

Wpływ architektury procesora na system operacyjny

Krzysztof Kąs, Paweł Bedyński, Krzysztof Niemkiewicz

Wydział Matematyki, Informatyki i Mechaniki Uniwersytetu Warszawskiego

07.11.2008

Plan Prezentacji

- 1 Część1 - Potokowanie, WMB (Paweł Bedyński)
- 2 Część 2 - Architektury 32-bitowe i 64-bitowe (Krzysztof Kąs)
- 3 Część 3 VLIW + Multiprocessing (Krzysztof Niemkiewicz)
- 4 Część 4 - Multiprocessing cd. (Krzysztof Niemkiewicz)

gdzie szukać informacji...



tu ...

```
cat /proc/cpuinfo
```

i tu...

```
/sys/devices/system/cpu/
```

i tutaj

```
linux/Documentation/cpu-freq/user-guide.txt
```

hot plug CPU

```
Documentation/hot-plug.txt
```

gdzie szukać informacji...



tu ...

cat /proc/cpuinfo

i tu...

/sys/devices/system/cpu/

i tutaj

linux/Documentation/cpu-freq/user-guide.txt

hot plug CPU

Documentation/hot-plug.txt

```
Processor : 0
vendor_id : GenuineIntel
cpu family : 6
model : 15
model name : Intel(R) Core(TM)2 Duo CPU    T5250  @ 1.50GHz
stepping : 13
cpu MHz : 1000.000
cache size : 2048 KB
physical id : 0
siblings : 2
core id : 0
cpu cores : 2
fdiv_bug : no
hlt_bug : no
f00f_bug : no
coma_bug : no
fpu : yes
fpu_exception : yes
cpuid level : 10
wp : yes
flags : fpu vme de pse tsc msr pae mce cx8 apic sep mtrr
pge mca cmov pat pse36 clflush dts acpi mmx
fxsr sse sse2 ss ht tm pbe nx lm constant_tsc
arch_perfmon pebs bts pni monitor ds_cpl
est tm2 ssse3 cx16 xtpr lahf_lm
bogomips : 2995.62
clflush size : 64
```

gdzie szukać informacji...



tu ...

cat /proc/cpuinfo

i tu...

/sys/devices/system/cpu/

i tutaj

linux/Documentation/cpu-freq/user-guide.txt

hot plug CPU

Documentation/hot-plug.txt

```
processor : 1
vendor_id : GenuineIntel
cpu family : 6
model : 15
model name : Intel(R) Core(TM)2 Duo CPU    T5250  @ 1.50GHz
stepping : 13
cpu MHz : 1000.000
cache size : 2048 KB
physical id : 0
siblings : 2
core id : 1
cpu cores : 2
fdiv_bug : no
hlt_bug : no
f00f_bug : no
coma_bug : no
fpu : yes
fpu_exception : yes
cpuid level : 10
wp : yes
flags : fpu vme de pse tsc msr pae mce cx8 apic sep mtrr
pge mca cmov pat pse36 clflush dts acpi mmx
fxsr sse sse2 ss ht tm pbe nx lm constant_tsc
arch_perfmon pebs bts pni monitor ds_cpl
est tm2 ssse3 cx16 xtpr lahf_lm
bogomips : 2992.63
clflush size : 64
```

gdzie szukać informacji...



tu ...

```
cat /proc/cpuinfo
```

i tu...

```
/sys/devices/system/cpu/
```

i tutaj

```
linux/Documentation/cpu-freq/user-guide.txt
```

hot plug CPU

```
Documentation/hot-plug.txt
```

gdzie szukać informacji...



tu ...

cat /proc/cpuinfo

i tu...

/sys/devices/system/cpu/

i tutaj

linux/Documentation/cpu-freq/user-guide.txt

hot plug CPU

Documentation/hot-plug.txt

```

/sys/devices/system/cpu$ ls
cpu0  cpu1  cpuidle  sched_mc_power_savings

/sys/devices/system/cpu/cpu0$ ls -R
.:
cache  cpufreq  cpuidle  crash_notes  topology
./cache:
index0  index1  index2
./cache/index0: (to samo w index1 i index 2)
coherency_line_size  physical_line_partition  type
level                 shared_cpu_map           ways_of_associativity
number_of_sets        size
./cpufreq:
affected_cpus  ondemand  scaling_driver  stats
cpuinfo_cur_freq  scaling_available_frequencies  scaling_governor
cpuinfo_max_freq  scaling_available_governors  scaling_max_freq
cpuinfo_min_freq  scaling_cur_freq  scaling_min_freq
./cpufreq/ondemand:
ignore_nice_load  sampling_rate  sampling_rate_min
powersave_bias  sampling_rate_max  up_threshold
./cpufreq/stats:
time_in_state  total_trans  trans_table
./cpuidle:
state0  state1  state2
./cpuidle/state0: (to samo w state 1 i state 2)
latency  name  power  time  usage
./topology:
core_id  core_siblings  physical_package_id  thread_siblings

```

gdzie szukać informacji...



tu ...

```
cat /proc/cpuinfo
```

i tu...

```
/sys/devices/system/cpu/
```

i tutaj

```
linux/Documentation/cpu-freq/user-guide.txt
```

hot plug CPU

```
Documentation/hot-plug.txt
```


gdzie szukać informacji...



tu ...

cat /proc/cpuinfo

i tu ...

/sys/devices/system/cpu/

i tutaj

linux/Documentation/cpu-freq/user-guide.txt

hot plug CPU

Documentation/hot-plug.txt

```

1  CPU frequency and voltage scaling code in the Linux(TM) kernel
2
3
4  L i n u x   C P U F r e q
5
6  U S E R   G U I D E
7
8
9  Dominik Brodowski <linux@brodo.de>
10
11
12
13  Clock scaling allows you to change the clock speed of the CPUs on the
14  fly. This is a nice method to save battery power, because the lower
15  the clock speed, the less power the CPU consumes.
16
17
18
19
20
21
22
23
24  The following processors for the x86 architecture are supported by cpufreq:
25
26  AMD Elan - SC400, SC410
27  AMD mobile K6-2+
28  AMD mobile K6-3+
29  AMD mobile Duron
30  AMD mobile Athlon
31  AMD Opteron
32  AMD Athlon 64
33  Cyrix Media GXm
34  Intel mobile PIII and Intel mobile PIII-M on certain chipsets
35  Intel Pentium 4, Intel Xeon
36  Intel Pentium M (Centrino)
37  National Semiconductors Geode GX
38  Transmeta Crusoe
39  Transmeta Efficeon
40  VIA Cyrix 3 / C3
41  various processors on some ACPI 2.0-compatible systems
42
43
44
45
46
47
48
49
50
51
52
53
54
55
56
57
58
59
60
61
62
63
64
65
66
67
68
69
70
71

```

gdzie szukać informacji...



tu ...

```
cat /proc/cpuinfo
```

i tu...

```
/sys/devices/system/cpu/
```

i tutaj

```
linux/Documentation/cpu-freq/user-guide.txt
```

hot plug CPU

```
Documentation/hot-plug.txt
```


Instrukcje "specjalne" dla różnych architektur - Intel MMX (1)

MMX 1997r.

(MultiMedia eXtensions lub Matrix Math eXtensions) to zestaw 57 instrukcji SIMD dla procesorów Pentium i zgodnych. Rozkazy MMX mogą realizować działania logiczne i arytmetyczne na liczbach całkowitych. Pierwotnie wprowadzone w 1997 przez Intela dla procesorów Pentium MMX, aktualnie dostępne również na procesory innych producentów.

- Rejestry MMX mają rozmiar 64 bitów, jest ich 8.

Rejestry MMX w assemblerze noszą nazwy `mm0`,
`mm1`,... `mm7`.

- Nowe typy danych - "packed"
 - 8 - 8 bitów (packed byte),
 - 4 - 16 bitów (packed word),
 - 2 - 32 bity (packed dword),
 - 1 - 64 bity (quad word).



Instrukcje "specjalne" dla różnych architektur - Intel MMX (1)

MMX 1997r.

(MultiMedia eXtensions lub Matrix Math eXtensions) to zestaw 57 instrukcji SIMD dla procesorów Pentium i zgodnych. Rozkazy MMX mogą realizować działania logiczne i arytmetyczne na liczbach całkowitych. Pierwotnie wprowadzone w 1997 przez Intela dla procesorów Pentium MMX, aktualnie dostępne również na procesory innych producentów.

- Rejestry MMX mają rozmiar 64 bitów, jest ich 8.

Rejestry MMX w assemblerze noszą nazwy **mm0**,
mm1,... **mm7**.

- Nowe typy danych - "packed"

- 8 - 8 bitów (packed byte),
- 4 - 16 bitów (packed word),
- 2 - 32 bity (packed dword),
- 1 - 64 bity (quad word).



Instrukcje "specjalne" dla różnych architektur - Intel MMX (2)

PCMPEGB = P-CMP-EQ-B

```
mm1 = |01001000|01001100|10000100|00101110|11000100|01001110|11100000|00000001|
+
mm2 = |00001000|01001100|10000100|10101111|11000100|01000000|11110111|00000001|
=
mm2 = |00000000|11111111|11111111|00000000|11111111|00000000|00000000|11111111|
```

operacje:

- arytmetyczne (np. PADDB, PADDW, PADDD)
- porównania (np. PCMPEQB, PCMPEQW, PCMPEQD)
- logiczne (np. POR, PAD)
- przesunięcia bitowe (np. PSLLW, PSLLD, PSLLQ)

Zastosowania:

- wyświetlanie grafiki trójwymiarowej: przekształcenia geometryczne, cieniowanie, teksturowanie;
- dekodowanie obrazów JPEG i PNG; filmów MPEG (transformacja DCT)

Instrukcje "specjalne" dla różnych architektur - Intel MMX (2)

PCMPEGB = P-CMP-EQ-B

```
mm1 = |01001000|01001100|10000100|00101110|11000100|01001110|11100000|00000001|
+
mm2 = |00001000|01001100|10000100|10101111|11000100|01000000|11110111|00000001|
=
mm2 = |00000000|11111111|11111111|00000000|11111111|00000000|00000000|11111111|
```

operacje:

- arytmetyczne (np. PADDB, PADDW, PADDD)
- porównania (np. PCMPEQB, PCMPEQW, PCMPEQD)
- logiczne (np. POR, PAD)
- przesunięcia bitowe (np. PSLLW, PSLLD, PSLLQ)

Zastosowania:

- wyświetlanie grafiki trójwymiarowej: przekształcania geometryczne, cieniowanie, teksturowanie;
- dekodowanie obrazów JPEG i PNG; filmów MPEG (transformacja DCT)

Instrukcje "specjalne" dla różnych architektur - Intel MMX (2)

PCMPEGB = P-CMP-EQ-B

```
mm1 = |01001000|01001100|10000100|00101110|11000100|01001110|11100000|00000001|
+
mm2 = |00001000|01001100|10000100|10101111|11000100|01000000|11110111|00000001|
=
mm2 = |00000000|11111111|11111111|00000000|11111111|00000000|00000000|11111111|
```

operacje:

- arytmetyczne (np. PADDB, PADDW, PADDD)
- porównania (np. PCMPEQB, PCMPEQW, PCMPEQD)
- logiczne (np. POR PAD)
- przesunięcia bitowe (np. PSLLW, PSLLD, PSLLQ)

Zastosowania:

- wyświetlanie grafiki trójwymiarowej: przekształcenia geometryczne, cieniowanie, teksturowanie;
- dekodowanie obrazów JPEG i PNG; filmów MPEG (transformacja DTC)

Instrukcje "specjalne" dla różnych architektur - SSE/SSE2/SSE3/SSSE3/SSE4

SSE - Streaming SIMD Extensions

SSE - jest nazwa zestawu instrukcji wprowadzonego w 1999 roku po raz pierwszy w procesorach Pentium III firmy Intel. <http://sourceware.org/gdb/papers/linux/linux-sse.html>

SSE2 - Streaming SIMD Extensions 2

Technologia ta została wprowadzona w procesorach rodziny Pentium 4 oraz Athlon 64. (20 XI 2000)

SSE3 - Streaming SIMD Extensions 3

oznaczany również przez firmę Intel jako Prescott New Instructions lub PNI (III 2001)

SSSE3 - Supplemental Streaming SIMD Extensions 3

Intel Core 2, AMD Phenom (2006)

SSE4 - Streaming SIMD Extensions 4

Intel Core 2, Penryn (VI 2007)

Instrukcje "specjalne" dla różnych architektur - SSE/SSE2/SSE3/SSSE3/SSE4

SSE - Streaming SIMD Extensions

SSE - jest nazwa zestawu instrukcji wprowadzonego w 1999 roku po raz pierwszy w procesorach Pentium III firmy Intel. <http://sourceware.org/gdb/papers/linux/linux-sse.html>

SSE2 - Streaming SIMD Extensions 2

Technologia ta została wprowadzona w procesorach rodziny Pentium 4 oraz Athlon 64. (20 XI 2000)

SSE3 - Streaming SIMD Extensions 3

oznaczany również przez firmę Intel jako Prescott New Instructions lub PNI (III 2001)

SSSE3 - Supplemental Streaming SIMD Extensions 3

Intel Core 2, AMD Phenom (2006)

SSE4 - Streaming SIMD Extensions 4

Intel Core 2, Penryn (VI 2007)

Instrukcje "specjalne" dla różnych architektur - SSE/SSE2/SSE3/SSSE3/SSE4

SSE - Streaming SIMD Extensions

SSE - jest nazwa zestawu instrukcji wprowadzonego w 1999 roku po raz pierwszy w procesorach Pentium III firmy Intel. <http://sourceware.org/gdb/papers/linux/linux-sse.html>

SSE2 - Streaming SIMD Extensions 2

Technologia ta została wprowadzona w procesorach rodziny Pentium 4 oraz Athlon 64. (20 XI 2000)

SSE3 - Streaming SIMD Extensions 3

oznaczany również przez firmę Intel jako Prescott New Instructions lub PNI (III 2001)

SSSE3 - Supplemental Streaming SIMD Extensions 3

Intel Core 2, AMD Phenom (2006)

SSE4 - Streaming SIMD Extensions 4

Intel Core 2, Penryn (VI 2007)

Instrukcje "specjalne" dla roznych architektur - SSE

C M P N E Q P S		x m m 0 , x m m 1							
	1 . 0		- 5 . 3		1 6 . 5		1 7 . 2		x m m 0
	7 . 0		- 5 . 3		1 6 . 5		1 7 . 3		x m m 1
=		=		=		=			
	1 1 1 . . 1 1 1 1		0 0 0 . . 0 0 0 0		0 0 0 . . 0 0 0 0		1 1 1 . . 1 1 1 1		x m m 0

Typy danych - 4 elementowy wektor (128bit)

- ➊ packed (rownoległe) : wykonując równocześnie 4 niezależne działania zmiennoprzecinkowe na odpowiadających sobie elementach wektorów;
- ➋ scalar (skalarne) : wykonując działanie tylko na pierwszych elementach wektorów.

Przesyłanie liczb zmiennoprzecinkowych między rejestrami i pamięcią

- **MOVAPS, MOVUPS** - Przesłanie 4 liczb zmiennoprzecinkowych pomiędzy rejestrem XMM, a pamięcią lub innym rejestrem XMM
- **MOVLHPS, MOVHLPS** - Przesłanie między rejestrami 64-bitów (2 liczb zmiennoprzecinkowych).

Instrukcje "specjalne" dla różnych architektur - SSE (2)

- SSE

Pamięć Podreeczna:

- Z = PREFETCH : poziomy 0, 1, 2 ; prefetchNTA (non-temporal)
- BEZ = MOVNTQ (z rejestru MMX), MOVNTPS (z rejestru SSE)

- SSE3

- **MONITOR** - ustala początkowy adres zakresu pamięci (rozmiar tego obszaru jest stały, zależny od modelu procesora), który następnie jest automatycznie monitorowany przez procesor
- **MWAIT** - powoduje wejście w zaprogramowaną sprężynę petla, która jest przypisana flagom, gdy flaga zostanie ustawiona, a więc gdy nastąpi zapis pod obserwowany adres.

Instrukcje "specjalne" dla różnych architektur - SSE (2)

- SSE

Pamięć Podreeczna:

- Z = PREFETCH : poziomy 0, 1, 2 ; prefetchNTA (non-temporal)
- BEZ = MOVNTQ (z rejestru MMX), MOVNTPS (z rejestru SSE)

- SSE3

- MONITOR - ustala początkowy adres zakresu pamięci (monitor tego obszaru jest stały, zależny od modelu procesora). Kiedy nastąpiła jest automatycznie monitorowany przez procesor.
- MWAIT - powoduje wejście w zaprogramowaną specyficzną pację, która jest przypisana flagom, gdy flagi zostaną ustawione, a więc gdy nastąpi jakiś pod obserwowany adres.

Instrukcje "specjalne" dla różnych architektur - SSE (2)

- SSE

Pamięć Podreeczna:

- Z = PREFETCH : poziomy 0, 1, 2 ; prefetchNTA (non-temporal)
- BEZ = MOVNTQ (z rejestru MMX), MOVNTPS (z rejestru SSE)

- SSE3

- MONITOR ustala początkowy adres zakresu pamięci (rozmiar tego obszaru jest stały, zależny od modelu procesora), który następnie jest automatycznie monitorowany przez procesor
- MWAIT powoduje wejście w zoptymalizowaną sprzętową pętlę, która jest przerywana dopiero, gdy flaga zostanie ustawiona, a więc gdy nastąpi zapis pod obserwowany adres.

Instrukcje "specjalne" dla różnych architektur - SSE (2)

- SSE

Pamięć Podreeczna:

- Z = PREFETCH : poziomy 0, 1, 2 ; prefetchNTA (non-temporal)
- BEZ = MOVNTQ (z rejestru MMX), MOVNTPS (z rejestru SSE)

- SSE3

- MONITOR ustala początkowy adres zakresu pamięci (rozmiar tego obszaru jest stały, zależny od modelu procesora), który następnie jest automatycznie monitorowany przez procesor
- MWAIT powoduje wejście w zoptymalizowaną sprzętową petle, która jest przerywana dopiero, gdy flaga zostanie ustawiona, a więc gdy nastąpi zapis pod obserwowany adres.

