

**Dr inż. hab. Siergiej Fialko, IMK-PK,**

**[https://torus.uck.pk.edu.pl/~fialko.sergiy/  
sfialko@pk.edu.pl](https://torus.uck.pk.edu.pl/~fialko.sergiy/sfialko@pk.edu.pl)**

Education

Research

Examin.

????

email: [sfialko@pk.edu.pl](mailto:sfialko@pk.edu.pl)

**Ogłoszenia:**

**Konsultacje**

Konsultacje, prowadzone jawnie - od oznaczonej godziny czekam 15 min. Jeśli nikt nie przyszedł - uważam, że konsultacja nie odbyła się.

| Dzien    | Godziny       | Tryb     | Status |
|----------|---------------|----------|--------|
| Wtorek   | 16:15 - 17:15 | MS TEAMS |        |
| Czwartek | 11:00 - 12:00 | p. 144/7 |        |

-----

**Materiały Dydaktyczne:**

**Linki aktualne:**

- [Modelowanie Zagadnień Technicznych](#) (MZT1.html)
- [JiPP cz. I. Studia stacjonarne i niestacjonarne](#) (CC.html)
- [JiPP cz. II. Studia stacjonarne i niestacjonarne.](#) (CPP\_z\_new.html)

Contact [Sergiy Fialko](#)

dr hab inż Sergiy Fialko  
IMK  
p. 144/7

## Wykłady:

Kurs MZT: <http://torus.uck.pk.edu.pl/~pk21>

- [Kurs MZT](#)(MZT.pdf)
- [w1](#) (w1.pdf)
- [w2](#) (w2.pdf)
- [w3](#) (w3.pdf)
- [w4](#) (w4.pdf)
- [w5](#) (w5.pdf)
- [w6](#) (w6.pdf)
- [w7](#) (w7.pdf)
- [w8](#) (w8.pdf)
- [w9](#) (w9.pdf)
- [w10](#) (w10.pdf)
- [w10\\_1](#) (w10\_1.pdf)

## Przykłady do wykładów:

- [MemTest\\_1](#) (MemTest\_1.zip)
- [MemTest\\_2](#) (MemTest\_2.zip)
- [Zrownoleglenie petli for za pomoca OpenMP](#) (parallel\_1.zip)
- [MultThreads](#) (MultThreads.pdf)
- [DoesNotDoIt](#) (DoesNotDoIt.zip)
- [SendMess](#) (SendMess.zip)
- [MatrixMultipl](#) (Mulm\_Stud.7z)

Laboratoria: Studia stacjonarne

- [Lab\\_1](#) (Lab\_1\_2.7z)
- [Lab\\_2](#) (Lab\_2.zip)
- [Lab\\_3 32 bit application](#) (Lab\_3\_1\_i\_ia32.7z)
- [Lab\\_3 64 bit application](#) (Lab\_3\_1\_i\_x64.7z)
- [Lab\\_3 64 bit application AVX](#) (Lab\_3\_1\_x64\_AVX.7z)
- [Lab\\_3 64 bit application AVX FMA3](#) (Lab\_3\_1\_x64\_AVX\_FMA.7z)
- [CPU\\_Z\\_x64](#) (cpuz64\_153.zip)
- [CPU\\_Z\\_ia32](#) (cpuz\_153.zip)
- [Obliczenie dot = X\\*Y w trybie wielowatkowym](#) (Lab\_dot\_prod.zip)
- [Obliczenie całki w trybie wielowatkowym](#) (Lab\_integr.7z)
- [Obliczenie C = A\\*B w trybie wielowatkowym](#) (L6.doc)
- [Macierze rzadkie: format skompresowany](#) (SparseMatr.7z)
- [Macierze rzadkie: format skompresowany](#) (SparseMatr\_mul.zip)

**Warunki zaliczenia:**

- 1. Zaliczenie wszystkich prac laboratoryjnych na podstawie obrony każdej pracy.**
- 2. Napisanie kolokwium zaliczeniowego, który obejmuje ostatnie dwa tematy.**

