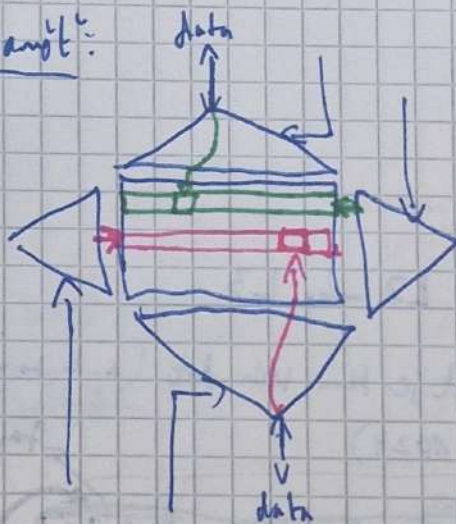
mít
se toho
dát
více
nejednou
(1 banka
množství,
tak můžeme
řek jím
at dít
něco dal-
šího...)

SRAM



Dvoubránová paměť:

(v TLB)
(a jiných malých pamětech)



=> Paralelně dva řádky najednou!!

(jinými než CPU)

Přes

Asociativní paměť

(např. TLB) ^{nebo prediktory skoků}
(~ nejvýše desítky řádků)

= SRAM rozdělena tak, že nějaká pam. jednotka (~ stovky bítů) bude doplněna komparátorem

-> zvrhne záznam

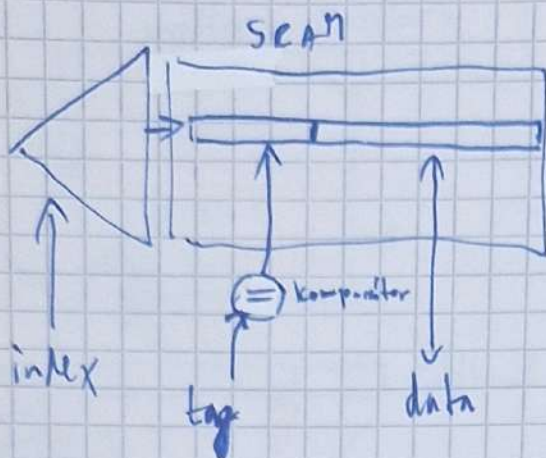
tag, ty komparátory

to paralelně

porovnají, a ten, kde to pasuje, nám dá data



Primo mapovani asoc. pameti



Ale FUJ - moc kolizi!

=> hash, tabulka

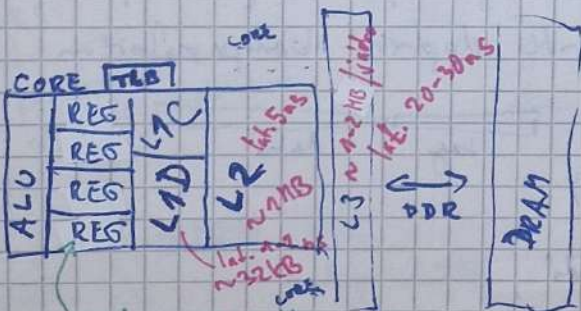
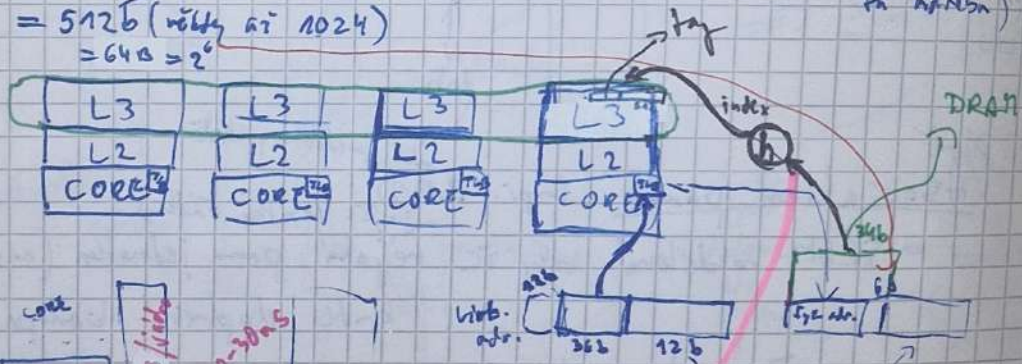
N-cestni asoc. pameti => cache

=> n-krat primo map. asoc. pameti - index pameti do vsech a ta, u kteri se tag, se pouzije

(casto index je nejaky hashem tagu)

LLC (last level cache) - L2 nebo L3

- 1-2 MB na jadro
- cache line = 512b (vety az 1024) = 64B = 2⁶



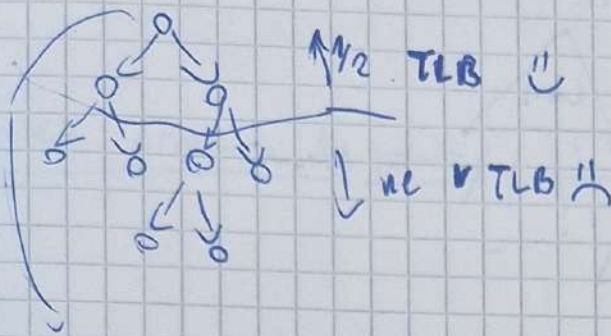
saky registru i hyperthreadingove vladno

- inkluzivita = kdyz je to v jedne urovni, tak je to i v tech nizsich
- dnes L1 inkl. s L2, ale L2 ne inkl. s L3