Instrukcje "specjalne" dla różnych architektur - przeszłość SSE5/AVX

AMD - SSE5

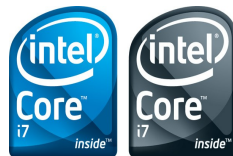
- 2009 - planowane pojawienie się pierwszego procesora
- 82 nowe rozkazy
- rozkazy 4 argumentowe (RISC)

$$w = y \pm (\pm x \cdot z)$$



INTEL - AVX

- 2010 - planowane pojawienie się pierwszego procesora
- ponad 250 nowych instrukcji
- 256-bitowe rejestry (ymm0 ... ymm15) (w dalszych wersjach 512 i 1024 bitowe)
- Sprzętowe wsparcie dla szyfrowania AES
<http://software.intel.com/en-us/articles/advanced-encryption-standard-aes-instructions-set>



Instrukcje "specjalne" dla różnych architektur - przeszłość SSE5/AVX

AMD - SSE5

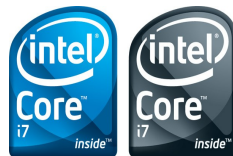
- 2009 - planowane pojawienie się pierwszego procesora
- 82 nowe rozkazy
- rozkazy 4 argumentowe (RISC)

$$w = y \pm (\pm x \cdot z)$$



INTEL - AVX

- 2010 - planowane pojawienie się pierwszego procesora
- ponad 250 nowych instrukcji
- 256 bitowe rejestry (ymm0 ... ymm15) (w dalszych wersjach 512 i 1024 bitowe)
- Sprzętowe wsparcie dla szyfrowania AES
<http://software.intel.com/en-us/articles/advanced-encryption-standard-aes-instructions-set>



RISC vs CISC

CISC lata 60-te/70-te

Complex Instruction Set Computers

- duża liczba instrukcji 100 -
- mała liczba rejestrów - specjalizowane
- duża liczba trybów adresowania
- złożone instrukcje wymagające kilku-kilkunastu cykli zegara
- Complex od złożoności poszczególnych instrukcji a nie od ich ilości (RISC często wykonuje więcej instrukcji niż CISC)

RISC lata 80-te

Reduced Instruction Set Computers

- ograniczona liczba instrukcji
- duża liczba rejestrów - uniwersalne
- redukcja trybów adresowania
- wszystkie rozkazy wykonują się w jednym cyklu maszynowym (pipelining)

RISC vs CISC

CISC lata 60-te/70-te

Complex Instruction Set Computers

- duża liczba instrukcji 100 -
- mała liczba rejestrów - specjalizowane
- duża liczba trybów adresowania
- złożone instrukcje wymagające kilku-kilkunastu cykli zegara
- Complex od złożoności poszczególnych instrukcji a nie od ich ilości (RISC często wykonuje więcej instrukcji niż CISC)

RISC lata 80-te

Reduced Instruction Set Computers

- ograniczona liczba instrukcji
- duża liczba rejestrów - uniwersalne
- redukcja trybów adresowania
- wszystkie rozkazy wykonują się w jednym cyklu maszynowym (pipelining)

RISC vs CISC

CISC lata 60-te/70-te

Complex Instruction Set Computers

- duża liczba instrukcji 100 -
- mała liczba rejestrów - specjalizowane
- duża liczba trybów adresowania
- złożone instrukcje wymagające kilku-kilkunastu cykli zegara
- Complex od złożoności poszczególnych instrukcji a nie od ich ilości (RISC często wykonuje więcej instrukcji niż CISC)

RISC lata 80-te

Reduced Instruction Set Computers

- ograniczona liczba instrukcji
- duża liczba rejestrów - uniwersalne
- redukcja trybów adresowania
- wszystkie rozkazy wykonują się w jednym cyklu maszynowym (pipelining)

RISC vs CISC

CISC lata 60-te/70-te

Complex Instruction Set Computers

- duża liczba instrukcji 100 -
- mała liczba rejestrów - specjalizowane
- duża liczba trybów adresowania
- złożone instrukcje wymagające kilku-kilkunastu cykli zegara
- Complex od złożoności poszczególnych instrukcji a nie od ich ilości (RISC często wykonuje więcej instrukcji niż CISC)

RISC lata 80-te

Reduced Instruction Set Computers

- ograniczona liczba instrukcji
- duża liczba rejestrów - uniwersalne
- redukcja trybów adresowania
- wszystkie rozkazy wykonują się w jednym cyklu maszynowym (pipelining)

RISC vs CISC

CISC lata 60-te/70-te

Complex Instruction Set Computers

- duża liczba instrukcji 100 -
- mała liczba rejestrów - specjalizowane
- duża liczba trybów adresowania
- złożone instrukcje wymagające kilku-kilkunastu cykli zegara
- Complex od złożoności poszczególnych instrukcji a nie od ich ilości (RISC często wykonuje więcej instrukcji niż CISC)

RISC lata 80-te

Reduced Instruction Set Computers

- ograniczona liczba instrukcji
- duża liczba rejestrów - uniwersalne
- redukcja trybów adresowania
- wszystkie rozkazy wykonują się w jednym cyklu maszynowym (pipelining)

RISC vs CISC

CISC lata 60-te/70-te

Complex Instruction Set Computers

- duża liczba instrukcji 100 -
- mała liczba rejestrów - specjalizowane
- duża liczba trybów adresowania
- złożone instrukcje wymagające kilku-kilkunastu cykli zegara
- Complex od złożoności poszczególnych instrukcji a nie od ich ilości (RISC często wykonuje więcej instrukcji niż CISC)

RISC lata 80-te

Reduced Instruction Set Computers

- ograniczona liczba instrukcji
- duża liczba rejestrów - uniwersalne
- redukcja trybów adresowania
- wszystkie rozkazy wykonują się w jednym cyklu maszynowym (pipelining)

RISC vs CICS problemy

CISC

Complex Instruction Set Computers

- dużo trybów odwołań do pamięci
- skomplikowane instrukcje
- duża liczba odwołań do pamięci (koszt czasu)

RISC

Reduced Instruction Set Computers

- 7 trybów odwołań do pamięci
- 7 trybów odwołań do pamięci
- 7 trybów odwołań do pamięci

RISC vs CICS problemy

CISC

Complex Instruction Set Computers

- dużo trybów odwołań do pamięci
- skomplikowane instrukcje
- duża liczba odwołań do pamięci (koszt czasu)

RISC

Reduced Instruction Set Computers

- 7 trybów odwołań do pamięci
- 7 trybów odwołań do pamięci
- 7 trybów odwołań do pamięci

RISC vs CICS problemy

CISC

Complex Instruction Set Computers

- dużo trybów odwołan do pamięci
- skomplikowane instrukcje
- duża liczba odwołań do pamięci (koszt czasu)

RISC

Reduced Instruction Set Computers

- 7 trybów odwołań do pamięci
- 7 trybów odwołań do pamięci
- 7 trybów odwołań do pamięci

RISC vs CISC problemy

CISC

Complex Instruction Set Computers

- dużo trybów odwołań do pamięci
- skomplikowane instrukcje
- duża liczba odwołań do pamięci (koszt czasu)

RISC

Reduced Instruction Set Computers

- 1 tryb odwołania do pamięci
- proste instrukcje
- mała liczba odwołań do pamięci

RISC vs CICS problemy

CISC

Complex Instruction Set Computers

- dużo trybów odwołań do pamięci
- skomplikowane instrukcje
- duża liczba odwołań do pamięci (koszt czasu)

RISC

Reduced Instruction Set Computers

- ?
- ?
- ?

RISC vs CISC czym jest x86?

- do 80386 procesory rodziny x86 były typowymi CISC-ami
- od 80486 wewnętrzne architektury coraz bardziej zaczynają przypominać RISC (od zewnątrz paskudne CISC-y)

RISC vs CISC czym jest x86?

- do 80386 procesory rodziny x86 były typowymi CISC-ami
- od 80486 wewnętrzne architektury coraz bardziej zaczynają przypominać RISC (od zewnątrz paskudne CISC-y)

RISC vs CISC czym jest x86?

- do 80386 procesory rodziny x86 były typowymi CISC-ami
- od 80486 wewnętrzne architektury coraz bardziej zaczynają przypominać RISC (od zewnątrz paskudne CISC-y)

RISC

- krótki czas wykonywania rozkazu
- duża częstotliwość zegara sugeruje, że operacje nie mogą być "obszerne"

RISC vs CISC czym jest x86?

- do 80386 procesory rodziny x86 były typowymi CISC-ami
- od 80486 wewnętrzne architektury coraz bardziej zaczynają przypominać RISC (od zewnątrz paskudne CISC-y)

RISC

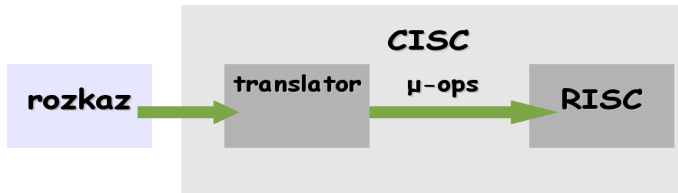
- krótki czas wykonywania rozkazu
- duża częstotliwość zegara sugeruje że operacje nie mogą być "obszerne"

CISC

- duża liczba instrukcji (bardzo duża)
- skomplikowane instrukcje ... i tak są rozbijane na mniejsze

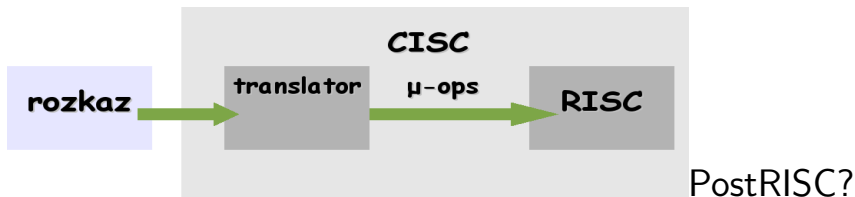
RISC vs CISC czym jest x86?

- do 80386 procesory rodziny x86 były typowymi CISC-ami
- od 80486 wewnętrzne architektury coraz bardziej zaczynają przypominać RISC (od zewnątrz paskudne CISC-y)



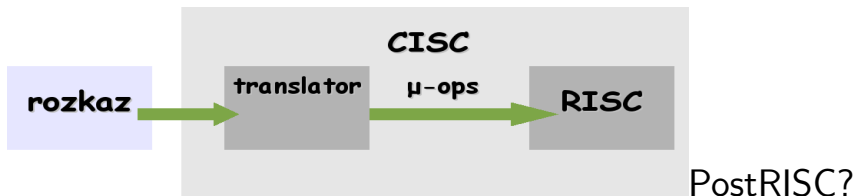
RISC vs CISC czym jest x86?

- do 80386 procesory rodziny x86 były typowymi CISC-ami
- od 80486 wewnętrzne architektury coraz bardziej zaczynają przypominać RISC (od zewnątrz paskudne CISC-y)



RISC vs CISC czym jest x86?

- do 80386 procesory rodziny x86 były typowymi CISC-ami
- od 80486 wewnętrzne architektury coraz bardziej zaczynają przypominać RISC (od zewnątrz paskudne CISC-y)



<http://community.tomshardware.pl/comments.php?DiscussionID=135563>

Pipelining - co to ?

etapy przetwarzania rozkazów:

- 1 pobranie rozkazu z pamieci operacyjnej
- 2 dekodowanie rozkazu, pobieranie danych z rejestrów
- 3 wykonanie instrukcji, lub obliczenie adresu
- 4 dostęp do rejestru
- 5 zapis wyniku

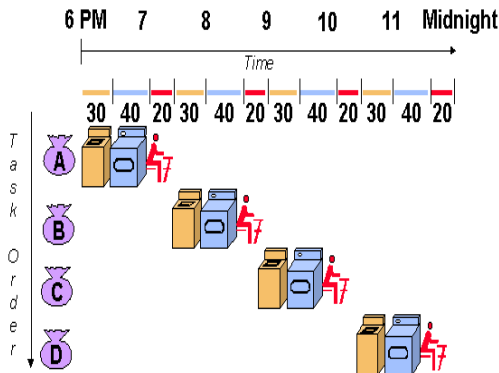
architektury RISC mają przewagę
ponieważ ich instrukcje są małe

Pipelining - co to ?

etapy przetwarzania rozkazów:

- 1 pobranie rozkazu z pamieci operacyjnej
- 2 dekodowanie rozkazu, pobieranie danych z rejestrow
- 3 wykonanie instrukcji, lub obliczenie adresu
- 4 dostep do rejestru
- 5 zapis wyniku

architektury RISC maja przewage
poniewaz ich instrukcje sa male

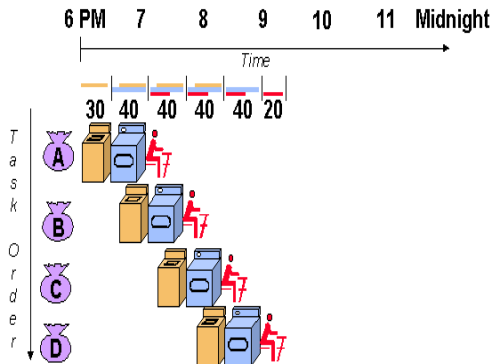


Pipelining - co to ?

etapy przetwarzania rozkazów:

- ❶ pobranie rozkazu z pamieci operacyjnej
- ❷ dekodowanie rozkazu, pobieranie danych z rejestrow
- ❸ wykonanie instrukcji, lub obliczenie adresu
- ❹ dostep do rejestru
- ❺ zapis wyniku

architektury RISC maja przewage
poniewaz ich instrukcje sa male



Pipelining - problemy

problem1

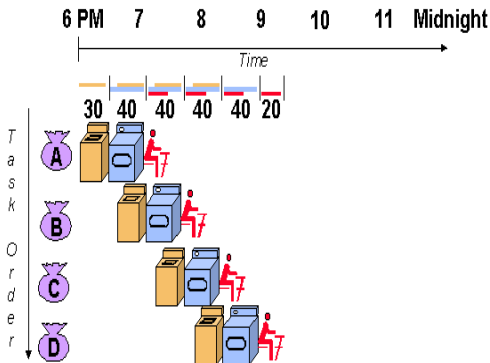
```
add $r3, $r2, $r1
add $r5, $r4, $r3
```

```
/* dodaj $r2 do $r1 i
wynik umiesc w $r3 */
```

problem2

```
Loop :
    add $r3, $r2, $r1
    sub $r6, $r5, $r4
    beq $r3, $r6, Loop
//dalsze instrukcje
```

```
/* beq to skok do Loop:
jesli $r3==$r6 */
```



Pipelining - problemy

problem1

```
add $r3, $r2, $r1
add $r5, $r4, $r3
```

```
/* dodaj $r2 do $r1 i
wynik umiesc w $r3 */
```

problem2

```
Loop :
    add $r3, $r2, $r1
    sub $r6, $r5, $r4
    beq $r3, $r6, Loop
//dalsze instrukcje
```

```
/* beq to skok do Loop:
jesli $r3==$r6 */
```

Rozwiązania:

- Zamiana kolejności
(w przykładzie 1 wepchanie innych NIEZALEŻNYCH instrukcji może uratować sprawę)
reorganizacja kolejności jest zadaniem kompilatora, ma on za zadanie zminimalizować liczbę przestojów
- Przewidywanie wyniku
na ogół pętle piszemy po to by wykonywały się więcej niż raz
procesor 'zgaduje' prawdopodobną ścieżkę i wykonuje ją w przód
- Przewidywanie dynamiczne
pamiętanie historii wyborów przy każdym rozgałęzieniu
nawet 90% skuteczność trafienia - algorytm Yeh'a
- Superskalarność
osobne potoki - powielone części sprzętu (ew. specjalizacja w typach obliczeń)

Pipelining - problemy

problem1

```
add $r3, $r2, $r1
add $r5, $r4, $r3
```

```
/* dodaj $r2 do $r1 i
wynik umiesc w $r3 */
```

problem2

Loop :

```
    add $r3, $r2, $r1
    sub $r6, $r5, $r4
    beq $r3, $r6, Loop
```

```
//dalsze instrukcje
```

```
/* beq to skok do Loop:
jesli $r3==$r6 */
```

Rozwiązania:

- Zamiana kolejności
(w przykładzie 1 wepchanie innych NIEZALEŻNYCH instrukcji może uratować sprawę)
reorganizacja kolejności jest zadaniem kompilatora, ma on za zadanie zminimalizować liczbę przestojów
- Przewidywanie wyniku
na ogół pętle piszemy po to by wykonywały się więcej niż raz
procesor 'zgaduje' prawdopodobną ścieżkę i wykonuje ją w przód
- Przewidywanie dynamiczne
pamiętanie historii wyborów przy każdym rozgałęzieniu
nawet 90% skuteczność trafienia - algorytm Yeh'a
- Superskalarność
osobne potoki - powielone części sprzętu (ew. specjalizacja w typach obliczeń)

Pipelining - problemy

problem1

```
add $r3, $r2, $r1
add $r5, $r4, $r3
```

```
/* dodaj $r2 do $r1 i
wynik umiesc w $r3 */
```

problem2

Loop :

```
    add $r3, $r2, $r1
    sub $r6, $r5, $r4
    beq $r3, $r6, Loop
```

```
//dalsze instrukcje
```

```
/* beq to skok do Loop:
jesli $r3==$r6 */
```

Rozwiązania:

- Zamiana kolejności
(w przykładzie 1 wepchanie innych NIEZALEŻNYCH instrukcji może uratować sprawę)
reorganizacja kolejności jest zadaniem kompilatora, ma on za zadanie zminimalizować liczbę przestojów
- Przewidywanie wyniku
na ogół pętle piszemy po to by wykonywały się więcej niż raz
procesor 'zgaduje' prawdopodobną ścieżkę i wykonuje ją w przód
- Przewidywanie dynamiczne
pamiętanie historii wyborów przy każdym rozgałęzieniu
nawet 90% skuteczność trafienia - algorytm Yeh'a
- Superskalarność
osobne potoki - powielone części sprzętu (ew. specjalizacja w typach obliczeń)