## Osobliwości przedmiotu

- W podanym kursie główna uwaga będzie przydzielona osobliwościom modelowania komputerowego zagadnień technicznych, czyli aspektom informatycznym, który powstają przy tworzeniu systemów obliczeniowych w zagadnieniach technicznych.
- Celą kursu jest zapoznanie z technikami i metodami opracowania algorytmów, często spotykanych przy modelowaniu zagadnień technicznych, przy czym główna uwaga przydzielona podstawom tworzenia takich metod na wielordzeniowych komputerach z pamięcią wspólną.

## **Kurs obejmuje:**

- **Osiągnięcie wysokiej wydajności typowych algorytmów obliczeniowych zagadnień technicznych przy oprogramowaniu sekwencyjnym, a również wielowątkowym dla komputerów wielordzeniowych (SMP) .**
- **Typowe techniki oprogramowania wielowątkowego dla algorytmów obliczeniowych, a również zrównoleglenia SIMD (wektoryzowanie obliczeń) oraz wspieranie potokowości (pipelines) współczesnych procesorów.**
- **Macierzy rzadkie. Metody bezpośrednie i iteracyjne rozwiązywania układów równań liniowych algebraicznych z macierzami rzadkimi symetrycznymi.**

## **Osobliwości kursu:**

- 1. Kurs zawiera elementy matematyki – algebry liniowej, algebry macierzy, teorii grafów, metod numerycznych, a również wymaga wiedzy z podstaw oprogramowania w języku C, C++, systemów operacyjnych i oprogramowania wielowątkowego.**
- 2. Brak podręczników. Źródła literaturowe są rozrzucane po wielu książkach (podręcznikach z matematyki, metod numerycznych, podręcznikach z informatyki), artykułach w czasopismach naukowych, stronach www .**

**Miara wydajności:**

**FLOPS = (ilość operacji zmiennoprzecinkowych)/(czas wykonania)**

**Jest naturalną miarą dla algorytmów, które powstają z zagadnień naukowo-technicznych (dominantowymi są operacji arytmetyczne zmiennoprzecinkowe), przecież nie jest przydatna do oceny wydajności innych rodzajów algorytmów (kompilatory, bazy danych, itp.)**

**Operacje zmiennoprzecinkowe mogą być szybkimi i wolnymi. Stosunek trwałości różnych operacji zmiennoprzecinkowych:**

| <b>Operacje zmiennoprzecinkowe</b>                    | <b>Stosunkowa trwałość</b> |
|-------------------------------------------------------|----------------------------|
| <b>dodawanie, odejmowanie, porównanie, mnożenie</b>   | <b>1</b>                   |
| <b>dzielenie, obliczenie pierwiastka kwadratowego</b> | <b>4</b>                   |
| <b>eksponenta, sinus, ...</b>                         | <b>8</b>                   |

**Pamięć główna systemu równoległego może być wspólną (shared memory) lub rozproszoną (distributed memory).**

**Charakterystyki pamięci:**

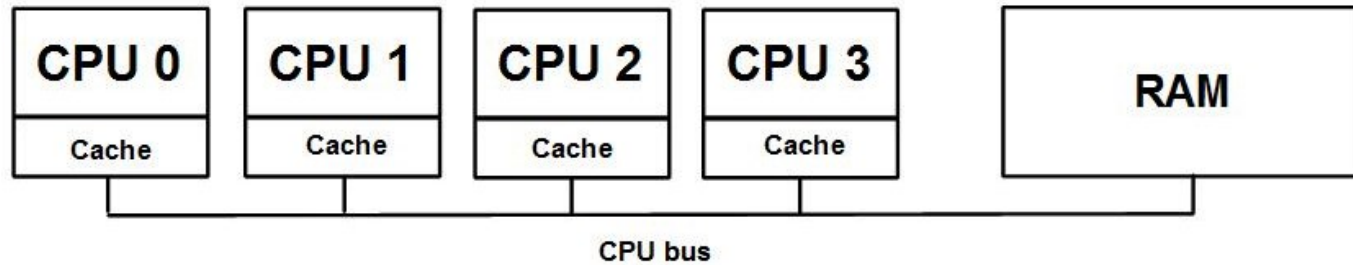
- **Objętość (MB), Latentność (Latency – ang.) (ns), Przepustowość (b/s).**
- **Wielopoziomowa pamięć podręczna (cache memory) – szybka pamięć dla przechowywania tymczasowych danych.**
- **bufory rozkazów i danych.**