- jonom se z toho konck ukousne => index (konkrtni ridky)
=> porovna se ty
=> tim mene k en ridky

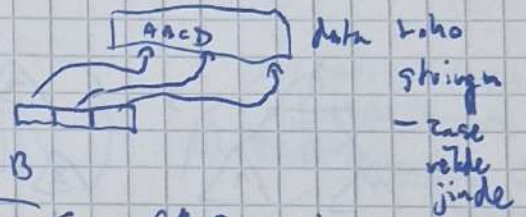
Stromy v cacheích

- TLB má ~1000 záznamů $\Rightarrow$ ~1000 stránek efektivně



1 uzel např.

static string $\rightarrow$ 24 B
2x unique ptr $\rightarrow$ 2 * 8 B



40 B < 64 B řádek

64 kB L1

11
1024 řádek

+ to nemusí být rovnou na řádky

$\Rightarrow$ takže to možná může být i přes dva řádky $\Rightarrow$ 128 čtení

32-bit čísla (31B) v poli:

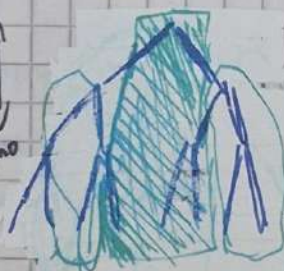


64 B = 10 * 4 B

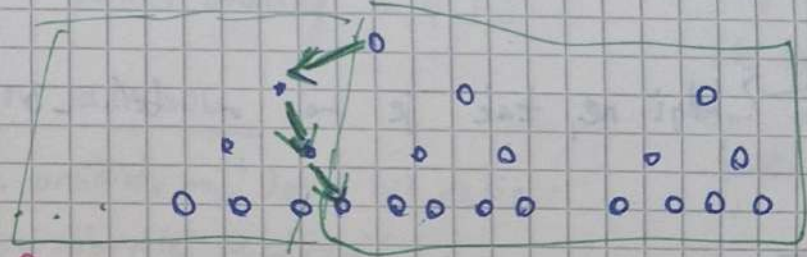
velikost bloku

když uzel zabírá 2-4 řádky!

#řádů cache
 $\log_2 C$
trvale mechanismus

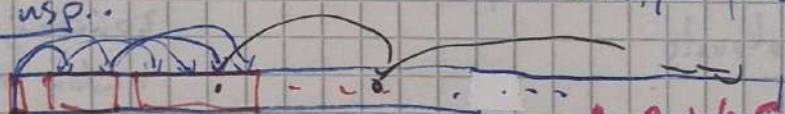


$\log_2 B$



cache nám pomáhají až takhle v těch hlubokých patrech dole

jiné usp.:



$\log B + \log C$
 $\log(B * C)$ trvale v cache

$\Rightarrow$ výhodnější pro paralelní procházení

řádů

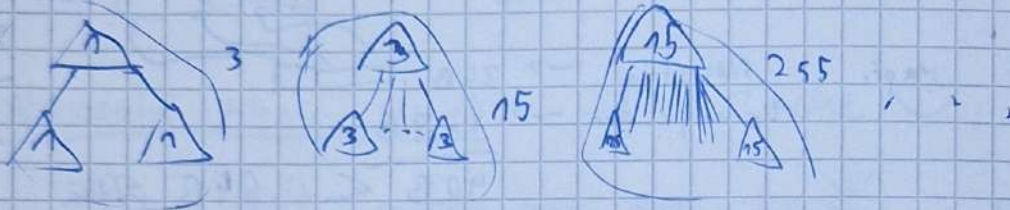
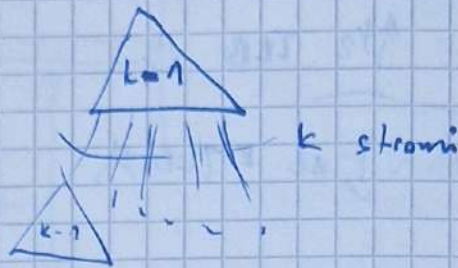
en řádek

van Emde Boasovy stromy

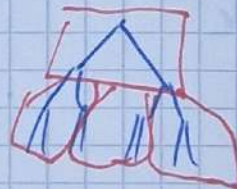
k přepážek v prostoru klíčů

strom o n uzlech = $(k+1)$ stromů po $(k-1)$ uzlech

$$(k^2 - 1 = n) \\ \Rightarrow k = \sqrt{n+1}$$



pro 15 prvků:



blízko ve stromě
 $\Rightarrow$ blízko
v paměti
"))

pokr. cachí

☞ když jádra přistupují na disj. adresy, tak vše, co je v L2,
(je i v L3 (a co není v L2 není ani v L3))
→ když ne, tak je to v L3 jen jednou "

$\Rightarrow$ zbytek tam místo na
další data, která se
do L2 nevejdou

! ty přístupy do cachí jsou

pipelineované - ty latence rozm-
něnají, že by se to po tu
dobu zablokovalo

(ale zase zápisy na data,
co jsou ve více L2,
zprísobují, že si je
budou hodně navzájem
vyhazovat ")

TLB

• 2 úrovně: - ~64 záznamů (pro 4 kB stránky) TLB1
 $\Rightarrow$ 256 kB adr. prost.

- ~1024 záznamů TLB2, lat. 2-3 ns
- společné pro hypertext. vládku

není-li záznam v TLB:

• page walk (2-4 přístupy do paměti)

↓
procházet stránk. tabulkou

↑
přes cache!