Pipelining - problemy

problem1

```
add $r3, $r2, $r1
add $r5, $r4, $r3
```

```
/* dodaj $r2 do $r1 i
wynik umiesc w $r3 */
```

problem2

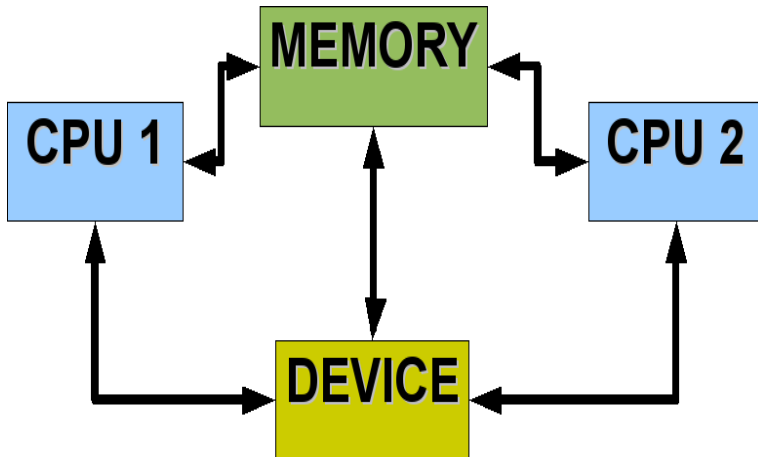
```
Loop :
    add $r3, $r2, $r1
    sub $r6, $r5, $r4
    beq $r3, $r6, Loop
//dalsze instrukcje
```

```
/* beq to skok do Loop:
jesli $r3==$r6 */
```

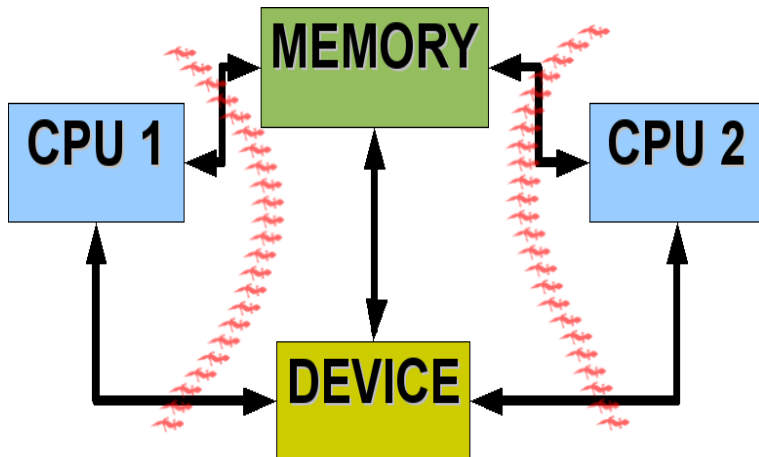
Rozwiązania:

- Zamiana kolejności
(w przykładzie 1 wepchanie innych NIEZALEŻNYCH instrukcji może uratować sprawę)
reorganizacja kolejności jest zadaniem kompilatora, ma on za zadanie zminimalizować liczbę przestojów
- Przewidywanie wyniku
na ogół pętle piszemy po to by wykonywały się więcej niż raz
procesor 'zgaduje' prawdopodobną ścieżkę i wykonuje ją w przód
- Przewidywanie dynamiczne
pamiętanie historii wyborów przy każdym rozgałęzieniu
nawet 90% skuteczność trafienia - algorytm Yeh'a
- Superskalarność
osobne potoki - powielone części sprzętu (ew. specjalizacja w typach obliczeń)

model



model



po co? (1)

A==1; B==2

CPU 1

- A=3
- B=4

```
STORE A=3,
STORE A=3,
STORE A=3,
STORE A=3,
STORE A=3,
STORE A=3,
STORE A=3,
STORE B=4,
STORE B=4, ...
...
```

CPU 2

- x=A
- y=B

```
x=LOAD A->3, y=LOAD B->4
y=LOAD B->4, x=LOAD A->3
STORE B=4, y=LOAD B->4
y=LOAD B->2, STORE B=4
STORE B=4, x=LOAD A->3
x=LOAD A->3, STORE B=4
x=LOAD A->3, y=LOAD B->4
```

razem 24 przeploty, a możliwe wyniki ...

po co? (1)

A==1; B==2

CPU 1

- A=3
- B=4

```
STORE A=3,      STORE B=4,
STORE A=3,      STORE B=4,
STORE A=3,      x=LOAD A->3,
STORE A=3,      x=LOAD A->3,
STORE A=3,      y=LOAD B->2,
STORE A=3,      y=LOAD B->2,
STORE B=4,      STORE A=3,
STORE B=4, ...
...
```

CPU 2

- x=A
- y=B

```
x=LOAD A->3,    y=LOAD B->4
y=LOAD B->4,    x=LOAD A->3
STORE B=4,      y=LOAD B->4
y=LOAD B->2,    STORE B=4
STORE B=4,      x=LOAD A->3
x=LOAD A->3,    STORE B=4
x=LOAD A->3,    y=LOAD B->4
```

razem 24 przeploty, a mozliwe wyniki ...

po co? (2)

A == 1, B == 2, C = 3, P == &A, Q == &C

CPU 1

- B=4
- P=&B

CPU 2

- Q=P
- D=*Q

(Q == &A) and (D == 1)

(Q == &B) and (D == 2)

(Q == &B) and (D == 4)

UWAGA: CPU2 nigdy nie załaduje C do D ponieważ CPU załaduje P do Q zanim rozpocznie ładowanie *Q

po co? (2)

A == 1, B == 2, C = 3, P == &A, Q == &C

CPU 1

- B=4
- P=&B

CPU 2

- Q=P
- D=*Q

(Q == &A) and (D == 1)

(Q == &B) and (D == 2)

(Q == &B) and (D == 4)

UWAGA: CPU2 nigdy nie załaduje C do D ponieważ CPU załaduje P do Q zanim rozpocznie ładowanie *Q

co mamy gwarantowane...

Z perspektywy pojedynczego procesora jego instrukcje 'zależne' wykonują się szeregowo

CPU 1

- $Q = P$; $D = *Q$;
- $Q = \text{LOAD } P$, $D = \text{LOAD } *Q$

'zachodzące' instrukcje zapisu/odczytu ('wewnątrz' jednego procesora) będą szeregowane

CPU 1

- $a = *X$; $*X = b$;
- $a = \text{LOAD } *X$, $\text{STORE } *X = b$

CPU 1

- $*X = c$; $d = *X$;
- $\text{STORE } *X = c$, $d = \text{LOAD } *X$

co mamy gwarantowane...

Z perspektywy pojedynczego procesora jego instrukcje 'zależne' wykonują się szeregowo

CPU 1

- $Q = P$; $D = *Q$;
- $Q = \text{LOAD } P$, $D = \text{LOAD } *Q$

'zachodzące' instrukcje zapisu/odczytu ('wewnątrz' jednego procesora) będą szeregowane

CPU 1

- $a = *X$; $*X = b$;
- $a = \text{LOAD } *X$, $\text{STORE } *X = b$

CPU 1

- $*X = c$; $d = *X$;
- $\text{STORE } *X = c$, $d = \text{LOAD } *X$

co mamy gwarantowane...

Z perspektywy pojedynczego procesora jego instrukcje 'zależne' wykonują się szeregowo

CPU 1

- $Q = P$; $D = *Q$;
- $Q = \text{LOAD } P$, $D = \text{LOAD } *Q$

'zachodzące' instrukcje zapisu/odczytu ('wewnątrz' jednego procesora) będą szeregowane

CPU 1

- $a = *X$; $*X = b$;
- $a = \text{LOAD } *X$, $\text{STORE } *X = b$

CPU 1

- $*X = c$; $d = *X$;
- $\text{STORE } *X = c$, $d = \text{LOAD } *X$

Co musimy, a czego nie możemy zakładać..

Nie możemy zakładać, że niezależne odczyty/zapisy wykonują się szeregowo
tzn. dla $X = *A; Y = *B; *D = Z$; możemy mieć:

```
X = LOAD *A,  Y = LOAD *B,  STORE *D = Z
X = LOAD *A,  STORE *D = Z, Y = LOAD *B
Y = LOAD *B,  X = LOAD *A,  STORE *D = Z
Y = LOAD *B,  STORE *D = Z, X = LOAD *A
STORE *D = Z, X = LOAD *A,  Y = LOAD *B
STORE *D = Z, Y = LOAD *B,  X = LOAD *A
```

Musimy zakładać, że zależne (overlapping) zapisy/odczyty mogą być mergowane lub pomijane

tzn. $X = *A; Y = *(A + 4)$; możemy mieć:

```
X = LOAD *A; Y = LOAD *(A + 4);
Y = LOAD *(A + 4); X = LOAD *A;
{X, Y} = LOAD {*A, *(A + 4)};
```

a dla $*A = X; *A = Y$; możemy mieć:

```
STORE *A = X; Y = LOAD *A;
STORE *A = Y = X;
```


Co musimy, a czego nie możemy zakładać..

Nie możemy zakładać, że niezależne odczyty/zapisy wykonują się szeregowo
tzn. dla $X = *A$; $Y = *B$; $*D = Z$; możemy mieć:

```
X = LOAD *A,  Y = LOAD *B,  STORE *D = Z
X = LOAD *A,  STORE *D = Z, Y = LOAD *B
Y = LOAD *B,  X = LOAD *A,  STORE *D = Z
Y = LOAD *B,  STORE *D = Z, X = LOAD *A
STORE *D = Z, X = LOAD *A,  Y = LOAD *B
STORE *D = Z, Y = LOAD *B,  X = LOAD *A
```

Musimy zakładać, że zależne (overlapping) zapisy/odczyty mogą być mergowane lub pomijane
tzn. $X = *A$; $Y = *(A + 4)$; możemy mieć:

```
X = LOAD *A; Y = LOAD *(A + 4);
Y = LOAD *(A + 4); X = LOAD *A;
{X, Y} = LOAD {*A, *(A + 4)};
```

a dla $*A = X$; $*A = Y$; możemy mieć:

```
STORE *A = X; Y = LOAD *A;
STORE *A = Y = X;
```

Co musimy, a czego nie możemy zakładać..

Nie możemy zakładać, że niezależne odczyty/zapisy wykonują się szeregowo
tzn. dla $X = *A$; $Y = *B$; $*D = Z$; możemy mieć:

```
X = LOAD *A,  Y = LOAD *B,  STORE *D = Z
X = LOAD *A,  STORE *D = Z, Y = LOAD *B
Y = LOAD *B,  X = LOAD *A,  STORE *D = Z
Y = LOAD *B,  STORE *D = Z, X = LOAD *A
STORE *D = Z, X = LOAD *A,  Y = LOAD *B
STORE *D = Z, Y = LOAD *B,  X = LOAD *A
```

Musimy zakładać, że zależne (overlapping) zapisy/odczyty mogą być mergowane lub pomijane
tzn. $X = *A$; $Y = *(A + 4)$; możemy mieć:

```
X = LOAD *A; Y = LOAD *(A + 4);
Y = LOAD *(A + 4); X = LOAD *A;
{X, Y} = LOAD {*A, *(A + 4)};
```

a dla $*A = X$; $*A = Y$; możemy mieć:

```
STORE *A = X; Y = LOAD *A;
STORE *A = Y = X;
```

Co to jest MEMORY BARRIER

Niezależne operacje na pamięci mogą być wykonywane efektywnie w losowej kolejności, ale może to stwarzać problemy (np. CPU-CPU lub I/O). Czasem jednak wymagane jest, aby kompilator i CPU były poinformowane o pewnych restrykcjach.

MEMORY BARRIER

Memory barriers nakładają częściowy porządek na operacjach po dwóch stronach bariery. Powodują, że sekwencja zdarzeń na pamięci staje się dla innych części systemu imitacją sytuacji, w której to CPU wykonuje te operacje w takiej, a nie innej kolejności.

Co to jest MEMORY BARRIER

Niezależne operacje na pamięci mogą być wykonywane efektywnie w losowej kolejności, ale może to stwarzać problemy (np. CPU-CPU lub I/O). Czasem jednak wymagane jest, by kompilator i CPU były poinformowane o pewnych restrykcjach.

MEMORY BARRIER

Memory barriers nakładają częściowy porządek na operacjach po dwóch stronach bariery. Powodują, że sekwencja zdarzeń na pamięci staje się dla innych części systemu imitacją sytuacji, w której to CPU wykonuje te operacje w takiej, a nie innej kolejności.

Co to jest MEMORY BARRIER

Niezależne operacje na pamięci mogą być wykonywane efektywnie w losowej kolejności, ale może to stwarzać problemy (np. CPU-CPU lub I/O). Czasem jednak wymagane jest, aby kompilator i CPU były poinformowane o pewnych restrykcjach.

MEMORY BARRIER

Memory barriers nakładają częściowy porządek na operacjach po dwóch stronach bariery. Powodują, że sekwencja zdarzeń na pamięci staje się dla innych części systemu imitacją sytuacji, w której to CPU wykonuje te operacje w takiej, a nie innej kolejności.

Bogactwo barier ...

- ❶ WMB - write (STORE) memory barrier
- ❷ data dependency barriers
- ❸ RMB - read (LOAD) memory barrier
- ❹ general memory barrier
- ❺ LOCK / UNLOCK

Czy wszędzie bariery są potrzebne ?

Jak się mają te gwarancje do różnych architektur?

linux/arch/alpha/include/asm/barrier.h

Bogactwo barier ...

- ❶ WMB - write (STORE) memory barrier
- ❷ data dependency barriers
- ❸ RMB - read (LOAD) memory barrier
- ❹ general memory barrier
- ❺ LOCK / UNLOCK

Czy wszędzie bariery są potrzebne ?

Jak się mają te gwarancje do różnych architektur?

linux/arch/alpha/include/asm/barrier.h

Bogactwo barier ...

- ❶ WMB - write (STORE) memory barrier
- ❷ data dependency barriers
- ❸ RMB - read (LOAD) memory barrier
- ❹ general memory barrier
- ❺ LOCK / UNLOCK

Czy wszędzie bariery są potrzebne ?

Jak się mają te gwarancje do różnych architektur?

linux/arch/alpha/include/asm/barrier.h

Bogactwo barier ...

- ❶ WMB - write (STORE) memory barrier
- ❷ data dependency barriers
- ❸ RMB - read (LOAD) memory barrier
- ❹ general memory barrier
- ❺ LOCK / UNLOCK

Czy wszędzie bariery są potrzebne ?

Jak się mają te gwarancje do różnych architektur?

linux/arch/alpha/include/asm/barrier.h

Bogactwo barier ...

- ❶ WMB - write (STORE) memory barrier
- ❷ data dependency barriers
- ❸ RMB - read (LOAD) memory barrier
- ❹ general memory barrier
- ❺ LOCK / UNLOCK

Czy wszędzie bariery są potrzebne ?

Jak się mają te gwarancje do różnych architektur?

linux/arch/alpha/include/asm/barrier.h

Bogactwo barier ...

- ❶ WMB - write (STORE) memory barrier
- ❷ data dependency barriers
- ❸ RMB - read (LOAD) memory barrier
- ❹ general memory barrier
- ❺ LOCK / UNLOCK

Czy wszędzie bariery są potrzebne ?

Jak się mają te gwarancje do różnych architektur?

`linux/arch/alpha/include/asm/barrier.h`

Bogactwo barier ...

- ❶ WMB - write (STORE) memory barrier
- ❷ data dependency barriers
- ❸ RMB - read (LOAD) memory barrier
- ❹ general memory barrier
- ❺ LOCK / UNLOCK

Czy wszędzie bariery są potrzebne ?

Jak się mają te gwarancje do różnych architektur?

[linux/arch/alpha/include/asm/barrier.h](#)

Czego nie możemy założyć o MEMORY BARRIER w LINUX KERNEL

NIE możemy założyć że

- 1 operacja dostępu do pamięci wykonana przed założeniem bariery będzie ukończona w chwili ukończenia operacji bariery (bariera może być postrzegana jako linia której niektóre instrukcje nie mogą przekraczać)
- 2 bariera na jednym CPU ma bezpośredni wpływ na drugi CPU lub inne urządzenia (co z pośrednim?)
- 3 jeden CPU będzie widział w prawidłowej kolejności efekty dostępu do pamięci innego CPU nawet jeśli ten drugi używa MEMORY BARRIER (odsylam do SMP BARRIER PAIRING)
- 4 inny hardware nie zmieni kolejności wykonywania instrukcji dostępu do pamięci

Czego nie możemy założyć o MEMORY BARRIER w LINUX KERNEL

NIE możemy założyć że

- ❶ operacja dostępu do pamięci wykonana przed założeniem bariery będzie ukończona w chwili ukończenia operacji bariery (bariera może być postrzegana jako linia której niektóre instrukcje nie mogą przekraczać)
- ❷ bariera na jednym CPU ma bezpośredni wpływ na drugi CPU lub inne urządzenia (co z pośrednim?)
- ❸ jeden CPU będzie widział w prawidłowej kolejności efekty dostępu do pamięci innego CPU nawet jeśli ten drugi używa MEMORY BARRIER (odsylam do SMP BARRIER PAIRING)
- ❹ inny hardware nie zmieni kolejności wykonywania instrukcji dostępu do pamięci

Czego nie możemy założyć o MEMORY BARRIER w LINUX KERNEL

NIE możemy założyć że

- 1 operacja dostępu do pamięci wykonana przed założeniem bariery będzie ukończona w chwili ukończenia operacji bariery (bariera może być postrzegana jako linia której niektóre instrukcje nie mogą przekraczać)
- 2 bariera na jednym CPU ma bezpośredni wpływ na drugi CPU lub inne urządzenia (co z pośrednim?)
- 3 jeden CPU będzie widział w prawidłowej kolejności efekty dostępu do pamięci innego CPU nawet jeśli ten drugi używa MEMORY BARRIER (odsylam do SMP BARRIER PAIRING)
- 4 inny hardware nie zmieni kolejności wykonywania instrukcji dostępu do pamięci

Czego nie możemy założyć o MEMORY BARRIER w LINUX KERNEL

NIE możemy założyć że

- 1 operacja dostępu do pamięci wykonana przed założeniem bariery będzie ukończona w chwili ukończenia operacji bariery (bariera może być postrzegana jako linia której niektóre instrukcje nie mogą przekraczać)
- 2 bariera na jednym CPU ma bezpośredni wpływ na drugi CPU lub inne urządzenia (co z pośrednim?)
- 3 jeden CPU będzie widział w prawidłowej kolejności efekty dostępu do pamięci innego CPU nawet jeśli ten drugi używa MEMORY BARRIER (odsylam do SMP BARRIER PAIRING)
- 4 inny hardware nie zmieni kolejności wykonywania instrukcji dostępu do pamięci

Czego nie możemy założyć o MEMORY BARRIER w LINUX KERNEL

NIE możemy założyć że

- 1 operacja dostępu do pamięci wykonana przed założeniem bariery będzie ukończona w chwili ukończenia operacji bariery (bariera może być postrzegana jako linia której niektóre instrukcje nie mogą przekraczać)
- 2 bariera na jednym CPU ma bezpośredni wpływ na drugi CPU lub inne urządzenia (co z pośrednim?)
- 3 jeden CPU będzie widział w prawidłowej kolejności efekty dostępu do pamięci innego CPU nawet jeśli ten drugi używa MEMORY BARRIER (odsylam do SMP BARRIER PAIRING)
- 4 inny hardware nie zmieni kolejności wykonywania instrukcji dostępu do pamięci

anomalie

Rozpatrzmy układ:

```

CPU 1          CPU 2
=====
{ A == 1, B == 2, C = 3, P == &A, Q == &C }
B = 4;
<write barrier>
P = &B

                Q = P;
                D = *Q;

```

W oczywisty sposób występuje tu data dependency i wydaje się że na koncu Q musi być albo &A lub &B oraz że

```

(Q == &A) implies (D == 1)
(Q == &B) implies (D == 4)

```

Jednak wiedza CPU2 o P może być uaktualniona PRZED wiedzą o B, co prowadzi do następującego przeplotu

```

(Q == &B) and (D == 2) ???