**Wspólna pamięć (SMP – symmetric multiprocessing) – procesory mają równoprawny dostęp (każdy z procesorów nie ma żadnej przewagi nad innymi) i równoznaczny (czas dostępu do dowolnej części pamięci jest taki samy).**

**Zalety: związek między procesami i synchronizacja są proste, oprogramowanie mniej skomplikowane, niż w wypadku architektury z pamięcią rozproszoną.**

**Wady: mają podstawowe ograniczenie na ilość procesorów – krytycznym jest przepustowość magistrali.**

**Wspólna pamięć. Wieloprocessorowe systemy z jednorodnym dostępem do pamięci - UMA (Uniform Memory Access)**



**System symetryczny wieloprocessorowy SMP (Symmetrical Multiprocessing)  
UMA**

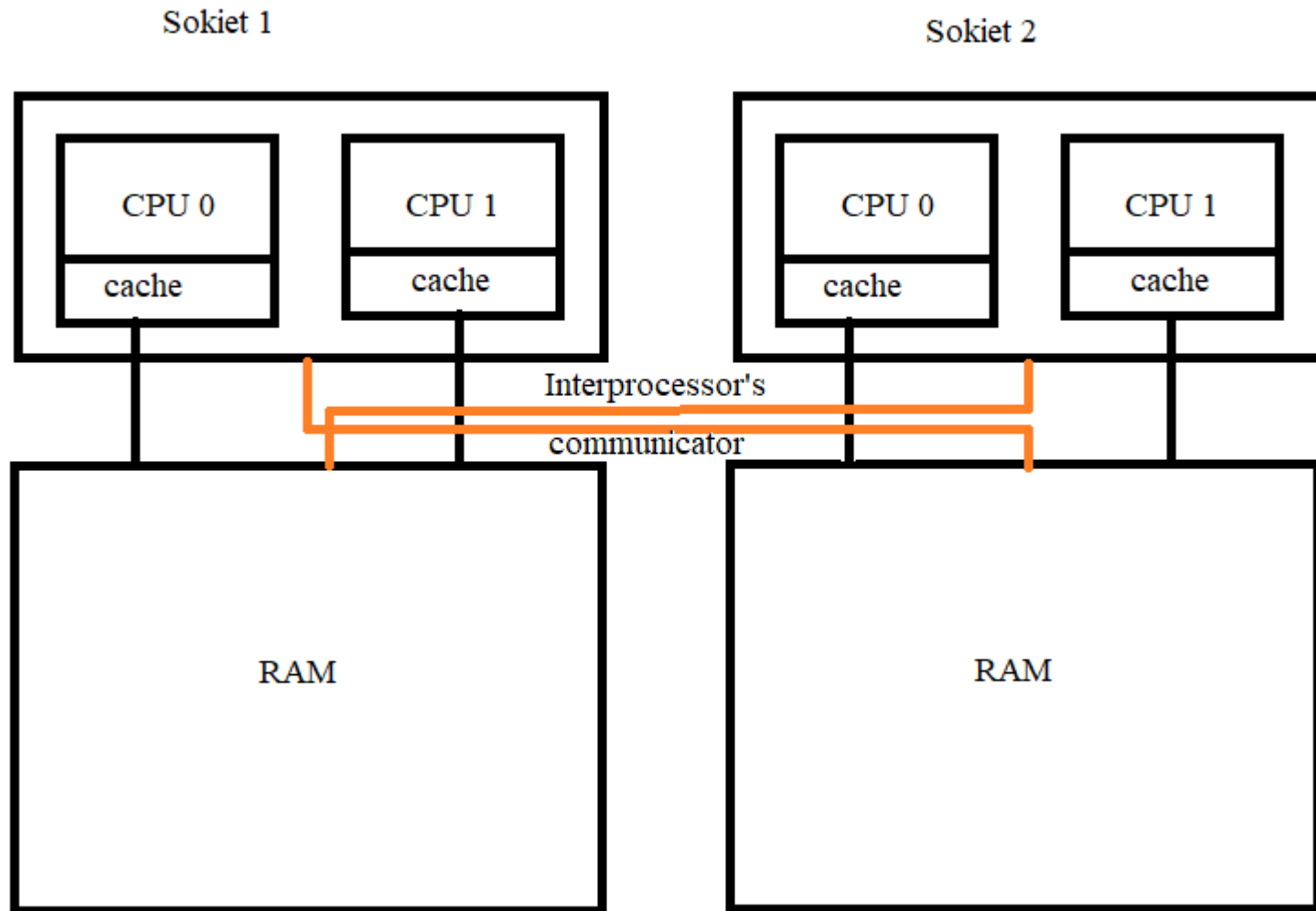
**Wieloprocessorowe systemy z niejednorodnym dostępem do pamięci - NUMA  
(Non - uniform Memory Access)**