• page Fault → OS

proces přístupu do paměti:

mov rmm1, [rbx+8*rax+1000]

1) spočte se ten výraz (64 bit)

virt. adresa (48 bit)

2) TLB1 36 bitů (normální) 12 bitů fyz. adr. (46 bitů)

TLB2

page walk

45 61501

adr. řádky
- klíč do cache

z toho

se vše použije jako
a dolní kousek (6 bitů)
jako index

ale ten

ani nebyl

změněn

tím překladem,

takže se to

udělá už

předtím

použije se

to jako

index do

L1D

⇒ z toho

je tag

a ten

se porovná

s klíčem

Algoritmy a cache

přístupy do paměti:

• s předvídatelnou adresou

- lin. příchody potom - procesor pozná a dělá

- lin. příchody s větším skokem (analize) -

• s náhodnou adresou

- hash. tabulky

⇒ velká latence

- dá se pomoci paralelismem

- bucket sort

- bin. vyhl.

- spojové struktury ⇒ pointer chasing !! FUJ

Přístupy s předvídatelnou adresou

- cache line ⇒ husté lin. příchody mají! dobré hit ratio

- write buffers ⇒ zápisy obvykle nezdvojnají

- HW a SW prefetching

- latence se skrývá paralelním vykonáváním jiné užitečné činnosti

paralelně
s tím
překladem

použije se

to jako

index do

L1D

⇒ z toho

je tag

a ten

se porovná

s klíčem

Cache-oblivious algorithms

cache-aware = alg. vyladěn pro konkr. parametry cache

cache-oblivious = víme, že cache $\exists$, ale neznáme její parametry



nepočítáme arithm. op., ale přesuny mezi hl. pamětí a cache^(#cache miss)

$$\rightarrow n \sim |\text{vstup}|, C \sim |\text{cache}| \Rightarrow O(f(n, C))$$

#cache miss se dá vyjádřit jako násobek arithm. složí. alg.

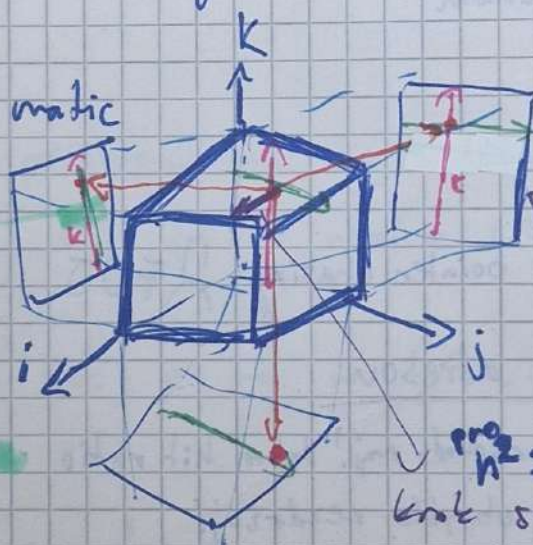
$$O(f(n, C)) = O(t(n) * g(n, C))$$

↓
jak často dochází ke cache missům

navíc pro velkou n často $O(g(n, C)) = O(g(C))$

u dobrých alg. s rost. C klesá $g(C)$ k nule

Pr.: nás. matice



pro $n^2 \gg C$
krok sem záměrně
celé nové řádky

→ vykazujeme si data, která
budeme potřebovat

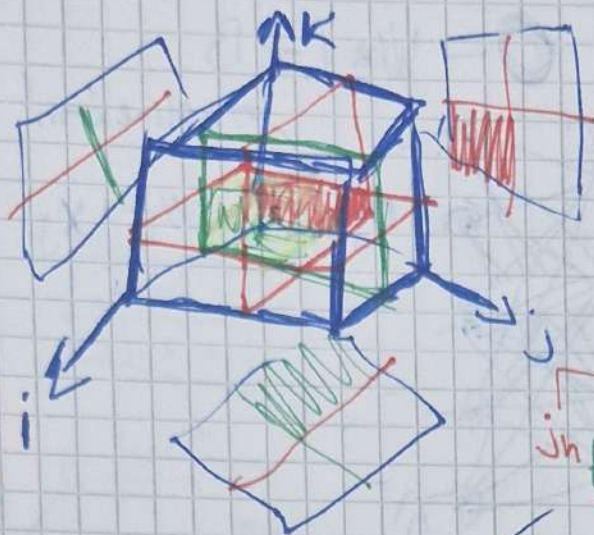
→ cache vůbec nepomáhá

data vel. n^2 ,
 n -krát je tam
stejně → n^3

$$t = O(n^3) \Rightarrow g(C) = 1$$

↓
celá matice
budem z hl. paměti
stahovat n krát

Le'pe: matice zmenšime



$j = j_h \cdot \frac{n}{2} + j_l$
 for i_h in $(0, m)$
 for j_h in $(0, m)$
 for k_h in $(0, m)$
 for i_l in $(0, \frac{n}{2})$
 for j_l in $(0, \frac{n}{2})$
 for k_l in $(0, \frac{n}{2})$

typicky se dělá

$m=2$ (a pak rekursivně tak, aby se to tím věřilo)

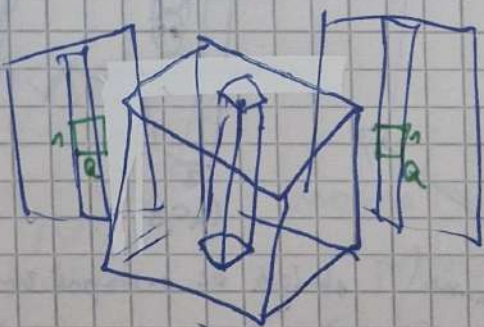
$O(C^{-1/2} * n^3) \rightarrow f(C) = C^{-1/2}$
 (11)

tedy $n' < C \cdot \frac{1}{3}$

161 cache
to může jít na 3

na nejvyšší úrovni:

vyjde to asi: $n' \approx 50$



$Q^2 + 2Q$ celkem do registrů

(např. AVX512: $32 \cdot 16$)