```

architektura DEC Alpha

anomalie

Rozpatrzmy układ:

```

CPU 1          CPU 2
=====
{ A == 1, B == 2, C = 3, P == &A, Q == &C }
B = 4;
<write barrier>
P = &B

                Q = P;
                D = *Q;

```

W oczywisty sposób występuje tu data dependency i wydaje się że na koncu Q musi być albo &A lub &B oraz że

```

(Q == &A) implies (D == 1)
(Q == &B) implies (D == 4)

```

Jednak wiedza CPU2 o P może być uaktualniona PRZED wiedzą o B, co prowadzi do następującego przeplotu

```
(Q == &B) and (D == 2) ????
```

architektura DEC Alpha

anomalie

Rozpatrzmy układ:

```

CPU 1          CPU 2
=====
{ A == 1, B == 2, C = 3, P == &A, Q == &C }
B = 4;
<write barrier>
P = &B

                Q = P;
                D = *Q;

```

W oczywisty sposób występuje tu data dependency i wydaje się że na koncu Q musi być albo &A lub &B oraz że

```

(Q == &A) implies (D == 1)
(Q == &B) implies (D == 4)

```

Jednak wiedza CPU2 o P może być updateowana PRZED wiedzą o B, co prowadzi do następującego przeplotu

(Q == &B) and (D == 2) ????

architektura DEC Alpha

anomalie - rozwiązanie

Rozwiązanie:

```

CPU 1          CPU 2
=====
{ A == 1, B == 2, C = 3, P == &A, Q == &C }
B = 4;
<write barrier>
P = &B

                Q = P;
                <data dependency barrier>
                D = *Q;
  
```

wprowadzenie data dependency barrier powoduje ze zajdzie jedna z tych 2 'oczekiwanych' sytuacji

MEMORY BARRIERS w LINUX KERNEL 1

COMPILER BARRIER - `barrier()`;

- prosta blokada powodująca że żadna operacja dostępu nie może przedostać się z jednej strony na drugą
- ta blokada jest osobna względem procesora więc ...
- procesor może później przeorganizować kolejność wykonywania instrukcji ?!

MEMORY BARRIERS w LINUX KERNEL 2

CPU MEMORY BARRIERS

TYPE	MANDATORY	SMP CONDITIONAL
=====	=====	=====
GENERAL	<code>mb()</code>	<code>smp_mb()</code>
WRITE	<code>wmb()</code>	<code>smp_wmb()</code>
READ	<code>rmb()</code>	<code>smp_rmb()</code>
DATA DEPENDENCY	<code>read_barrier_depends()</code>	<code>smp_read_barrier_depends()</code>

oraz, trochę bardziej zaawansowane:

- `set_mb(var, value)`
- `set_wbm(var, value)`

przypisują wartość do zmiennej a potem (w drugim przypadku) ustawiają write-barrier

MEMORY BARRIERS w LINUX KERNEL 3

CPU MEMORY BARRIERS

```
(*) smp_mb__before_atomic_dec();  
(*) smp_mb__after_atomic_dec();  
(*) smp_mb__before_atomic_inc();  
(*) smp_mb__after_atomic_inc();
```

Zobaczmy na przykładzie:

```
obj->dead = 1;  
smp_mb__before_atomic_dec();  
atomic_dec(&obj->ref_count);
```

To powoduje ze flaga 'dead' zostanie ustawiona zanim zmieni sie licznik. wiecej informacji w *Documentation\atomic_ops.txt*

```
(*) smp_mb__before_clear_bit(void);  
(*) smp_mb__after_clear_bit(void);
```

Zabezpiecza operacje przed wywołaniem cleara przed 'przeskoczeniem' po clearu

```
smp_mb__before_clear_bit();  
clear_bit( ... );
```


MEMORY BARRIERS w LINUX KERNEL 4

Posrednie wywołanie MEMORY BARRIER

Niektóre z funkcji jądra linuxa pośrednio wywołuje Memory Barriers (locking, scheduling, memory allocation)
scheduling, memory allocation

Co może być gwarantowane ?

Co z LOCKING Functions?

Jądro linuxa ma kilka konstrukcji typu LOCK m.in.

- spin locks
- mutexy
- semafony

MEMORY BARRIERS w LINUX KERNEL 4

Posrednie wywołanie MEMORY BARRIER

Niektóre z funkcji jądra linuxa pośrednio wywołuje Memory Barriers (locking, scheduling, memory allocation)

scheduling, memory allocation - full memory barrier

Co może być gwarantowane ?

Co z LOCKING Functions?

Jądro linuxa ma kilka konstrukcji typu LOCK m.in.

- spin locks
- mutexy
- semafony

MEMORY BARRIERS w LINUX KERNEL 4

Posrednie wywołanie MEMORY BARRIER

Niektóre z funkcji jądra linuxa pośrednio wywołuje Memory Barriers (locking, scheduling, memory allocation)

scheduling, memory allocation - full memory barrier

Co może być gwarantowane ?

Co z LOCKING Functions?

Jądro linuxa ma kilka konstrukcji typu LOCK m.in.

- spin locks
- mutexy
- semafony

MEMORY BARRIERS w LINUX KERNEL 4

Posrednie wywołanie MEMORY BARRIER

Niektóre z funkcji jądra linuxa pośrednio wywołuje Memory Barriers (locking, scheduling, memory allocation)

scheduling, memory allocation - full memory barrier

Co może być gwarantowane ?

Co z LOCKING Functions?

Jądro linuxa ma kilka konstrukcji typu LOCK m.in.

- spin locks
- mutexy
- semafony

MEMORY BARRIERS w LINUX KERNEL 5

LOCKING FUNCTIONS

- ❶ implikacje operacji LOCK
- ❷ implikacje operacji UNLOCK
- ❸ implikacje LOCK(1) ; LOCK(2)
- ❹ implikacje LOCK ; UNLOCK
- ❺ implikacje LOCK'a który sie nie powiodł

wiec (1),(2),(4) UNLOCK z następującym po nim LOCK są równoważne full barrier
natomiast LOCK z następującym po nim UNLOCK już nie

UWAGA!

Następstwem tego że LOCK i UNLOCK mogą być barierami 'w jedną stronę' jest to że efekty instrukcji spoza sekcji krytycznej mogą 'przesiaknąć' do wnętrza sekcji krytycznej

MEMORY BARRIERS w LINUX KERNEL 5

LOCKING FUNCTIONS

- 1 implikacje operacji LOCK
- 2 implikacje operacji UNLOCK
- 3 implikacje LOCK(1) ; LOCK(2)
- 4 implikacje LOCK ; UNLOCK
- 5 implikacje LOCK'a który sie nie powiodl

wiec (1),(2),(4) UNLOCK z nastepujacym po nim LOCK sa rownowazne full barrier
natomiast LOCK z nastepujacym po nim UNLOCK juz nie

UWAGA!

Nastepstwem tego ze LOCK i UNLOCK moga byc barierami 'w jedna strone' jest to
ze efekty instrukcji spoza sekcji krytycznej moga 'przesiaknac' do wnetrza sekcji
krytycznej

MEMORY BARRIERS w LINUX KERNEL 5

LOCKING FUNCTIONS

- 1 implikacje operacji LOCK
- 2 implikacje operacji UNLOCK
- 3 implikacje LOCK(1) ; LOCK(2)
- 4 implikacje LOCK ; UNLOCK
- 5 implikacje LOCK'a który sie nie powiodl

wiec (1),(2),(4) UNLOCK z nastepujacym po nim LOCK sa rownowazne full barrier
natomiast LOCK z nastepujacym po nim UNLOCK juz nie

UWAGA!

Nastepstwem tego ze LOCK i UNLOCK moga byc barierami 'w jedna strone' jest to
ze efekty instrukcji spoza sekcji krytycznej moga 'przesiaknac' do wnetrza sekcji
krytycznej

MEMORY BARRIERS w LINUX KERNEL 5

LOCKING FUNCTIONS

- 1 implikacje operacji LOCK
- 2 implikacje operacji UNLOCK
- 3 implikacje LOCK(1) ; LOCK(2)
- 4 implikacje LOCK ; UNLOCK
- 5 implikacje LOCK'a który sie nie powiodl

wiec (1),(2),(4) UNLOCK z nastepujacym po nim LOCK sa rownowazne full barrier
natomiast LOCK z nastepujacym po nim UNLOCK juz nie

UWAGA!

Nastepstwem tego ze LOCK i UNLOCK moga byc barierami 'w jedna strone' jest to
ze efekty instrukcji spoza sekcji krytycznej moga 'przesiaknac' do wnetrza sekcji
krytycznej

MEMORY BARRIERS w LINUX KERNEL 5

LOCKING FUNCTIONS

- ❶ implikacje operacji LOCK
- ❷ implikacje operacji UNLOCK
- ❸ implikacje LOCK(1) ; LOCK(2)
- ❹ implikacje LOCK ; UNLOCK
- ❺ implikacje LOCK'a który sie nie powiodl

wiec (1),(2),(4) UNLOCK z nastepujacym po nim LOCK sa rownowazne full barrier
natomiast LOCK z nastepujacym po nim UNLOCK juz nie

UWAGA!

Nastepstwem tego ze LOCK i UNLOCK moga byc barierami 'w jedna strone' jest to
ze efekty instrukcji spoza sekcji krytycznej moga 'przesiaknac' do wnetrza sekcji
krytycznej

MEMORY BARRIERS w LINUX KERNEL 5

LOCKING FUNCTIONS

- ❶ implikacje operacji LOCK
- ❷ implikacje operacji UNLOCK
- ❸ implikacje LOCK(1) ; LOCK(2)
- ❹ implikacje LOCK ; UNLOCK
- ❺ implikacje LOCK'a który sie nie powiodl

wiec (1),(2),(4) UNLOCK z następującym po nim LOCK sa rownowazne full barrier
natomiast LOCK z następującym po nim UNLOCK juz nie

UWAGA!

Następstwem tego ze LOCK i UNLOCK moga byc barierami 'w jedna strone' jest to
ze efekty instrukcji spoza sekcji krytycznej moga 'przesiaknac' do wnetrza sekcji
krytycznej

MEMORY BARRIERS w LINUX KERNEL 5

LOCKING FUNCTIONS

- ❶ implikacje operacji LOCK
- ❷ implikacje operacji UNLOCK
- ❸ implikacje LOCK(1) ; LOCK(2)
- ❹ implikacje LOCK ; UNLOCK
- ❺ implikacje LOCK'a który sie nie powiodl

wiec (1),(2),(4) UNLOCK z nastepujacym po nim LOCK sa rownowazne full barrier
natomiast LOCK z nastepujacym po nim UNLOCK juz nie

UWAGA!

Nastepstwem tego ze LOCK i UNLOCK moga byc barierami 'w jedna strone' jest to
ze efekty instrukcji spoza sekcji krytycznej moga 'przesiaknac' do wnetrza sekcji
krytycznej

MEMORY BARRIERS w LINUX KERNEL 6

LOCKING FUNCTIONS - przykłady

```
*A = a; *B = b;  
LOCK  
*C = c; *D = d;  
UNLOCK  
*E = e; *F = f;
```

Następująca kolejność zdarzeń jest możliwa

```
LOCK, {*F,*A}, *E, {*C,*D}, *B, UNLOCK
```

{*F,*A} oznacza jednoczesny dostęp.

Ale żadna z tych już nie

{*F,*A}, *B,	LOCK, *C, *D,	UNLOCK, *E
*A, *B, *C,	LOCK, *D,	UNLOCK, *E, *F
*A, *B,	LOCK, *C,	UNLOCK, *D, *E, *F
*B,	LOCK, *C, *D,	UNLOCK, {*F,*A}, *E

MEMORY BARRIERS w LINUX KERNEL 6

LOCKING FUNCTIONS - przykłady

```
*A = a; *B = b;  
LOCK  
*C = c; *D = d;  
UNLOCK  
*E = e; *F = f;
```

Następująca kolejność zdarzeń jest możliwa

LOCK, {*F,*A}, *E, {*C,*D}, *B, UNLOCK

{*F,*A} oznacza jednoczesny dostęp.

Ale żadna z tych już nie

{*F,*A}, *B,	LOCK, *C, *D,	UNLOCK, *E
*A, *B, *C,	LOCK, *D,	UNLOCK, *E, *F
*A, *B,	LOCK, *C,	UNLOCK, *D, *E, *F
*B,	LOCK, *C, *D,	UNLOCK, {*F,*A}, *E

MEMORY BARRIERS w LINUX KERNEL 6

LOCKING FUNCTIONS - przykłady

```
*A = a; *B = b;  
LOCK  
*C = c; *D = d;  
UNLOCK  
*E = e; *F = f;
```

Następująca kolejność zdarzeń jest możliwa

```
LOCK, {*F,*A}, *E, {*C,*D}, *B, UNLOCK
```

{*F,*A} oznacza jednoczesny dostęp.

Ale żadna z tych już nie

{*F,*A}, *B,	LOCK, *C, *D,	UNLOCK, *E
*A, *B, *C,	LOCK, *D,	UNLOCK, *E, *F
*A, *B,	LOCK, *C,	UNLOCK, *D, *E, *F
*B,	LOCK, *C, *D,	UNLOCK, {*F,*A}, *E

Plan Prezentacji

- 1 Część1 - Potokowanie, WMB (Paweł Bedyński)
- 2 Część 2 - Architektury 32-bitowe i 64-bitowe (Krzysztof Kąs)
- 3 Część 3 VLIW + Multiprocessing (Krzysztof Niemkiewicz)
- 4 Część 4 - Multiprocessing cd. (Krzysztof Niemkiewicz)

Czy 64-bitowe adresowanie to dzisiaj konieczność?



Czy 64-bitowe adresowanie to dzisiaj konieczność?

<i>system operacyjny</i>	<i>pamięć dostępna dla jednej aplikacji</i>	<i>maksymalna obsługiwana wielkość pamięci RAM</i>
<i>Windows XP Professional (32-bit)</i>	<i>4 GB</i>	<i>4 GB</i>
<i>Windows XP Professional (64-bit)</i>	<i>16 TB</i>	<i>128 GB</i>
<i>Windows Server 2003 Standard (32-bit)</i>	<i>4 GB</i>	<i>4 GB</i>
<i>Windows Server 2003 Datacenter (32-bit)</i>	<i>4 GB</i>	<i>128 GB (przy wsparciu sprzętu)</i>
<i>Windows Server 2003 Datacenter (64-bit)</i>	<i>16 TB</i>	<i>512 GB</i>

Przydzielanie pamięci dla aplikacji w systemie Windows

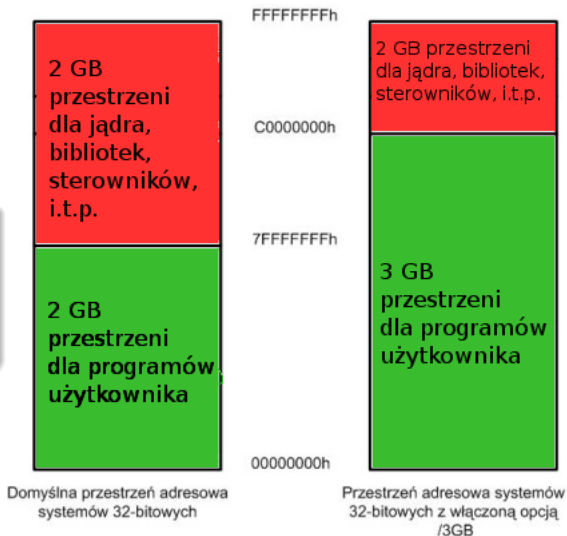
W 32-bitowym Windowsie
mamy standardowo do
dyspozycji jeszcze mniej niż 4
GB pamięci...

... bo system operacyjny też
jej potrzebuje.

Przydzielanie pamięci dla aplikacji w systemie Windows

W 32-bitowym Windowsie mamy standardowo do dyspozycji jeszcze mniej niż 4 GB pamięci...

... bo system operacyjny też jej potrzebuje.



Przydzielanie pamięci dla aplikacji w systemie Windows

W 32-bitowym Windowsie mamy standardowo do dyspozycji jeszcze mniej niż 4 GB pamięci...

... bo system operacyjny też jej potrzebuje.

Istnieją jednak mechanizmy pozwalające zaadresować dużo więcej niż 4 GB, bez konieczności przechodzenia na tryb 64-bitowy.

PAE - Physical Address Extension

- wprowadzone w Pentium Pro
- rozszerza fizyczną przestrzeń adresową w 32-bitowej architekturze poprzez zastosowanie 36-bitowych adresów
- wymagane wsparcie systemu operacyjnego (np. Windows XP, Windows Vista, Windows Server, Linux 2.6)
- można zaadresować do 64 GB pamięci RAM

PAE - Physical Address Extension

- wprowadzone w Pentium Pro
- rozszerza fizyczną przestrzeń adresową w 32-bitowej architekturze poprzez zastosowanie 36-bitowych adresów
- wymagane wsparcie systemu operacyjnego (np. Windows XP, Windows Vista, Windows Server, Linux 2.6)
- można zaadresować do 64 GB pamięci RAM

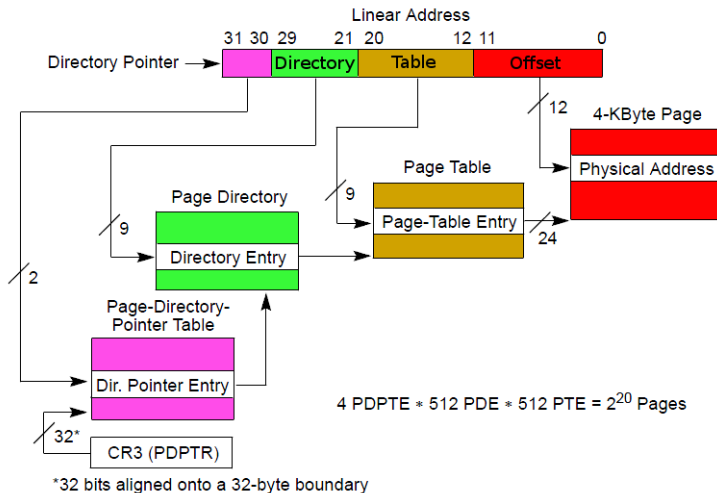
PAE - Physical Address Extension

- wprowadzone w Pentium Pro
- rozszerza fizyczną przestrzeń adresową w 32-bitowej architekturze poprzez zastosowanie 36-bitowych adresów
- wymagane wsparcie systemu operacyjnego (np. Windows XP, Windows Vista, Windows Server, Linux 2.6)
- można zaadresować do 64 GB pamięci RAM

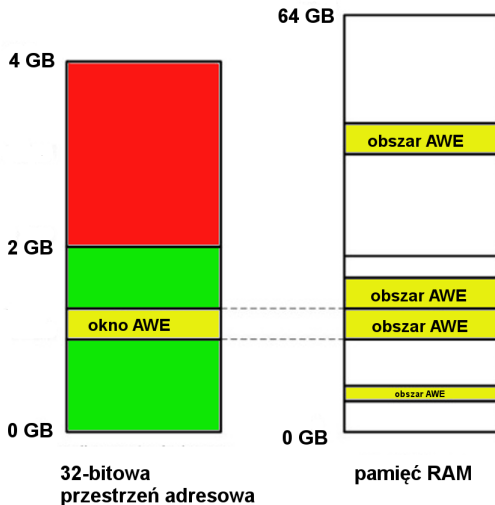
PAE - Physical Address Extension

- wprowadzone w Pentium Pro
- rozszerza fizyczną przestrzeń adresową w 32-bitowej architekturze poprzez zastosowanie 36-bitowych adresów
- wymagane wsparcie systemu operacyjnego (np. Windows XP, Windows Vista, Windows Server, Linux 2.6)
- można zaadresować do 64 GB pamięci RAM

PAE - trójpzomowe stronicowanie



AWE - Address Windowing Extensions



AWE - Address Windowing Extensions

Ograniczenia:

- procesy nie mogą współdzielić stron pamięci
- nie można tej samej ramce w pamięci RAM przypisać kilku różnych adresów logicznych w ramach jednego procesu
- jeśli wiele aplikacji używa AWE, to mogą one wyczerpać całą dostępną pamięć fizyczną, uniemożliwiając kolejnym korzystanie z tego mechanizmu
- aplikacje używające AWE nie mogą być emulowane na innej platformie.

AWE - Address Windowing Extensions

Ograniczenia:

- procesy nie mogą współdzielić stron pamięci
- nie można tej samej ramce w pamięci RAM przypisać kilku różnych adresów logicznych w ramach jednego procesu
- jeśli wiele aplikacji używa AWE, to mogą one wyczerpać całą dostępną pamięć fizyczną, uniemożliwiając kolejnym korzystanie z tego mechanizmu
- aplikacje używające AWE nie mogą być emulowane na innej platformie.

AWE - Address Windowing Extensions

Ograniczenia:

- procesy nie mogą współdzielić stron pamięci
- nie można tej samej ramce w pamięci RAM przypisać kilku różnych adresów logicznych w ramach jednego procesu
- jeśli wiele aplikacji używa AWE, to mogą one wyczerpać całą dostępną pamięć fizyczną, uniemożliwiając kolejnym korzystanie z tego mechanizmu
- aplikacje używające AWE nie mogą być emulowane na innej platformie.

AWE - Address Windowing Extensions

Ograniczenia:

- procesy nie mogą współdzielić stron pamięci
- nie można tej samej ramce w pamięci RAM przypisać kilku różnych adresów logicznych w ramach jednego procesu
- jeśli wiele aplikacji używa AWE, to mogą one wyczerpać całą dostępną pamięć fizyczną, uniemożliwiając kolejnym korzystanie z tego mechanizmu
- aplikacje używające AWE nie mogą być emulowane na innej platformie.

AWE - Address Windowing Extensions

Ograniczenia:

- procesy nie mogą współdzielić stron pamięci
- nie można tej samej ramce w pamięci RAM przypisać kilku różnych adresów logicznych w ramach jednego procesu
- jeśli wiele aplikacji używa AWE, to mogą one wyczerpać całą dostępną pamięć fizyczną, uniemożliwiając kolejnym korzystanie z tego mechanizmu
- aplikacje używające AWE nie mogą być emulowane na innej platformie.

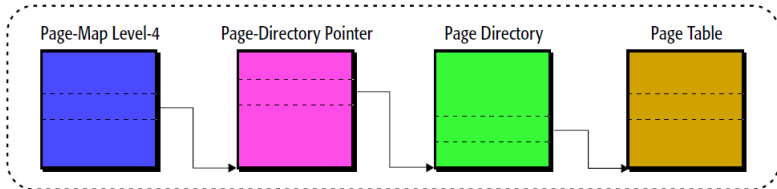
AWE - Address Windowing Extensions

Alternatywą w Linuksie dla AWE jest funkcja systemowa `mmap()` (która mapuje pliki lub urządzenia do pamięci).

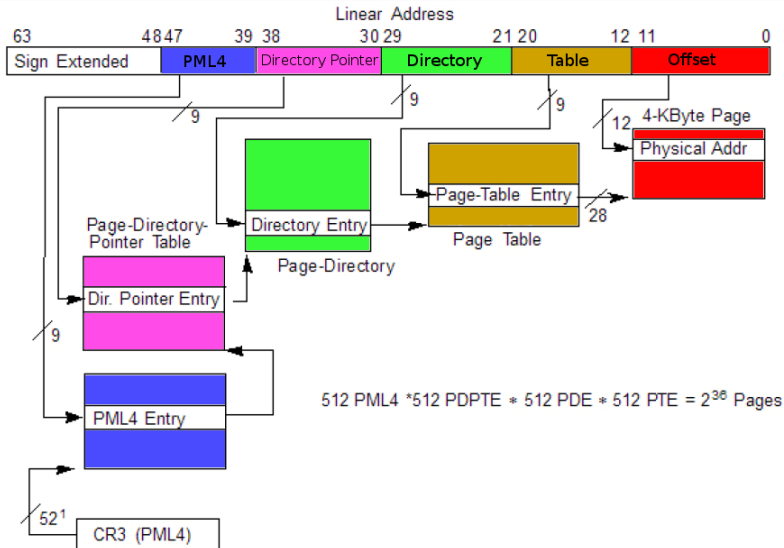
Adresowanie w trybie 64-bitowym

Adresowanie w trybie 64-bitowym

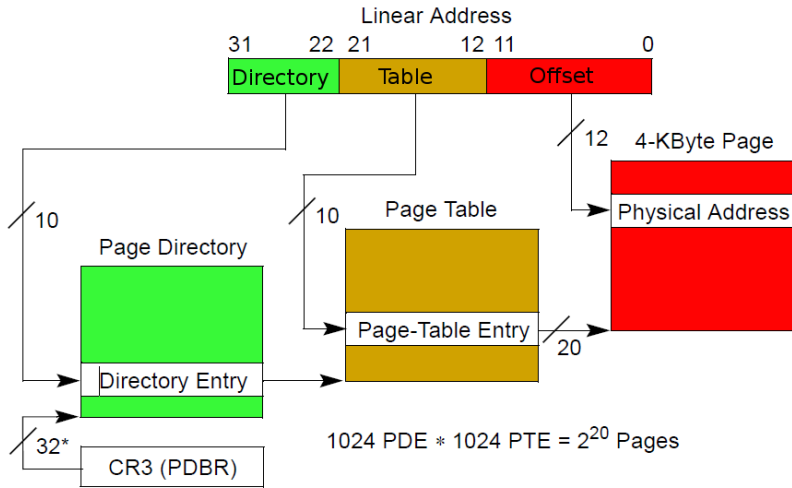
W trybie IA-32e, przy stronie 4 KB, stosuje się czteropoziomowe stronicowanie:



Adresowanie w trybie 64-bitowym



Dla porównania w trybie 32-bitowym:



*32 bits aligned onto a 4-KByte boundary.

64-bity a segmentacja

- istnieje "płaska" 64-bitowa liniowa przestrzeń adresowa, bez segmentacji
- rejestry CS, DS, ES, SS są traktowane jako zero
- rejestry FS, GS pełnią rolę tylko pomocniczą

64-bity a segmentacja

- istnieje "płaska" 64-bitowa liniowa przestrzeń adresowa, bez segmentacji
- rejestry CS, DS, ES, SS są traktowane jako zero
- rejestry FS, GS pełnią rolę tylko pomocniczą

64-bity a segmentacja

- istnieje "płaska" 64-bitowa liniowa przestrzeń adresowa, bez segmentacji
- rejestry CS, DS, ES, SS są traktowane jako zero
- rejestry FS, GS pełnią rolę tylko pomocniczą

Wielkość stron

Domyślnie 4 KB (Small Page Size).

Jest możliwa większa wartość w trybie PSE (Page Size Extensions):

- 4 MB dla x86
- 2 MB dla x64 i x86 w trybie PAE

Wielkość stron

Domyślnie 4 KB (Small Page Size).

Jest możliwa większa wartość w trybie PSE (Page Size Extensions):

- 4 MB dla x86
- 2 MB dla x64 i x86 w trybie PAE

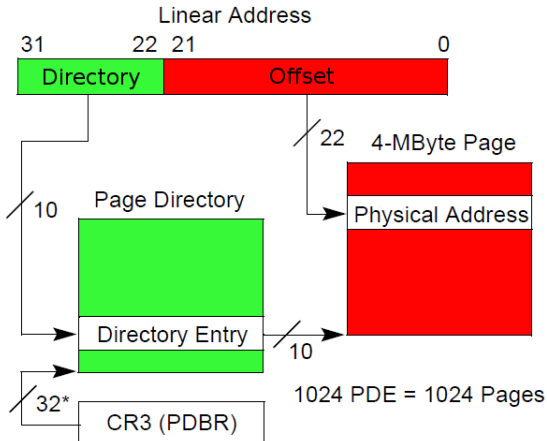
Wielkość stron

Domyślnie 4 KB (Small Page Size).

Jest możliwa większa wartość w trybie PSE (Page Size Extensions):

- 4 MB dla x86
- 2 MB dla x64 i x86 w trybie PAE

Wielkość stron - PSE w trybie 32-bitowym:



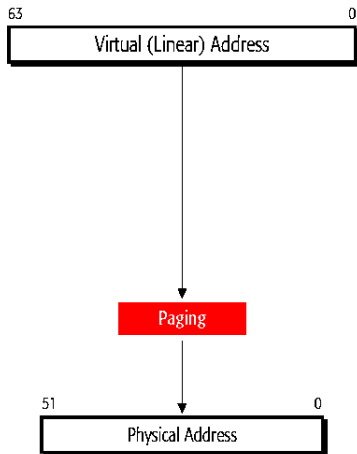
*32 bits aligned onto a 4-KByte boundary.

Tryby pracy procesora 64-bitowego

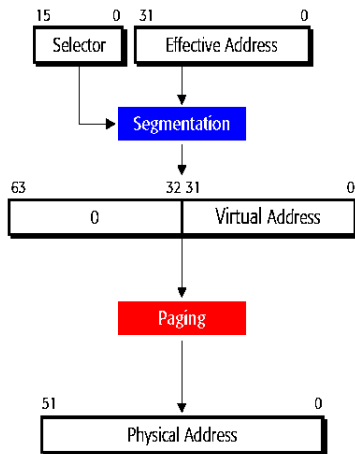
- long mode
- legacy mode

long mode

64-Bit Mode

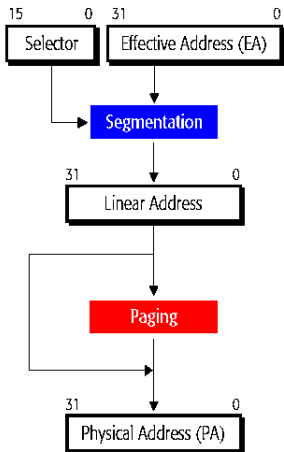


Compatibility Mode

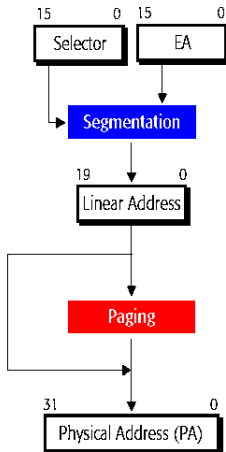


legacy mode

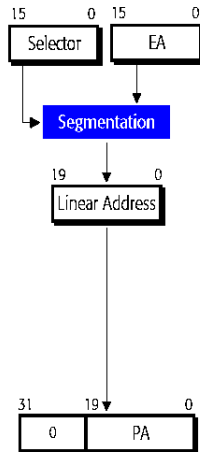
Protected Mode



Virtual-8086 Mode



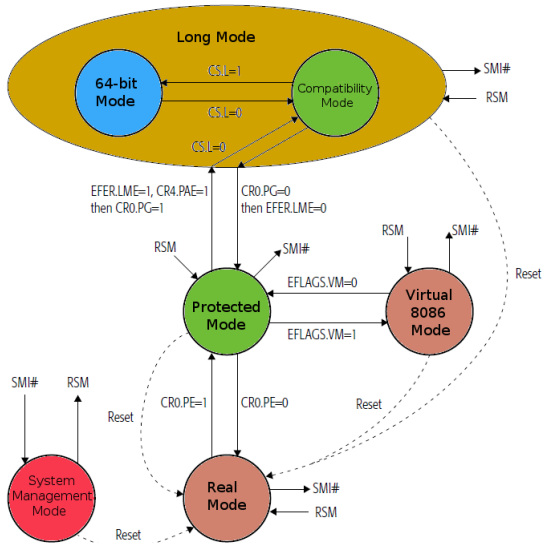
Real Mode



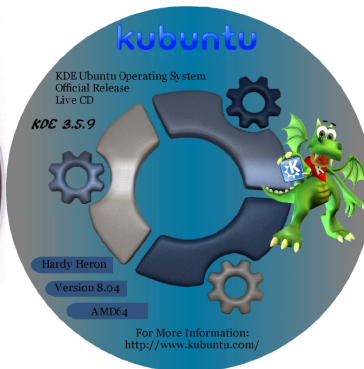
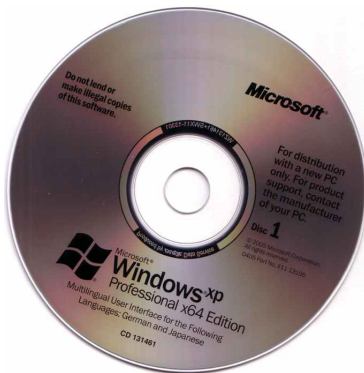
Tryby pracy procesora 64-bitowego - podsumowanie

Operating Mode		Operating System Required	Application Recompile Required	Defaults		Register Extensions	Typical
				Address Size (bits)	Operand Size (bits)		GPR Width (bits)
Long Mode	64-Bit Mode	64-bit OS	yes	64	32	yes	64
	Compatibility Mode		no	32		no	32
				16	16		16
Legacy Mode	Protected Mode	Legacy 32-bit OS	no	32	32	no	32
	Virtual-8086 Mode			16	16		
	Real Mode	Legacy 16-bit OS		16	16	no	16

Tryby pracy procesora 64-bitowego - podsumowanie



64-bitowe systemy operacyjne - przykłady



64-bitowe edycje systemu Windows ...

- ... mają problemy z 32-bitowymi sterownikami
- dlatego logo "Certified for Vista" może uzyskać tylko sprzęt, do którego zostaną dostarczone sterowniki 32 i 64-bitowe

64-bitowe edycje systemu Windows ...

... posiadają następujące ograniczenia pamięciowe:

- maksymalnie 16 TB logicznej przestrzeni adresowej (8 TB dla procesów jądra i 8 TB dla procesów użytkownika)
- nie więcej niż 128 GB (Windows XP) lub 1 TB (Windows Server 2003) pamięci RAM

64-bitowe edycje systemu Windows ...

... umożliwiają skorzystanie z technologii WOW64 (Windows-on-Windows 64-bit)

- ale nie można łączyć kodu 32-bitowego z 64-bitowym

64-bitowy Linux

- był jednym z pierwszych 64-bitowych systemów operacyjnych
- jest w nim możliwa 64-bitowa rekompilacja programów, które wykorzystują inne, 32-bitowe programy

64-bitowy Linux

- był jednym z pierwszych 64-bitowych systemów operacyjnych
- jest w nim możliwa 64-bitowa rekompilacja programów, które wykorzystują inne, 32-bitowe programy

64-bitowy Mac OS X

- jądro uruchamiane jest jako 32-bitowy proces, dzięki czemu działają 32-bitowe sterowniki urządzeń
- jednocześnie wspierane są 64-bitowe procesy użytkownika
- technologia "universal binary" - wszystkie wersje programu w jednym pliku

64-bitowy Mac OS X

- jądro uruchamiane jest jako 32-bitowy proces, dzięki czemu działają 32-bitowe sterowniki urządzeń
- jednocześnie wspierane są 64-bitowe procesy użytkownika
- technologia "universal binary" - wszystkie wersje programu w jednym pliku

64-bitowy Mac OS X

- jądro uruchamiane jest jako 32-bitowy proces, dzięki czemu działają 32-bitowe sterowniki urządzeń
- jednocześnie wspierane są 64-bitowe procesy użytkownika
- technologia "universal binary" - wszystkie wersje programu w jednym pliku



Spójność pamięci cache w procesorach Intel'a i AMD

- Procesor nie zapewnia, że dane znajdujące się w jego pamięci podręcznej są zawsze spójne z pamięcią RAM.
- Musi więc to zapewnić system operacyjny.
- Pomocna jest instrukcja WBINVD.