$\rightarrow Q \approx 20$

$\frac{L1 \gg R}{ALU} = \frac{2}{Q} = \frac{1}{10} \rightarrow$ 1 přesun na 10 operací

DCNN

$$O = W * I + B$$

Fully connected vrstvy:



$$o_i = f(\sum w_{ij} \cdot x_j + b_i)$$

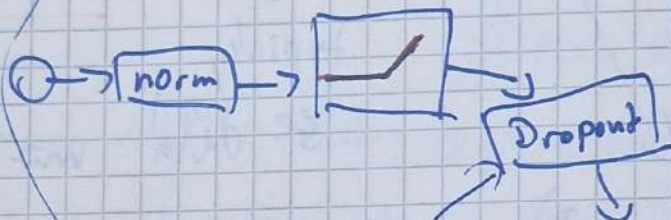
nelineární!



zde se učí!

ještě normalizace, aby ty hodnoty byly kolem 0

(str. hodn. 0, rozptyl 1)



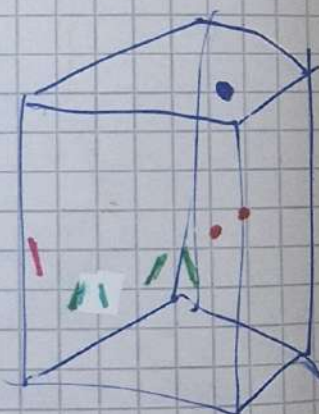
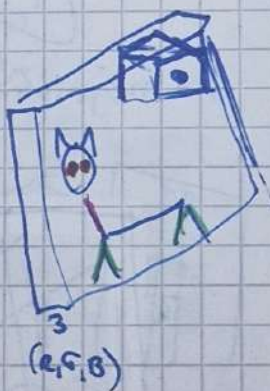
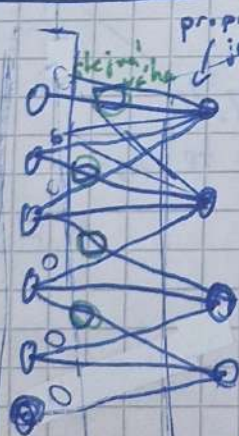
jednou za čas se náhodně data vynulují! (v době učení)

$$3 \cdot 7 \cdot 7 \cdot 9 = 147 \text{ parametrů}$$

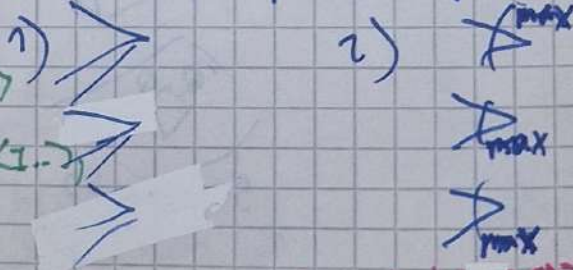
$$3 \cdot 7 \cdot 7 \cdot 9 = 147 \text{ arit. op.}$$

Konvoluční vrstvy:

propojen jen se sousedy



obrázek se postupně zmenšuje:



$$F(I) = \int \dots$$

lépe: $(F(\text{stř. integr. konst. } I, \dots))$

void forward - 1x pro celý ⁽¹⁶⁾ batch obrázek $\in [0, 15]$
 $\Rightarrow$ 0 1 rozměr navíc (vnější for cyklus)

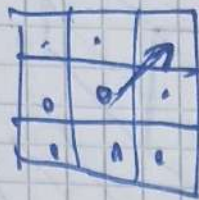
voli 6 dalších

1. optimalizace

n. je to nejsložitější

Zbytek je jednoduché

z kombinovat do jednoho (a vektorializovat)
 (tuhle je další optimalizace)



(3x3)
 kernel
 k_h, k_w
 s_h, s_w
 t

$\in [0, 2]$

hor 2, ver 2 - výstupní range

inv, wtr, outv
 - tenzory

4-rozměrný

outb = outv[b]
 3 rozměry

JAK

Změna

zanorenosti cyklu

+ vektorizace

jsou tam pomocné funkce (viz dokumentace)

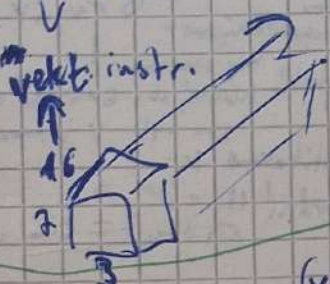
až to bude mít 1 rozměr, tak na tom jde udělat `data()` $\Rightarrow$ vekt. instr.!!