Spójność pamięci cache w procesorach Intel'a i AMD

- Procesor nie zapewnia, że dane znajdujące się w jego pamięci podręcznej są zawsze spójne z pamięcią RAM.
- Musi więc to zapewnić system operacyjny.
- Pomocna jest instrukcja WBINVD.

Spójność pamięci cache w procesorach Intel i AMD

- Procesor nie zapewnia, że dane znajdujące się w jego pamięci podręcznej są zawsze spójne z pamięcią RAM.
- Musi więc to zapewnić system operacyjny.
- Pomocna jest instrukcja WBINVD.

MTRR - memory type range registers

- Rejestry MTRR pozwalają powiązać typ pamięci z zakresem przestrzeni adresowej.
- Dzięki temu procesor może optymalizować operacje na różnych typach pamięci, np. RAM, ROM, mapowane do pamięci urządzenia I/O.
- Strong Uncacheable (UC) - lokacje w pamięci nie są nigdy sprowadzane do pamięci cache.
- W systemach wieloprocessorowych system operacyjny musi utrzymywać spójność MTRR (bo oczywiście różne procesory muszą używać tych samych wartości MTRR).
- Procesory AMD i Intelu nie zapewniają bowiem sprzętowego wsparcia dla utrzymania tej spójności.

MTRR - memory type range registers

- Rejestry MTRR pozwalają powiązać typ pamięci z zakresem przestrzeni adresowej.
- Dzięki temu procesor może optymalizować operacje na różnych typach pamięci, np. RAM, ROM, mapowane do pamięci urządzenia I/O.
- Strong Uncacheable (UC) - lokacje w pamięci nie są nigdy sprowadzane do pamięci cache.
- W systemach wieloprocessorowych system operacyjny musi utrzymywać spójność MTRR (bo oczywiście różne procesory muszą używać tych samych wartości MTRR).
- Procesory AMD i Intelu nie zapewniają bowiem sprzętowego wsparcia dla utrzymania tej spójności.

MTRR - memory type range registers

- Rejestry MTRR pozwalają powiązać typ pamięci z zakresem przestrzeni adresowej.
- Dzięki temu procesor może optymalizować operacje na różnych typach pamięci, np. RAM, ROM, mapowane do pamięci urządzenia I/O.
- Strong Uncacheable (UC) - lokacje w pamięci nie są nigdy sprowadzane do pamięci cache.
- W systemach wieloprocesorowych system operacyjny musi utrzymywać spójność MTRR (bo oczywiście różne procesory muszą używać tych samych wartości MTRR).
- Procesory AMD i Intelu nie zapewniają bowiem sprzętowego wsparcia dla utrzymania tej spójności.

MTRR - memory type range registers

- Rejestry MTRR pozwalają powiązać typ pamięci z zakresem przestrzeni adresowej.
- Dzięki temu procesor może optymalizować operacje na różnych typach pamięci, np. RAM, ROM, mapowane do pamięci urządzenia I/O.
- Strong Uncacheable (UC) - lokacje w pamięci nie są nigdy sprowadzane do pamięci cache.
- W systemach wieloprocessorowych system operacyjny musi utrzymywać spójność MTRR (bo oczywiście różne procesory muszą używać tych samych wartości MTRR).
- Procesory AMD i Intelu nie zapewniają bowiem sprzętowego wsparcia dla utrzymania tej spójności.

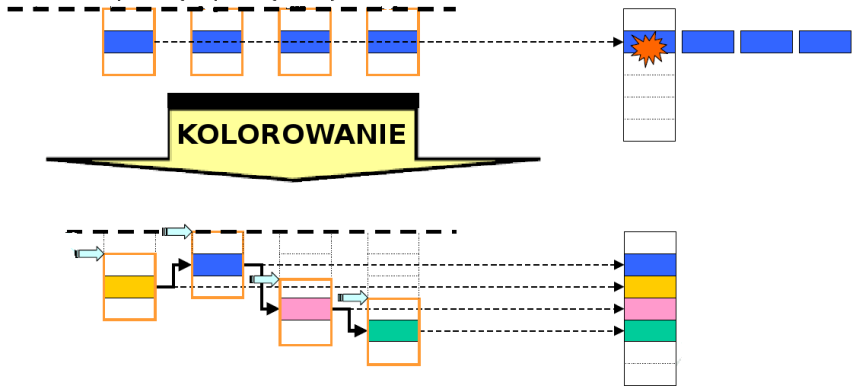
MTRR - memory type range registers

- Rejestry MTRR pozwalają powiązać typ pamięci z zakresem przestrzeni adresowej.
- Dzięki temu procesor może optymalizować operacje na różnych typach pamięci, np. RAM, ROM, mapowane do pamięci urządzenia I/O.
- Strong Uncacheable (UC) - lokacje w pamięci nie są nigdy sprowadzane do pamięci cache.
- W systemach wieloprocessorowych system operacyjny musi utrzymywać spójność MTRR (bo oczywiście różne procesory muszą używać tych samych wartości MTRR).
- Procesory AMD i Intel'a nie zapewniają bowiem sprzętowego wsparcia dla utrzymania tej spójności.

Kolorowanie - przykład: przyspieszanie schedulera w Linuksie

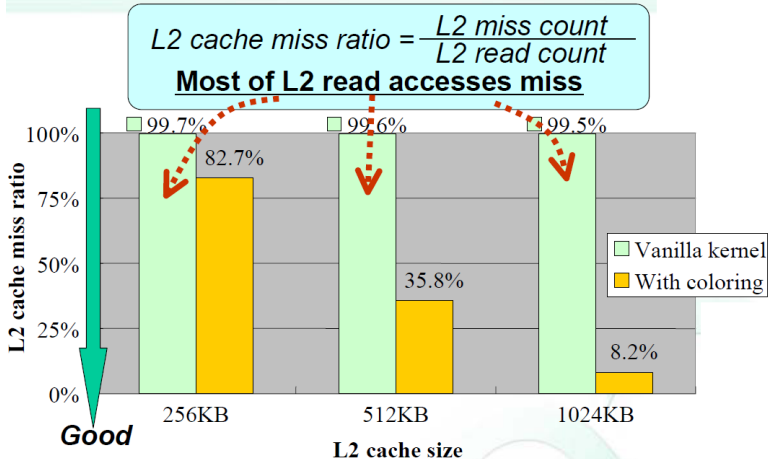
Kolorowanie - przykład: przyspieszanie schedulera w Linuksie

wyrównanie danych do wielokrotności
wielkości pamięci podręcznej

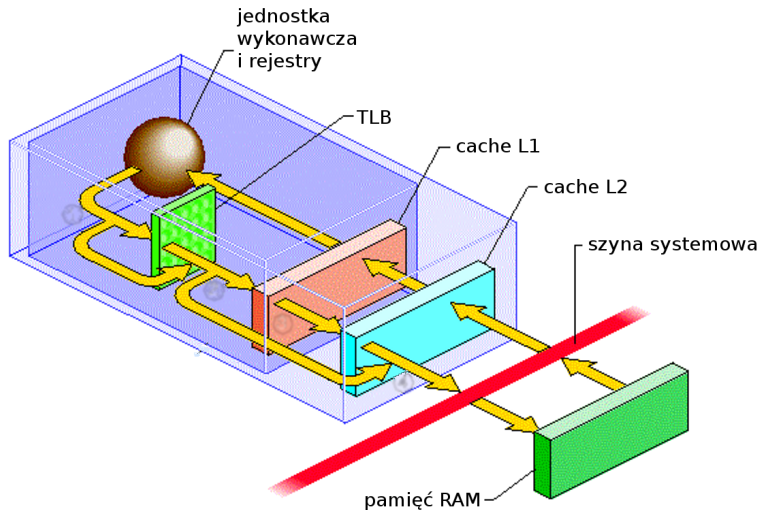


Kolorowanie - przykład: przyspieszanie schedulera w Linuksie

Różnica w wydajności jest szczególnie widoczna w przypadku Linuksa 2.4 (test przeprowadzony w laboratorium firmy Fujitsu)



TLB - przypomnienie



TLB - spójność

- za każdym razem, kiedy katalog lub tablica stron ulega zmianie, system operacyjny musi natychmiast unieważnić odpowiedni wpis w TLB (poprzez instrukcję INVLPG)
- konieczność modyfikacji TLB zachodzi między innymi podczas przełączania kontekstu

System operacyjny może radzić sobie z tym problemem na dwa sposoby:

- przez opróżnianie całej zawartości TLB przy zmianie kontekstu
- przez zapisywanie identyfikatorów przestrzeni adresowej (ASID - address space identifiers) razem z każdym wpisem w TLB.

TLB - spójność

- za każdym razem, kiedy katalog lub tablica stron ulega zmianie, system operacyjny musi natychmiast unieważnić odpowiedni wpis w TLB (poprzez instrukcję INVLPG)
- konieczność modyfikacji TLB zachodzi między innymi podczas przełączania kontekstu

System operacyjny może radzić sobie z tym problemem na dwa sposoby:

- przez opróżnianie całej zawartości TLB przy zmianie kontekstu
- przez zapisywanie identyfikatorów przestrzeni adresowej (ASID - address space identifiers) razem z każdym wpisem w TLB.

TLB - spójność

- za każdym razem, kiedy katalog lub tablica stron ulega zmianie, system operacyjny musi natychmiast unieważnić odpowiedni wpis w TLB (poprzez instrukcję INVLPG)
- konieczność modyfikacji TLB zachodzi między innymi podczas przełączania kontekstu

System operacyjny może radzić sobie z tym problemem na dwa sposoby:

- przez opróżnianie całej zawartości TLB przy zmianie kontekstu
- przez zapisywanie identyfikatorów przestrzeni adresowej (ASID - address space identifiers) razem z każdym wpisem w TLB.

TLB - spójność

- za każdym razem, kiedy katalog lub tablica stron ulega zmianie, system operacyjny musi natychmiast unieważnić odpowiedni wpis w TLB (poprzez instrukcję INVLPG)
- konieczność modyfikacji TLB zachodzi między innymi podczas przełączania kontekstu

System operacyjny może radzić sobie z tym problemem na dwa sposoby:

- przez opróżnianie całej zawartości TLB przy zmianie kontekstu
- przez zapisywanie identyfikatorów przestrzeni adresowej (ASID - address space identifiers) razem z każdym wpisem w TLB.

TLB - spójność

- za każdym razem, kiedy katalog lub tablica stron ulega zmianie, system operacyjny musi natychmiast unieważnić odpowiedni wpis w TLB (poprzez instrukcję INVLPG)
- konieczność modyfikacji TLB zachodzi między innymi podczas przełączania kontekstu

System operacyjny może radzić sobie z tym problemem na dwa sposoby:

- przez opróżnianie całej zawartości TLB przy zmianie kontekstu
- przez zapisywanie identyfikatorów przestrzeni adresowej (ASID - address space identifiers) razem z każdym wpisem w TLB.

TLB - spójność

- za każdym razem, kiedy katalog lub tablica stron ulega zmianie, system operacyjny musi natychmiast unieważnić odpowiedni wpis w TLB (poprzez instrukcję INVLPG)
- konieczność modyfikacji TLB zachodzi między innymi podczas przełączania kontekstu

System operacyjny może radzić sobie z tym problemem na dwa sposoby:

- przez opróżnianie całej zawartości TLB przy zmianie kontekstu
- przez zapisywanie identyfikatorów przestrzeni adresowej (ASID - address space identifiers) razem z każdym wpisem w TLB.

Obsługa chybień w TLB

Są dwie możliwości:

- sprzętowe zarządzanie TLB (np. w procesorach IA32) - chybiecie obsługiwane przez procesor
- programowe zarządzanie TLB (np. w procesorach IA64) - chybiecie przekazywane jako wyjątek systemowi operacyjnemu

Obsługa chybień w TLB

Są dwie możliwości:

- sprzętowe zarządzanie TLB (np. w procesorach IA32) - chybiecie obsługiwane przez procesor
- programowe zarządzanie TLB (np. w procesorach IA64) - chybiecie przekazywane jako wyjątek systemowi operacyjnemu

Obsługa chybień w TLB

Są dwie możliwości:

- sprzętowe zarządzanie TLB (np. w procesorach IA32) - chybiecie obsługiwane przez procesor
- programowe zarządzanie TLB (np. w procesorach IA64) - chybiecie przekazywane jako wyjątek systemowi operacyjnemu

Zasięg bufora TLB

- Zwiększając n -krotnie rozmiar strony, zwiększamy n -krotnie zasięg TLB ...
- ... w rezultacie osiągając wyższy współczynnik trafień w TLB.
- Windows XP dla procesorów zgodnych z IA32 używa zawsze stron wielkości 4 KB.

Zasięg bufora TLB

- Zwiększając n -krotnie rozmiar strony, zwiększamy n -krotnie zasięg TLB ...
- ... w rezultacie osiągając wyższy współczynnik trafień w TLB.
- Windows XP dla procesorów zgodnych z IA32 używa zawsze stron wielkości 4 KB.

Zasięg bufora TLB

- Zwiększając n -krotnie rozmiar strony, zwiększamy n -krotnie zasięg TLB ...
- ... w rezultacie osiągając wyższy współczynnik trafień w TLB.
- Windows XP dla procesorów zgodnych z IA32 używa zawsze stron wielkości 4 KB.

Zasięg bufora TLB

- Zwiększając n -krotnie rozmiar strony, zwiększamy n -krotnie zasięg TLB ...
- ... w rezultacie osiągając wyższy współczynnik trafień w TLB.
- Windows XP dla procesorów zgodnych z IA32 używa zawsze stron wielkości 4 KB.

Plan Prezentacji

- 1 Część1 - Potokowanie, WMB (Paweł Bedyński)
- 2 Część 2 - Architektury 32-bitowe i 64-bitowe (Krzysztof Kąs)
- 3 Część 3 VLIW + Multiprocessing (Krzysztof Niemkiewicz)
- 4 Część 4 - Multiprocessing cd. (Krzysztof Niemkiewicz)

Very Long Instruction Word

Bardzo długie instrukcje

- W jednej instrukcji kodujemy wiele atomowych operacji które procesor wykona równolegle
- Przykład (arch. SHARC - Super Harvard Architecture Single-Chip Computer)

```
r12=r0+r4, r6=r0+r12, r0=da(r0,r0), r4=pa(r6,r0);
```

- Zbiór możliwych instrukcji zależy od liczby jednostek wykonawczych
- Rozwiązywanie problemów związanych z zrównoleglaniem kodu, wyborem które instrukcje będą wykonywane jako następne (if-y) itp jest przerzucone na kompilator

Zalety

- Uproszczony hardware
- Zapewne lepsze przewidywania co do wyboru gałęzi kodu
- Jeśli mamy dobry kompilator to powinniśmy mieć lepsze rezultaty niż dla potokowania
- Niezgodność z innymi architekturami, złożone kompilatory
- Często instrukcje są w znacznej mierze puste

Very Long Instruction Word

Bardzo długie instrukcje

- W jednej instrukcji kodujemy wiele atomowych operacji które procesor wykona równolegle
- Przykład (arch. SHARC - Super Harvard Architecture Single-Chip Computer)

```
f12=f0*f4, f8=f8+f12, f0=dm(i0,m3), f4=pm(i8,m9);
```

- Zbiór możliwych instrukcji zależy od liczby jednostek wykonawczych
- Rozwiązywanie problemów związanych z zrównoleglaniem kodu, wyborem które instrukcje będą wykonywane jako następne (if-y) itp jest przerzucone na kompilator

Zalety

- Uproszczony hardware
- Zapewne lepsze przewidywania co do wyboru gałęzi kodu
- Jeśli mamy dobry kompilator to powinniśmy mieć lepsze rezultaty niż dla potokowania
- Niezgodność z innymi architekturami, złożone kompilatory
- Czasem instrukcje są w znacznej mierze puste

Very Long Instruction Word

Bardzo długie instrukcje

- W jednej instrukcji kodujemy wiele atomowych operacji które procesor wykona równolegle
- Przykład (arch. SHARC - Super Harvard Architecture Single-Chip Computer)

```
f12=f0*f4, f8=f8+f12, f0=dm(i0,m3), f4=pm(i8,m9);
```

- Zbiór możliwych instrukcji zależy od liczby jednostek wykonawczych
- Rozwiązywanie problemów związanych z zrównoleglaniem kodu, wyborem które instrukcje będą wykonywane jako następne (if-y) itp jest przerzucone na kompilator

Zalety

- Uproszczony hardware
- Zapewne lepsze przewidywania co do wyboru gałęzi kodu
- Jeśli mamy dobry kompilator to powinniśmy mieć lepsze rezultaty niż dla potokowania
- Niezgodność z innymi architekturami, złożone kompilatory
- Czasem instrukcje są w znacznej mierze puste

Very Long Instruction Word

Bardzo długie instrukcje

- W jednej instrukcji kodujemy wiele atomowych operacji które procesor wykona równolegle
- Przykład (arch. SHARC - Super Harvard Architecture Single-Chip Computer)

```
f12=f0*f4, f8=f8+f12, f0=dm(i0,m3), f4=pm(i8,m9);
```

- Zbiór możliwych instrukcji zależy od liczby jednostek wykonawczych
- Rozwiązywanie problemów związanych z zrównoleglaniem kodu, wyborem które instrukcje będą wykonywane jako następne (if-y) itp jest przerzucone na kompilator

Zalety i wady

- Uproszczony hardware
- Zapewne lepsze przewidywania co do wyboru gałęzi kodu
- Jeśli mamy dobry kompilator to powinniśmy mieć lepsze rezultaty niż dla potokowania
- Niezgodność z innymi architekturami, złożone kompilatory
- Czasem instrukcje są w znacznej mierze puste

Very Long Instruction Word

Bardzo długie instrukcje

- W jednej instrukcji kodujemy wiele atomowych operacji które procesor wykona równolegle
- Przykład (arch. SHARC - Super Harvard Architecture Single-Chip Computer)

```
f12=f0*f4, f8=f8+f12, f0=dm(i0,m3), f4=pm(i8,m9);
```

- Zbiór możliwych instrukcji zależy od liczby jednostek wykonawczych
- Rozwiązywanie problemów związanych z zrównoleglaniem kodu, wyborem które instrukcje będą wykonywane jako następne (if-y) itp jest przerzucone na kompilator

Zalety i wady

- Uproszczony hardware
- Zapewne lepsze przewidywania co do wyboru gałęzi kodu
- Jeśli mamy dobry kompilator to powinniśmy mieć lepsze rezultaty niż dla potokowania
- Niezgodność z innymi architekturami, złożone kompilatory
- Czasem instrukcje są w znacznej mierze puste

Wieloprocessorowość, wielordzeniowość

Wieloprocessorowość

- Komputery z wieloma procesorami
- Głównie superkomputery o naprawdę dużej liczbie procesorów i specjalnej budowie
- Specjalne SO
- Wieloprocessorowość jest raczej rzadko spotykana w komputerach klasy PC

Wielordzeniowość

- Dość nowe rozwiązanie
- Projektowane dla komputerów klasy PC
- Obecnie powszechnie stosowane
- Popularne SO muszą je obsługiwać aby pozostać popularnymi

Wieloprocessorowość, wielordzeniowość

Wieloprocessorowość

- Komputery z wieloma procesorami
- Głównie superkomputery o naprawdę dużej liczbie procesorów i specjalnej budowie
- Specjalne SO
- Wieloprocessorowość jest raczej rzadko spotykana w komputerach klasy PC

Wielordzeniowość

- Dość nowe rozwiązanie
- Projektowane dla komputerów klasy PC
- Obecnie powszechnie stosowane
- Popularne SO muszą je obsługiwać aby pozostać popularnymi

Wieloprocessorowość, wielordzeniowość

Wieloprocessorowość

- Komputery z wieloma procesorami
- Głównie superkomputery o naprawdę dużej liczbie procesorów i specjalnej budowie
- Specjalne SO
- Wieloprocessorowość jest raczej rzadko spotykana w komputerach klasy PC

Wielordzeniowość

- Dość nowe rozwiązanie
- Projektowane dla komputerów klasy PC
- Obecnie powszechnie stosowane
- Popularne SO muszą je obsługiwać aby pozostać popularnymi

Dlaczego wiele rdzeni?

Problemy z jednordzeniowymi procesorami

- Prawo Moore'a
- Granice możliwości technicznych

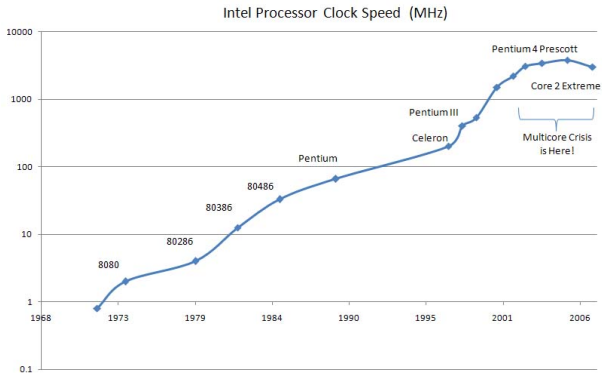
Wiele rdzeni

- Mniejsze wydzielanie ciepła i zużycie energii
- Teoretycznie n -razy szybszy procesor przy tej samej prędkości rdzenia
- Lepsze od kilku procesorów bo wspólny cache

Problemy

- Mało programów równoległych
- „Often it's cheaper to buy proper computer than to rewrite software“
- Zapychanie pasma dostępu do pamięci i urządzeń (np. karty sieciowej)

Dlaczego wiele rdzeni?



Prędkość pojedynczego procesora/rdzenia w funkcji czasu

Dlaczego wiele rdzeni?

Problemy z jednordzeniowymi procesorami

- Prawo Moore'a
- Granice możliwości technicznych

Wiele rdzeni

- Mniejsze wydzielanie ciepła i zużycie energii
- Teoretycznie n -razy szybszy procesor przy tej samej prędkości rdzenia
- Lepsze od kilku procesorów bo wspólny cache

Problemy

- Mało programów równoległych
- „Often it's cheaper to buy proper computer than to rewrite software“
- Zapychanie pasma dostępu do pamięci i urządzeń (np. karty sieciowej)

Dlaczego wiele rdzeni?

Problemy z jednordzeniowymi procesorami

- Prawo Moore'a
- Granice możliwości technicznych

Wiele rdzeni

- Mniejsze wydzielanie ciepła i zużycie energii
- Teoretycznie n-razy szybszy procesor przy tej samej prędkości rdzenia
- Lepsze od kilku procesorów bo wspólny cache

Problemy

- Mało programów równoległych
- „Often it's cheaper to buy proper computer than to rewrite software“
- Zapychanie pasma dostępu do pamięci i urządzeń (np. karty sieciowej)

Dlaczego wiele rdzeni?

Problemy z jednordzeniowymi procesorami

- Prawo Moore'a
- Granice możliwości technicznych

Wiele rdzeni

- Mniejsze wydzielanie ciepła i zużycie energii
- Teoretycznie n-razy szybszy procesor przy tej samej prędkości rdzenia
- Lepsze od kilku procesorów bo wspólny cache

Problemy

- Mało programów równoległych
- „Often it's cheaper to buy proper computer than to rewrite software“
- Zapychanie pasma dostępu do pamięci i urządzeń (np karty sieciowej)

Dlaczego wiele rdzeni?

Problemy z jednordzeniowymi procesorami

- Prawo Moore'a
- Granice możliwości technicznych

Wiele rdzeni

- Mniejsze wydzielanie ciepła i zużycie energii
- Teoretycznie n-razy szybszy procesor przy tej samej prędkości rdzenia
- Lepsze od kilku procesorów bo wspólny cache

Problemy

- Mało programów równoległych
- „Often it's cheaper to buy proper computer than to rewrite software“
- Zapychanie pasma dostępu do pamięci i urządzeń (np karty sieciowej)

Architektura asynchroniczna (1970-80)

AMP,ASMP,Master - Slave

- Każdy procesor jest traktowany niezależnie
- Procesory mogą być różnego przeznaczenia, różnie taktowane
- Często mają dostęp do różnych urządzeń peryferyjnych
- Procesy muszą specyfikować na którym procesorze mają pracować
- Łatwo o wąskie gardło, nierównomierne rozłożenie obciążenia
- Master - Slave czyli jeden główny procesor z SO, pozostałe procesory do szybkich obliczeń

Współczesne przykłady

- Procesor Cell(PS3 ale nie tylko...)
- GPU

Architektura asynchroniczna (1970-80)

AMP,ASMP,Master - Slave

- Każdy procesor jest traktowany niezależnie
- Procesory mogą być różnego przeznaczenia, różnie taktowane
- Często mają dostęp do różnych urządzeń preferencyjnych
- Procesy muszą specyfikować na którym procesorze mają pracować
- Łatwo o wąskie gardło, nierównomierne rozłożenie obciążenia
- Master - Slave czyli jeden główny procesor z SO, pozostałe procesory do szybkich obliczeń

Współczesne przykłady

- Procesor Cell(PS3 ale nie tylko...)
- GPU

Architektura asynchroniczna (1970-80)

AMP,ASMP,Master - Slave

- Każdy procesor jest traktowany niezależnie
- Procesory mogą być różnego przeznaczenia, różnie taktowane
- Często mają dostęp do różnych urządzeń preferencyjnych
- Procesy muszą specyfikować na którym procesorze mają pracować
- Łatwo o wąskie gardło, nierównomierne rozłożenie obciążenia
- Master - Slave czyli jeden główny procesor z SO, pozostałe procesory do szybkich obliczeń

Współczesne przykłady

- Procesor Cell(PS3 ale nie tylko...)
- GPU

Architektura asynchroniczna (1970-80)

AMP,ASMP,Master - Slave

- Każdy procesor jest traktowany niezależnie
- Procesory mogą być różnego przeznaczenia, różnie taktowane
- Często mają dostęp do różnych urządzeń preferencyjnych
- Procesy muszą specyfikować na którym procesorze mają pracować
- Łatwo o wąskie gardło, nierównomierne rozłożenie obciążenia
- Master - Slave czyli jeden główny procesor z SO, pozostałe procesory do szybkich obliczeń

Współczesne przykłady

- Procesor Cell(PS3 ale nie tylko...)
- GPU

Architektura asynchroniczna (1970-80)

AMP,ASMP,Master - Slave

- Każdy procesor jest traktowany niezależnie
- Procesory mogą być różnego przeznaczenia, różnie taktowane
- Często mają dostęp do różnych urządzeń preferencyjnych
- Procesy muszą specyfikować na którym procesorze mają pracować
- Łatwo o wąskie gardło, nierównomierne rozłożenie obciążenia
- Master - Slave czyli jeden główny procesor z SO, pozostałe procesory do szybkich obliczeń

Współczesne przykłady

- Procesor Cell(PS3 ale nie tylko...)
- GPU

Architektura asynchroniczna (1970-80)

AMP,ASMP,Master - Slave

- Każdy procesor jest traktowany niezależnie
- Procesory mogą być różnego przeznaczenia, różnie taktowane
- Często mają dostęp do różnych urządzeń preferencyjnych
- Procesy muszą specyfikować na którym procesorze mają pracować
- Łatwo o wąskie gardło, nierównomierne rozłożenie obciążenia
- Master - Slave czyli jeden główny procesor z SO, pozostałe procesory do szybkich obliczeń

Współczesne przykłady

- Procesor Cell(PS3 ale nie tylko...)
- GPU

SMP,BMP

SMP

- Wszystkie procesory traktowane równorzędnie
- Procesy mogą swobodnie migrować między procesorami
- SO powinien zapewniać równomierne obciążenie procesorów
- Obecnie najczęściej stosowane
- W linuxie od 1.2.33

Cache

- Poważnym problemem technicznym jest synchronizacja cache
- Zazwyczaj zapewniana sprzętowo, ale ...
- BMP - może jednak mieć wolności?

SMP,BMP

SMP

- Wszystkie procesory traktowane równorzędnie
- Procesy mogą swobodnie migrować między procesorami
- SO powinien zapewniać równomierne obciążenie procesorów
- Obecnie najczęściej stosowane
- W linuxie od 1.2.33

Cache

- Poważnym problemem technicznym jest synchronizacja cache
- Zazwyczaj zapewniana sprzętowo, ale ...
- BMP - może jednak mieć wolności?

SMP,BMP

SMP

- Wszystkie procesory traktowane równorzędnie
- Procesy mogą swobodnie migrować między procesorami
- SO powinien zapewniać równomierne obciążenie procesorów
- Obecnie najczęściej stosowane
- W linuxie od 1.2.33

Cache

- Poważnym problemem technicznym jest synchronizacja cache
- Zazwyczaj zapewniana sprzętowo, ale ...
- BMP - może jednak mieć wolności?

SMP,BMP

SMP

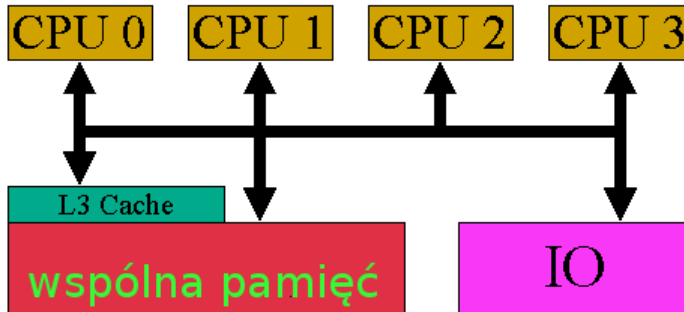
- Wszystkie procesory traktowane równorzędnie
- Procesy mogą swobodnie migrować między procesorami
- SO powinien zapewniać równomierne obciążenie procesorów
- Obecnie najczęściej stosowane
- W linuxie od 1.2.33

Cache

- Poważnym problemem technicznym jest synchronizacja cache
- Zazwyczaj zapewniana sprzętowo, ale ...
- BMP - może jednak mniej wolności?

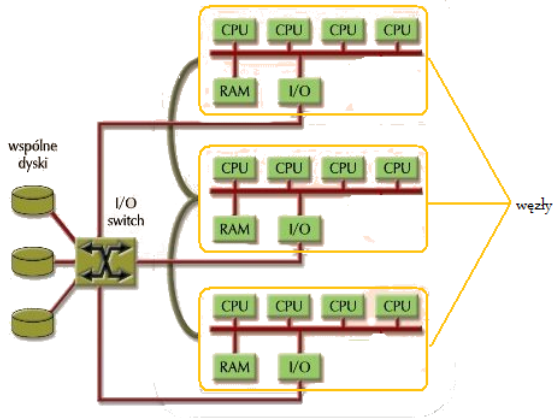
SMP

UMA, czyli Uniform Memory Access



NUMA, ccNUMA

Non-Uniform Memory Access lub Non-Uniform Memory Architecture, czasami też
DSM - Distributed Shared Memory



Obecnie w praktyce ccNUMA.

Utrzymywanie spójności pamięci podręcznych

Źródła kłopotów

- lokalne pamięci podręczne
- buforowanie operacji pisania

Modele spójności - podobnie jak na BD

- Spójność ścisła
- Spójność sekwencyjna
- Spójność słaba, na wyjściu, na wejściu

Implementacja

- Algorytm scentralizowanego zarządcy
- Dynamiczny algorytm zarządcy rozproszonego
- Jawne blokowanie zasobów
- Linda

Utrzymywanie spójności pamięci podręcznych

Źródła kłopotów

- lokalne pamięci podręczne
- buforowanie operacji pisania

Modele spójności - podobnie jak na BD

- Spójność ścista
- Spójność sekwencyjna
- Spójność słaba, na wyjściu, na wejściu

Implementacja

- Algorytm scentralizowanego zarządcy
- Dynamiczny algorytm zarządcy rozproszonego
- Jawne blokowanie zasobów
- Linda

Utrzymywanie spójności pamięci podręcznych

Źródła kłopotów

- lokalne pamięci podręczne
- buforowanie operacji pisania

Modele spójności - podobnie jak na BD

- Spójność ścisła
- Spójność sekwencyjna
- Spójność słaba, na wyjściu, na wejściu

Implementacja

- Algorytm scentralizowanego zarządcy
- Dynamiczny algorytm zarządcy rozproszonego
- Jawne blokowanie zasobów
- Linda

Plan Prezentacji

- 1 Część1 - Potokowanie, WMB (Paweł Bedyński)
- 2 Część 2 - Architektury 32-bitowe i 64-bitowe (Krzysztof Kąs)
- 3 Część 3 VLIW + Multiprocessing (Krzysztof Niemkiewicz)
- 4 Część 4 - Multiprocessing cd. (Krzysztof Niemkiewicz)

Przykre wspomnienia z PW



Filozofia

- Problem pięciu filozofów (raczej rzadko spotykany w SO)
- Wyłączność dostępu do zasobów
- Problem czytelników i pisarzy
- Zapewnienie „normalnej” semantyki programów na wielu procesorach

„Normalna” semantyka

```
x = 0;  
x = x + 1;  
if (x == 0) printf("Spokojnie, to tylko PW");
```

Przykre wspomnienia z PW



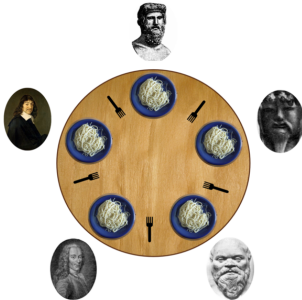
Filozofia

- Problem pięciu filozofów (raczej rzadko spotykany w SO)
- Wyłączność dostępu do zasobów
- Problem czytelników i pisarzy
- Zapewnienie „normalnej” semantyki programów na wielu procesorach

„Normalna” semantyka

```
x = 5;  
x = x + 1;  
if (x!=6)printf("Spokojnie, to tylko PW");
```

Przykre wspomnienia z PW



Filozofia

- Problem pięciu filozofów (raczej rzadko spotykany w SO)
- Wyłączność dostępu do zasobów
- Problem czytelników i pisarzy
- Zapewnienie „normalnej” semantyki programów na wielu procesorach

„Normalna” semantyka

$x = 5;$

$x = x + 1;$

if ($x \neq 6$) printf("Spokojnie, to tylko PW");

Rozwiązania w linuxie

Znane nam rozwiązania w linuxie

- `atomic_t`, `atomic_read(v)`, `atomic_set(v, i)`, `atomic_dec_and_test(v)` ...
- blokowanie przerwań
- `spinlock_t`
- `rw_lock_t`
- semafony(< 2.6.24), mutexy(> 2.6.16) – „my” jesteśmy akurat pominięty

Rozwiązania w linuxie

Znane nam rozwiązania w linuxie

- `atomic_t`, `atomic_read(v)`, `atomic_set(v, i)`, `atomic_dec_and_test(v)` ...
- blokowanie przerwań
- `spinlock_t`
- `rw_lock_t`
- semafony(< 2.6.24), mutexy(> 2.6.16) – „my” jestelmy akurat pomigdy

Rozwiązania w linuxie

Znane nam rozwiązania w linuxie

- `atomic_t`, `atomic_read(v)`, `atomic_set(v, i)`, `atomic_dec_and_test(v)` ...
- blokowanie przerwań
- `spinlock_t`
- `rw_lock_t`
- semafony(< 2.6.24), mutexy(> 2.6.16) – „my” jesteśmy akurat pominięci

Rozwiązania w linuxie

Znane nam rozwiązania w linuxie

- `atomic_t`, `atomic_read(v)`, `atomic_set(v, i)`, `atomic_dec_and_test(v)` ...
- blokowanie przerwań
- `spinlock_t`
- `rw_lock_t`
- semafony(< 2.6.24), mutexy(> 2.6.16) - „my” jesteśmy akurat pomiędzy

Rozwiązania w linuxie

Znane nam rozwiązania w linuxie

- `atomic_t`, `atomic_read(v)`, `atomic_set(v, i)`, `atomic_dec_and_test(v)` ...
- blokowanie przerwań
- `spinlock_t`
- `rw_lock_t`
- semafony(< 2.6.24), mutexy(> 2.6.16) - „my“ jesteśmy akurat pomiędzy

Rozwiązania w Windows

Porady dla twórców sterowników dla Windows

- Różne poziomy IRQL
- słowo kluczowe volatile - mało skuteczne, semantyka nie jest dokładnie wyspecyfikowana
- operacje atomowe, analogicznie jak w linuxie
- InterlockedXxx(szybkie) i ExInterlockedXxx
- Spinlocki
- Mutex, fast mutex, and executive resource