(možno upravit weights v permutation-policy)

tagged::co: $x \rightarrow x'$
 $x' \rightarrow x$

+ operator $\sim$ na hodnotách

konkr. vstava:



hor-batch = 7
 thr-size = 3

static-loop

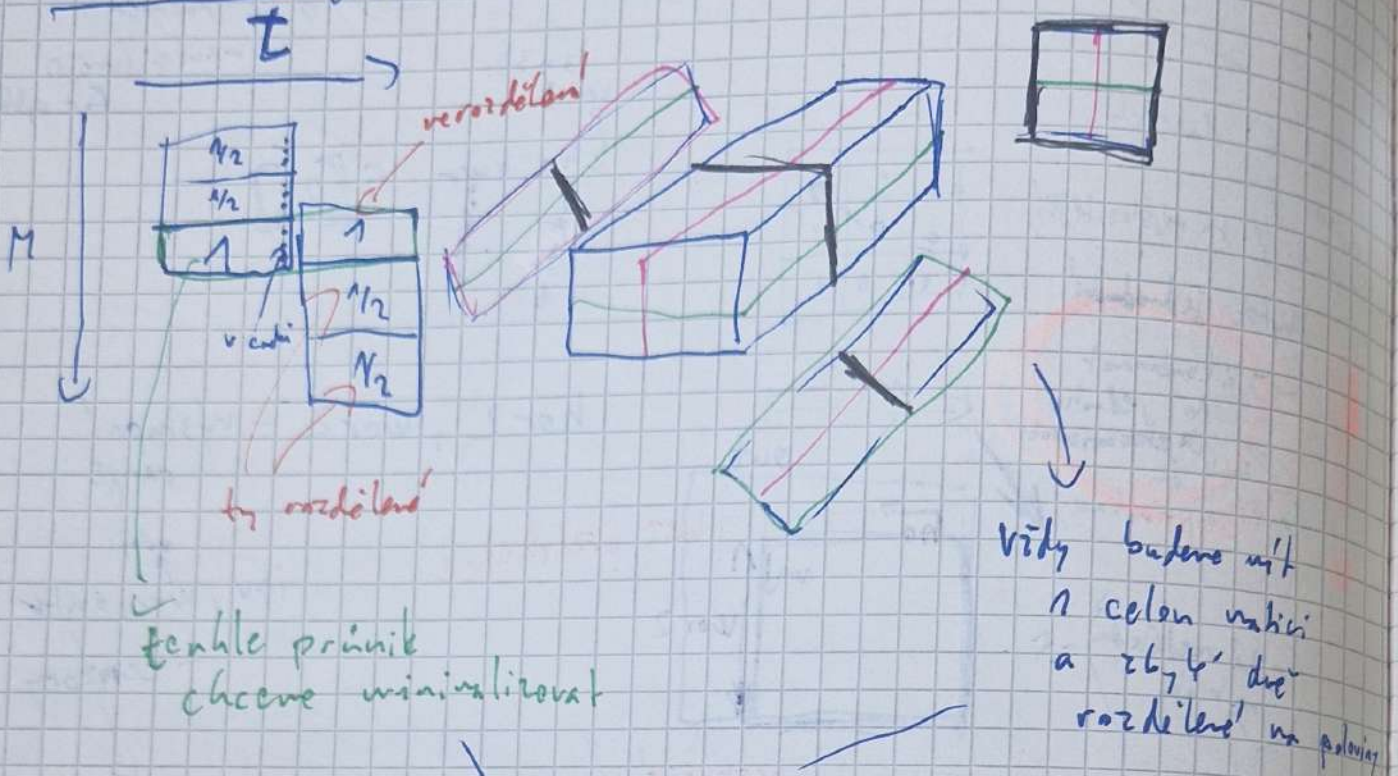
(Vš V-m $\leftarrow$ vektor)

perfectly-shrunk < lanes > (car) $\rightarrow$ parciální specializace

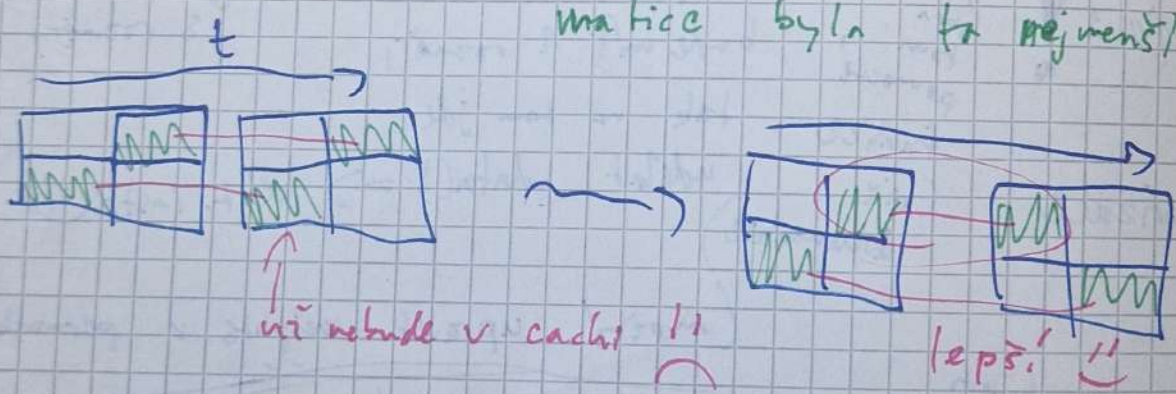
možno nejvnitřnější cyklus přepsat na forky bez těch typových indexů, aby to překladač autovektorializoval (nebo vektorializovat ručně - možná bude potřeba parciální specializace)