Rozwiązania w Windows

Porady dla twórców sterowników dla Windows

- Różne poziomy IRQL
 - słowo kluczowe `volatile` - mało skuteczne, semantyka nie jest dokładnie wyspecyfikowana
 - operacje atomowe, analogicznie jak w linuxie
 - `InterlockedXxx`(szybkie) i `ExInterlockedXxx`
 - Spinlocki
 - Mutex, fast mutex, and executive resource

Rozwiązania w Windows

Porady dla twórców sterowników dla Windows

- Różne poziomy IRQL
- słowo kluczowe `volatile` - mało skuteczne, semantyka nie jest dokładnie wyspecyfikowana
- operacje atomowe, analogicznie jak w linuxie
- `InterlockedXxx`(szybkie) i `ExInterlockedXxx`
- Spinlocki
- Mutex, fast mutex, and executive resource

Rozwiązania w Windows

Porady dla twórców sterowników dla Windows

- Różne poziomy IRQL
- słowo kluczowe `volatile` - mało skuteczne, semantyka nie jest dokładnie wyspecyfikowana
- operacje atomowe, analogicznie jak w linuxie
 - `InterlockedXxx`(szybkie) i `ExInterlockedXxx`
 - Spinlocki
 - Mutex, fast mutex, and executive resource

Rozwiązania w Windows

Porady dla twórców sterowników dla Windows

- Różne poziomy IRQL
- słowo kluczowe `volatile` - mało skuteczne, semantyka nie jest dokładnie wyspecyfikowana
- operacje atomowe, analogicznie jak w linuxie
- `InterlockedXxx`(szybkie) i `ExInterlockedXxx`
- Spinlocki
- `Mutex`, `fast mutex`, and `executive resource`

Rozwiązania w Windows

Porady dla twórców sterowników dla Windows

- Różne poziomy IRQL
- słowo kluczowe volatile - mało skuteczne, semantyka nie jest dokładnie wyspecyfikowana
- operacje atomowe, analogicznie jak w linuxie
- InterlockedXxx(szybkie) i ExInterlockedXxx
- Spinlocki
- Mutex, fast mutex, and executive resource

Scheduler w linuxie - stary i nowy

Stary(jądra 2.4.x)

- analogiczny jak dla jednego procesora
- jest jedna globalna kolejka procesów gotowych(wtedy był jeszcze Big Kernel Lock)
- wolny procesor bierze najlepszy proces z kolejki
- brak jakichkolwiek rozwiązań ustalających optymalne rozłożenie procesów

Nowy(jądra 2.6.x)

- osobna kolejka dla każdego procesora
- Load balancing co 200ms(albo 1ms)
- jak dobierać procesy między procesorami?
- pamiętajmy o potrzebie synchronizacji cache

Scheduler w linuxie - stary i nowy

Stary(jądra 2.4.x)

- analogiczny jak dla jednego procesora
- jest jedna globalna kolejka procesów gotowych(wtedy był jeszcze Big Kernel Lock)
- wolny procesor bierze najlepszy proces z kolejki
- brak jakichkolwiek rozwiązań ustalających optymalne rozłożenie procesów

Nowy(jądra 2.6.x)

- osobna kolejka dla każdego procesora
- Load balancing co 200ms(albo 1ms)
- jak dobierać procesy między procesorami?
- pamiętajmy o potrzebie synchronizacji cache

Scheduler w linuxie - stary i nowy

Stary(jądra 2.4.x)

- analogiczny jak dla jednego procesora
- jest jedna globalna kolejka procesów gotowych(wtedy był jeszcze Big Kernel Lock)
- wolny procesor bierze najlepszy proces z kolejki
- brak jakichkolwiek rozwiązań ustalających optymalne rozłożenie procesów

Nowy(jądra 2.6.x)

- osobna kolejka dla każdego procesora
- Load balancing co 200ms(albo 1ms)
- jak dobierać procesy między procesorami?
- pamiętajmy o potrzebie synchronizacji cache

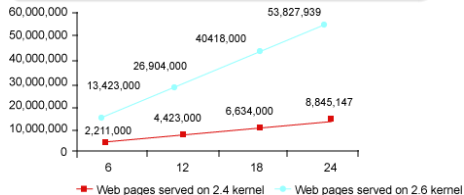
Scheduler w linuxie - stary i nowy

Stary(jądra 2.4.x)

- analogiczny jak dla jednego procesora
- jest jedna globalna kolejka procesów gotowych (wtedy był jeszcze Big Kernel Lock)
- wolny procesor bierze najlepszy proces z kolejki
- brak jakichkolwiek rozwiązań ustalających optymalne rozłożenie procesów

Nowy(jądra 2.6.x)

- osobna kolejka dla każdego procesora
- Load balancing co 200ms (albo 1ms)
- jak dobierać procesy między procesorami?
- pamiętać o potrzebie synchronizacji cache



Load balancing

Procedura wyrównywania obciążenia procesorów

- wyszukanie najbardziej zatłoczonej kolejki
- wybór kolejki expired lub active
- wybór procesu o najmniejszej wartości priorytetu który nie jest wykonywany i nie ma danych w cache'u
- przeniesienie procesu
- i tak dalej aż te dwie kolejki nie będą zbalansowane

Load balancing

Procedura wyrównywania obciążenia procesorów

- wyszukanie najbardziej zatłoczonej kolejki
- wybór kolejki expired lub active
- wybór procesu o najmniejszej wartości priorytetu który nie jest wykonywany i nie ma danych w cache'u
- przeniesienie procesu
- i tak dalej aż te dwie kolejki nie będą zbalansowane

Load balancing

Procedura wyrównywania obciążenia procesorów

- wyszukanie najbardziej zatłoczonej kolejki
- wybór kolejki expired lub active
- wybór procesu o najmniejszej wartości priorytetu który nie jest wykonywany i nie ma danych w cache'u
- przeniesienie procesu
- i tak dalej aż te dwie kolejki nie będą zbalansowane

Load balancing

Procedura wyrównywania obciążenia procesorów

- wyszukanie najbardziej zatłoczonej kolejki
- wybór kolejki expired lub active
- wybór procesu o najmniejszej wartości priorytetu który nie jest wykonywany i nie ma danych w cache'u
- przeniesienie procesu
- i tak dalej aż te dwie kolejki nie będą zbalansowane

Load balancing

Procedura wyrównywania obciążenia procesorów

- wyszukanie najbardziej zatłoczonej kolejki
- wybór kolejki expired lub active
- wybór procesu o najmniejszej wartości priorytetu który nie jest wykonywany i nie ma danych w cache'u
- przeniesienie procesu
- i tak dalej aż te dwie kolejki nie będą zbalansowane

Load balancing

Procedura wyrównywania obciążenia procesorów

- wyszukanie najbardziej zatłoczonej kolejki
- wybór kolejki expired lub active
- wybór procesu o najmniejszej wartości priorytetu który nie jest wykonywany i nie ma danych w cache'u
- przeniesienie procesu
- i tak dalej aż te dwie kolejki nie będą zbalansowane

Ogólne wymagania

Prawo Amdahl'a

$$W = \frac{1}{(1-P) + \frac{P}{S}}$$

- W - czas wykonania programu z wieloma procesorami w stosunku do czasu na pojedynczym procesorze
- P - część programu którą można wykonać równolegle
- S - liczba procesorów logicznych do dyspozycji
- Zakładamy, że nasze procesory logiczne są takiej wydajności jak ten jeden pierwotny

Inne ograniczenia

- Dostęp do szyny danych
- Dostęp do pamięci
- Dostęp do struktur jądra
- Dostęp do dysku, sieci itd.
- W praktyce dobrze jeśli dla dwóch rdzeni mamy 60% szybsze wykonanie

Ogólne wymagania

Prawo Amdahl'a

$$W = \frac{1}{(1-P) + \frac{P}{S}}$$

- W - czas wykonania programu z wieloma procesorami w stosunku do czasu na pojedynczym procesorze
- P - część programu którą można wykonać równolegle
- S - liczba procesorów logicznych do dyspozycji
- Zakładamy, że nasze procesory logiczne są takiej wydajności jak ten jeden pierwotny

Inne ograniczenia

- Dostęp do szyny danych
- Dostęp do pamięci
- Dostęp do struktur jądra
- Dostęp do dysku, sieci itd.
- W praktyce dobrze jeśli dla dwóch rdzeni mamy 60% szybsze wykonanie

Ogólne wymagania

Prawo Amdahl'a

$$W = \frac{1}{(1-P) + \frac{P}{S}}$$

- W - czas wykonania programu z wieloma procesorami w stosunku do czasu na pojedynczym procesorze
- P - część programu którą można wykonać równolegle
- S - liczba procesorów logicznych do dyspozycji
- Zakładamy, że nasze procesory logiczne są takiej wydajności jak ten jeden pierwotny

Inne ograniczenia

- Dostęp do szyny danych
- Dostęp do pamięci
- Dostęp do struktur jądra
- Dostęp do dysku, sieci itd.
- W praktyce dobrze jeśli dla dwóch rdzeni mamy 60% szybsze wykonanie

Ogólne wymagania

Prawo Amdahl'a

$$W = \frac{1}{(1-P) + \frac{P}{S}}$$

- W - czas wykonania programu z wieloma procesorami w stosunku do czasu na pojedynczym procesorze
- P - część programu którą można wykonać równolegle
- S - liczba procesorów logicznych do dyspozycji
- Zakładamy, że nasze procesory logiczne są takiej wydajności jak ten jeden pierwotny

Inne ograniczenia

- Dostęp do szyny danych
- Dostęp do pamięci
- Dostęp do struktur jądra
- Dostęp do dysku, sieci itd.
- W praktyce dobrze jeśli dla dwóch rdzeni mamy 60% szybsze wykonanie

Ogólne wymagania

Prawo Amdahl'a

$$W = \frac{1}{(1-P) + \frac{P}{S}}$$

- W - czas wykonania programu z wieloma procesorami w stosunku do czasu na pojedynczym procesorze
- P - część programu którą można wykonać równolegle
- S - liczba procesorów logicznych do dyspozycji
- Zakładamy, że nasze procesory logiczne są takiej wydajności jak ten jeden pierwotny

Inne ograniczenia

- Dostęp do szyny danych
- Dostęp do pamięci
- Dostęp do struktur jądra
- Dostęp do dysku, sieci itd.
- W praktyce dobrze jeśli dla dwóch rdzeni mamy 60% szybsze wykonanie

Grupy aplikacji przystosowanych do wielowątkowości...

Aplikacje biznesowe

- serwery WWW i serwery aplikacji, np Apache, Tomcat, JBoss...
- SAPy
- Bazy danych

Naukowe

- SETI, fizyka cząstek elementarnych
- alg. genetyczne (AI), Blast

Inne

- łamanie haseł, zwłaszcza brute-force
- ...

Grupy aplikacji przystosowanych do wielowątkowości...

Aplikacje biznesowe

- serwery WWW i serwery aplikacji, np Apache, Tomcat, JBoss...
- SAPy
- Bazy danych

Naukowe

- SETI, fizyka cząstek elementarnych
- alg. genetyczne (AI), Blast

Inne

- łamanie haseł, zwłaszcza brute-force
- ...

Grupy aplikacji przystosowanych do wielowątkowości...

Aplikacje biznesowe

- serwery WWW i serwery aplikacji, np Apache, Tomcat, JBoss...
- SAPy
- Bazy danych

Naukowe

- SETI, fizyka cząstek elementarnych
- alg. genetyczne (AI), Blast

Inne

- łamanie haseł, zwłaszcza brute-force
- ...

Grupy aplikacji przystosowanych do wielowątkowości...

Aplikacje biznesowe

- serwery WWW i serwery aplikacji, np Apache, Tomcat, JBoss...
- SAPy
- Bazy danych

Naukowe

- SETI, fizyka cząstek elementarnych
- alg. genetyczne (AI), Blast

Inne

- Łamanie haseł, zwłaszcza brute-force
- ...

... i inne typy aplikacji które też chcą skorzystać z nowych możliwości



Jeszcze inne :)...

- Bioshock, UT3 , Call Of Duty 4, Company of Heroes ,Crysis

... i inne typy aplikacji które też chcą skorzystać z nowych możliwości



Jeszcze inne :)...

- Bioshock, UT3 , Call Of Duty 4, Company of Heroes ,Crysis

SISD, SIMD, MISD, MIMD

Podział architektur ze względu na liczbę potoków

- SISD(Single Instruction, Single Data) - komputery z jednym procesorem logicznym
- MISD(Multiple Instruction, Single Data) - rzadko spotykane
- SIMD(Single Instruction, Multiple Data) - jednostki wektorowe, GPU
- MIMD(Multiple Instruction, Multiple Data) - komputery z wieloma procesorami, systemy rozproszone

Obecnie ten podział traci na znaczeniu

- Prawie każdy komputer jest MIMD
- Prawie każdy komputer zawiera jakieś elementy typu SIMD

SISD, SIMD, MISD, MIMD

Podział architektur ze względu na liczbę potoków

- SISD(Single Instruction, Single Data) - komputery z jednym procesorem logicznym
- MISD(Multiple Instruction, Single Data) - rzadko spotykane
- SIMD(Single Instruction, Multiple Data) - jednostki wektorowe, GPU
- MIMD(Multiple Instruction, Multiple Data) - komputery z wieloma procesorami, systemy rozproszone

Obecnie ten podział traci na znaczeniu

- Prawie każdy komputer jest MIMD
- Prawie każdy komputer zawiera jakieś elementy typu SIMD

SISD, SIMD, MISD, MIMD

Podział architektur ze względu na liczbę potoków

- SISD(Single Instruction, Single Data) - komputery z jednym procesorem logicznym
- MISD(Multiple Instruction, Single Data) - rzadko spotykane
- SIMD(Single Instruction, Multiple Data) - jednostki wektorowe, GPU
- MIMD(Multiple Instruction, Multiple Data) - komputery z wieloma procesorami, systemy rozproszone

Obecnie ten podział traci na znaczeniu

- Prawie każdy komputer jest MIMD
- Prawie każdy komputer zawiera jakieś elementy typu SIMD

SISD, SIMD, MISD, MIMD

Podział architektur ze względu na liczbę potoków

- SISD(Single Instruction, Single Data) - komputery z jednym procesorem logicznym
- MISD(Multiple Instruction, Single Data) - rzadko spotykane
- SIMD(Single Instruction, Multiple Data) - jednostki wektorowe, GPU
- MIMD(Multiple Instruction, Multiple Data) - komputery z wieloma procesorami, systemy rozproszone

Obecnie ten podział traci na znaczeniu

- Prawie każdy komputer jest MIMD
- Prawie każdy komputer zawiera jakieś elementy typu SIMD

SISD, SIMD, MISD, MIMD

Podział architektur ze względu na liczbę potoków

- SISD(Single Instruction, Single Data) - komputery z jednym procesorem logicznym
- MISD(Multiple Instruction, Single Data) - rzadko spotykane
- SIMD(Single Instruction, Multiple Data) - jednostki wektorowe, GPU
- MIMD(Multiple Instruction, Multiple Data) - komputery z wieloma procesorami, systemy rozproszone

Obecnie ten podział traci na znaczeniu

- Prawie każdy komputer jest MIMD
- Prawie każdy komputer zawiera jakieś elementy typu SIMD

SISD, SIMD, MISD, MIMD

Podział architektur ze względu na liczbę potoków

- SISD(Single Instruction, Single Data) - komputery z jednym procesorem logicznym
- MISD(Multiple Instruction, Single Data) - rzadko spotykane
- SIMD(Single Instruction, Multiple Data) - jednostki wektorowe, GPU
- MIMD(Multiple Instruction, Multiple Data) - komputery z wieloma procesorami, systemy rozproszone

Obecnie ten podział traci na znaczeniu

- Prawie każdy komputer jest MIMD
- Prawie każdy komputer zawiera jakieś elementy typu SIMD

SISD, SIMD, MISD, MIMD

Podział architektur ze względu na liczbę potoków

- SISD(Single Instruction, Single Data) - komputery z jednym procesorem logicznym
- MISD(Multiple Instruction, Single Data) - rzadko spotykane
- SIMD(Single Instruction, Multiple Data) - jednostki wektorowe, GPU
- MIMD(Multiple Instruction, Multiple Data) - komputery z wieloma procesorami, systemy rozproszone

Obecnie ten podział traci na znaczeniu

- Prawie każdy komputer jest MIMD
- Prawie każdy komputer zawiera jakieś elementy typu SIMD

SISD, SIMD, MISD, MIMD

Podział architektur ze względu na liczbę potoków

- SISD(Single Instruction, Single Data) - komputery z jednym procesorem logicznym
- MISD(Multiple Instruction, Single Data) - rzadko spotykane
- SIMD(Single Instruction, Multiple Data) - jednostki wektorowe, GPU
- MIMD(Multiple Instruction, Multiple Data) - komputery z wieloma procesorami, systemy rozproszone

Obecnie ten podział traci na znaczeniu

- Prawie każdy komputer jest MIMD
- Prawie każdy komputer zawiera jakieś elementy typu SIMD

Systemy rozproszone

Systemy rozproszone

- Sieciowy system operacyjny - współdzielone pliki
- Rozproszony system operacyjny (wiele komputerów) - komunikaty
- Rozproszony system operacyjny (wieloprocessorowy) - współdzielona pamięć
- Ciąg dalszy nastąpi...

Systemy rozproszone

Systemy rozproszone

- Sieciowy system operacyjny - współdzielone pliki
- Rozproszony system operacyjny (wiele komputerów) - komunikaty
- Rozproszony system operacyjny (wieloprocessorowy) - współdzielona pamięć
- Ciąg dalszy nastąpi...

Systemy rozproszone

Systemy rozproszone

- Sieciowy system operacyjny - współdzielone pliki
- Rozproszony system operacyjny (wiele komputerów) - komunikaty
- Rozproszony system operacyjny (wieloprocessorowy) - współdzielona pamięć
- Ciąg dalszy nastąpi...

Systemy rozproszone

Systemy rozproszone

- Sieciowy system operacyjny - współdzielone pliki
- Rozproszony system operacyjny (wiele komputerów) - komunikaty
- Rozproszony system operacyjny (wieloprocessorowy) - współdzielona pamięć
- Ciąg dalszy nastąpi...

Systemy rozproszone

Systemy rozproszone

- Sieciowy system operacyjny - współdzielone pliki
- Rozproszony system operacyjny (wiele komputerów) - komunikaty
- Rozproszony system operacyjny (wieloprocessorowy) - współdzielona pamięć
- Ciąg dalszy nastąpi...