car / lanes

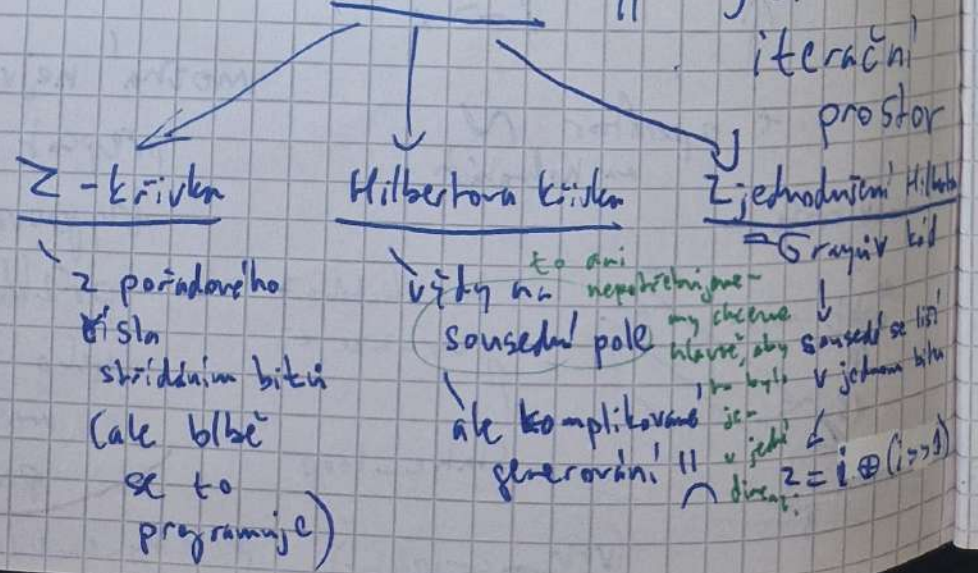
Ideální alg. - násobení matic



takže dělit tak, aby ta nedělená matice byla ta nejmenší



hledáme vhodnou křivku vyplňující

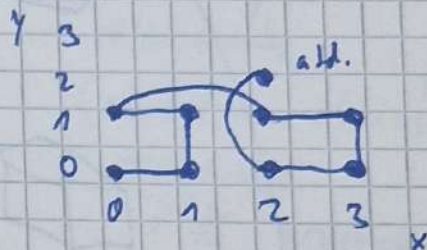


y	x
0	0 0 0 0
0	0 0 0 1
0	0 0 1 0
0	0 0 1 1
0	0 1 0 0
0	0 1 0 1
0	0 1 1 0
0	0 1 1 1
1	0 0 0 0
1	0 0 0 1
1	0 0 1 0
1	0 0 1 1
1	0 1 0 0
1	0 1 0 1
1	0 1 1 0
1	0 1 1 1

2 - křivka

y	x
0	0 0 0 0
0	0 0 0 1
0	0 0 1 1
0	0 0 1 0
0	0 1 1 0
0	0 1 1 1
0	0 1 0 1
0	0 1 0 0
1	1 1 0 0
1	1 1 0 1
1	1 1 1 1
1	1 1 1 0
1	1 0 1 0
1	1 0 1 1

Gray



Skoky jen v jedné dimenzi (1)

Cache matematicky

předp. 1 vláknem, divíme se na nějaký podproblém, Cache miss

T := doba běhu toho podproblému
 C := velikost cache

$m(t_1, t_2) := \#$ různých adres, na které jsme přistoupili v čas. intervalu (t_1, t_2)

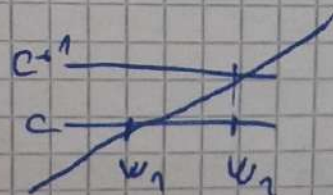
když vybereme t_1, t_2
 $a \quad m(t_1, t_2) \geq C$
 $\Rightarrow$ cache miss

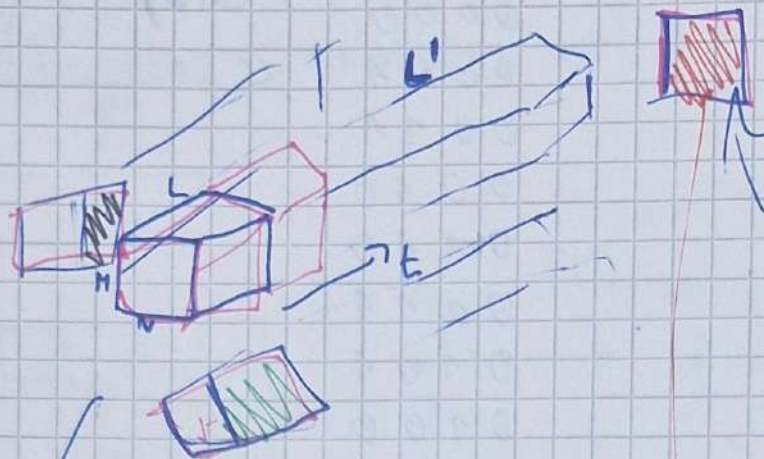
$m(w) :=$ počet $m(t, t+w)$ přes $\forall t \in [0, T]$

~~$$|A| = 3 \cdot 8 \cdot 8 = 192$$~~

$$X(C) = \frac{\partial m(w)}{\partial w} (m^{-1}(C))$$

↑
frekvence cache miss





Porád stejná
 $\Rightarrow$ chceme, aby zůstaly
 v cache

$\Downarrow$
 zaručíme, aby se tam
 vešly 2 by soustředěny
 krychličky

 $= n \cdot n \cdot l$

 +  $= n \cdot l + n \cdot l + n \cdot n$

Frekvence cache missů $= \frac{1}{n} + \frac{1}{n} + \frac{1}{l}$

co nejmenší - ideálně 1

$2 \cdot n \cdot l + 2 \cdot n \cdot l + n \cdot n \leq C$

$n = N$, $n \approx \sqrt{C}$

velikost L_0 - registra

a podobně pak
 pro vyšší úroveň

cache - aby se tam vešly 2 bytů celá kniha

$n \cdot l' + 2 \cdot n \cdot l' + 2 \cdot n \cdot n \leq C$

velikost L_1

CPU

Tranziistory

- CMOS $\Rightarrow$ polem řízený tranzistor

Metal-Oxide-Semiconductor

↓
jsou tam
obě
polarity
(n-ty p-ty)

vrstvy:

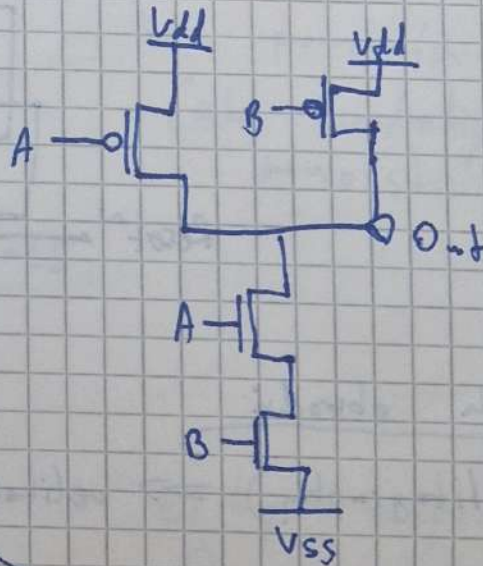
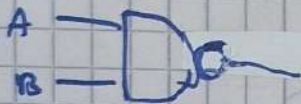


chová se jako
elektronky
(napětí na řídicí
elektrodě
řídí průchodnost
tého tranzistoru)

- Paněti používají odlišné technologie

Hradla

NAND - $\neg(A \& B)$



Spotřeba energie:

- ty dráty se chovají
jako malinké
kondenzátory

$\Rightarrow$ chvilka trvá
než se vybijí
 $\Rightarrow$ spotřebou
zbytečné energie (nJ)

platí Dennardovo škálování - čím menší,
tím menší
kap. Konden-
zátorů

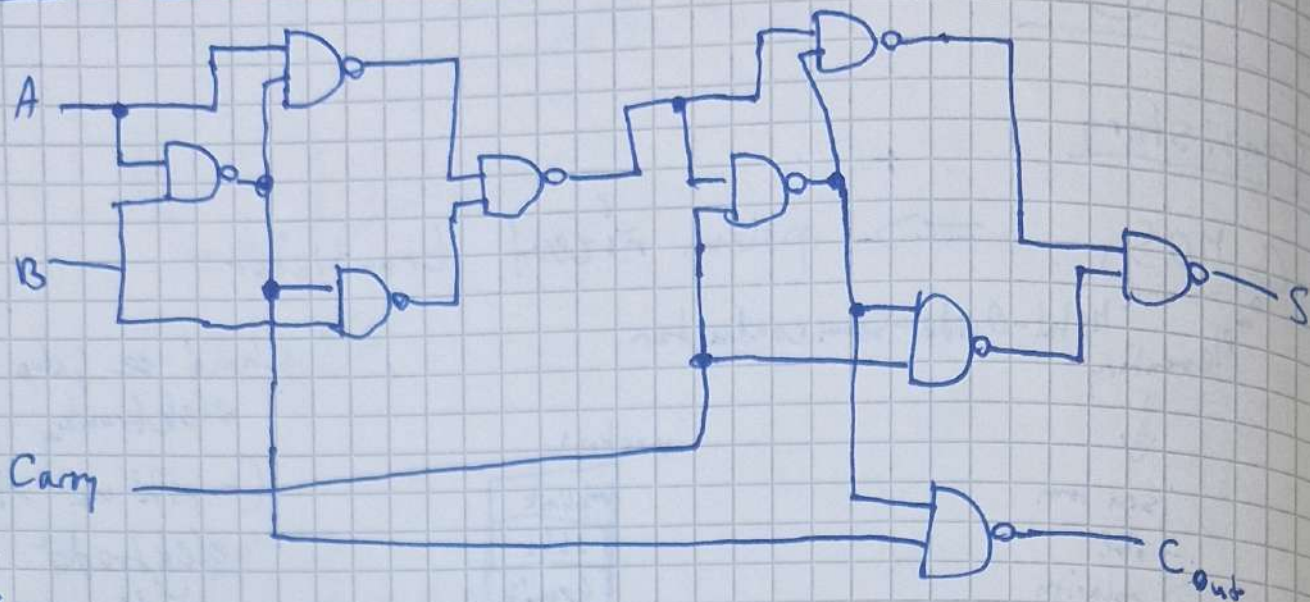
$\leftarrow$
klesá
spotřeba
energie

- když se zavřou jedny ty
tranziistory, tak chvilka trvá, než
se zavřou i ty další $\Rightarrow$ malý zkrat

- ty tranzistory nejsou nikdy zavřené úplně $\rightarrow$ něco tam pořád trochu leze $\Rightarrow$ spotřeba energie

tohle je
největší
problém

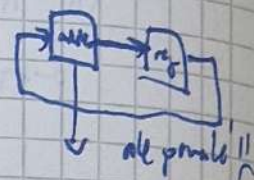
1-bitová úplná sčítacka



N-bitová sčítacka



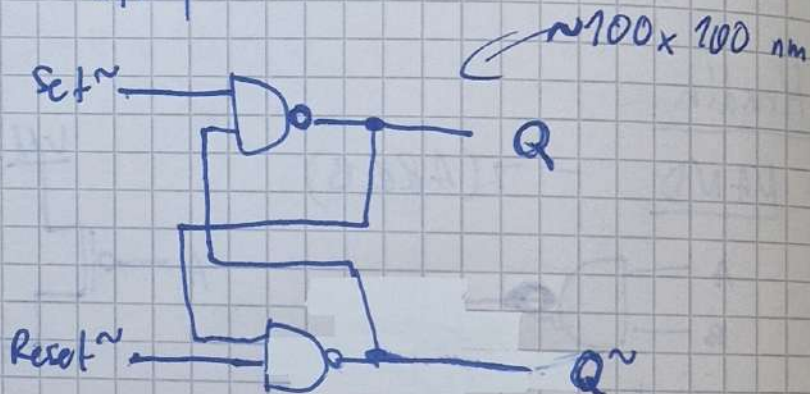
XX



Kombinační a sekvenční obvody

acyklický
graf

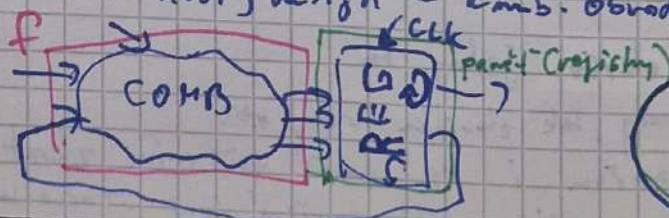
cykly



Design levels

- 1) process (lithography) $\Rightarrow$ velikosti součástek, množství prvků, ...
- 2) physical design $\Rightarrow$ masky $\Rightarrow$ frekvence, ...
- 3) transistor level $\Rightarrow$ síť tranzistorů
- 4) logic gate level $\Rightarrow$ síť hradel

5a) RTL (register transfer level) design - komb. obvody + registry



$$r = f(r)$$

5b) System-level design - System-C, Verilog, VHDL
 → vlastně takové progr. jazyky - programy takové,
 že umožňují
 sekvenční vykonání
 ⇒ to se převede
 na RTL
 by nástroje
 vyvíjí třeba
 firma Synopsys

Násobička

- školní alg.

1-bit. násobička - jenom AND hradla

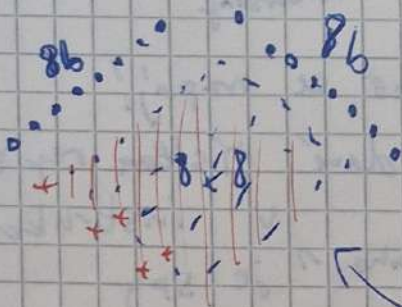
→ těžké je to sečíst výsledky

$O(N^2)$ zpoždění

s pomocnými registry

⇒ $O(N)$ když to uděláme sekvenčně
 bez pomocných registrů
 (ale kvadratická
 množství seřteček)

Wallacova násobička - dnes



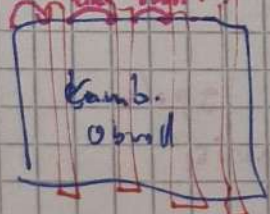
- 64 * AND

tabulka by
 měla být
 celou postavenou

triky s uspořádáním

⇒ postupně se
 redukuje #bitů
 (ze 3 bitů
 se udělají
 2)

n. počet registry



pipelining

$O(\log_2 N \times N)$ vrstev

např. pro 64b je to 21 vrstev

Obecně

- sekvencní impl. $\Rightarrow$ min. tranzistorů, velká latence (mnoho dělán)
- čistě kombinační impl. $\Rightarrow$ hodně tranzistorů, malá latence
- kombinační impl. rozdělení na fáze (stage) $\Rightarrow$ pipelining

$\Rightarrow$ větší propustnost

paralelizace:

- pipeliny

- SIMD

- multithreading

} nezlepšují latenci,
ale zvyšují
propustnost

Out-of-order execution

instr. $\Rightarrow$ mikroopence

out-of-order
fronty

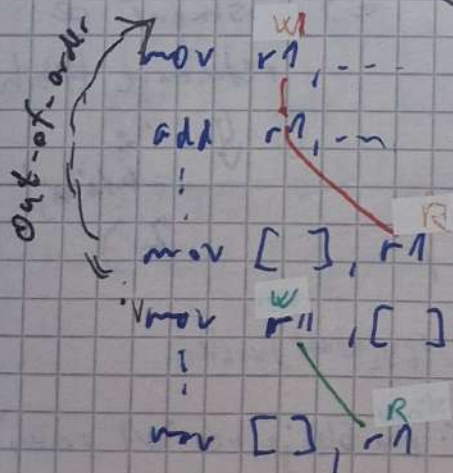
= zápis ^{to píše} in-order $\rightarrow$ retirement jednotek

Renamer

lady se dělá i z toho retirement = store forwarding

výk. jednotky

fyzických registrů je ~ 160 , na ně se mapují
architekturní registry (ty dávají číslům registrů
v instrukci)
16 (z toho 1 je SP)
četlost.



zruší nepřepojení na jiný reg.

antidependence

ale přitom to spolu
vůbec neresolvují

$\rightarrow$ může se
vykonávat
out-of-
order

Hyperthreading

- výk. jednotky obsluhují více vláken ($\Rightarrow$ při branch mis prediction se vyhadzuje celá pipeline, ale pokud by instr. daného vlákna ")

Branch prediction

- na podmíněná skoky nutné,
ale používá se i na nepodmíněné
- už ve front-endu se emuluje
HLK zásobník na návratové
adresy z return ")