Wielu procesorów są łączone jedną magistralą z tą samą pamięcią główną.

Technika zrównoleglenia dla SMP systemów zwykle polega na wielowątkowości.

Stan pamięci podręcznej każdego procesora ma spełnić warunku spójności danych.

**Wspólna pamięć. Wieloprocessorowe systemy z jednorodnym dostępem do pamięci - NUMA (Non-Uniform Memory Access)**



**Wielu procesorów, każdy ma własną fizyczną pamięć, dostęp do której jest szybki, i wspólną wirtualną pamięć. Jeśli procesor adresuje się do pamięci innego procesora, generuje się page fault, i jak obsługa sytuacji wyjątkowej, załadowuje się strona pamięci odległej. Dostęp do pamięci odległej jest wolny.**

## Cache-konflikty

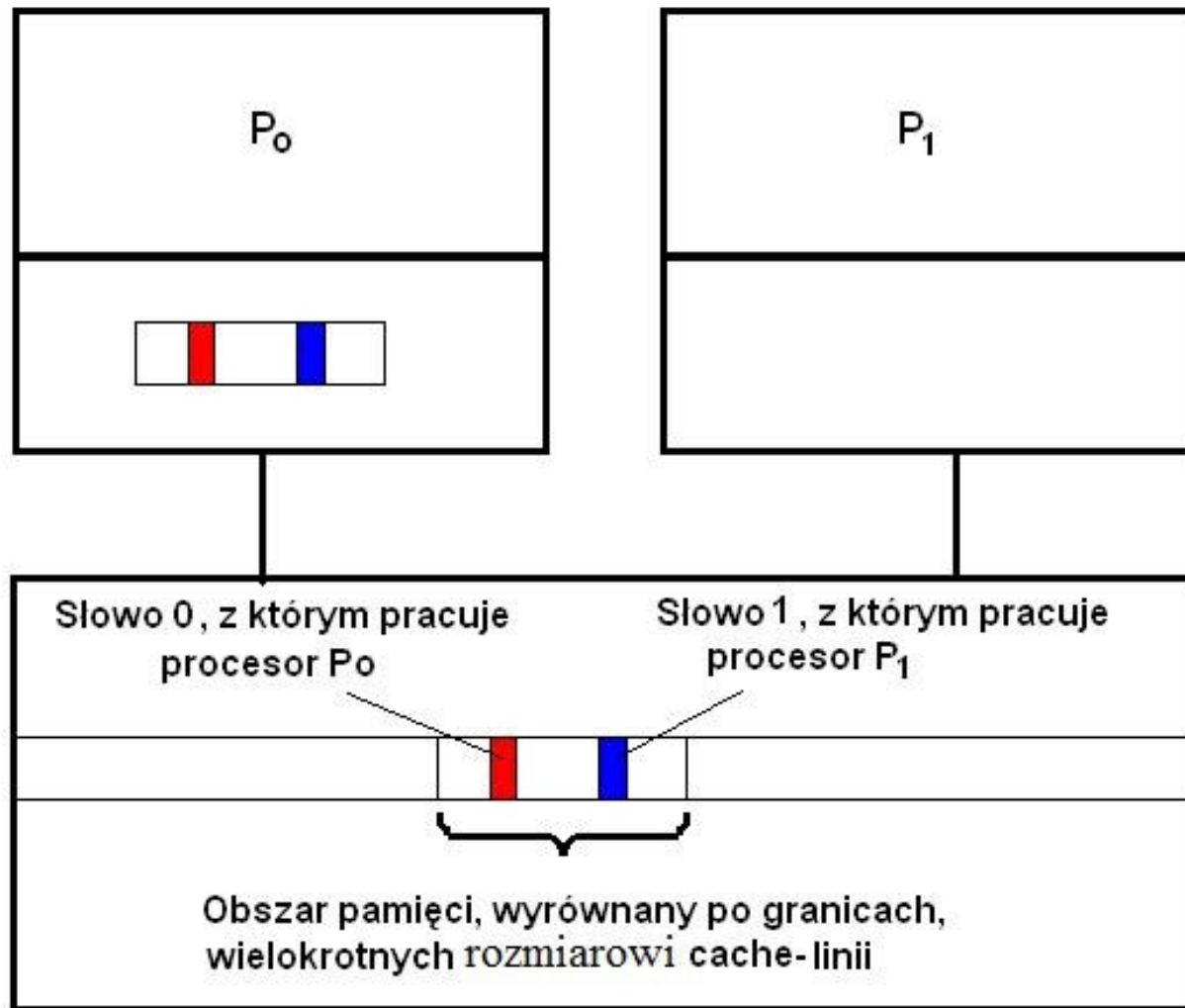
- Przy czytaniu danych (słowo, bajt) procesor najpierw sprawdza stan swojej pamięci podręcznej. Jeśli te dane są w pamięci podręcznej (read hit – trafienie), oni będą umieszczone w rejestry procesora. Jeśli tych danych (słowa, bajta) nie ma w pamięci podręcznej (read miss – chybienie), będzie przeładowana cała cache linia pamięci podręcznej.
- **Przy pobraniu tych danych (bajta, słowa) z pamięci głównej będą pobrane nie tylko to słowo, bajt, a i bajty sąsiednie tak, żeby wypełnić całą cache linie pamięci podręcznej.**
- W zależności od typu procesora pamięć podręczna składa się z linii o kilku słów (32 B, 64 B, 128 B). Dane z pamięci głównej będą umieszczone w pamięci podręcznej przed umieszczeniem w rejestrach procesora. Mówimy o odwzorowaniu bloków pamięci głównej w cache-linie pamięci podręcznej.
- Dane, które będą umieszczone w cache-linie, muszą być wyrównane po granicach, wielokrotnych 32 B, 64 B, 128 B odpowiednio.

- **Ideologia cache-linii polega na tym, że zadania najczęściej pracują ze zbiorem sąsiednich słów (bajtów), więc przy odczycie tych słów (bajtów) często wynika stan read hit, i procesor pobiera te dane w rejestry z cache-linii, a nie przez wolny kanał magistral – pamięć główna.**
- **Przecież cache-linii istotnie komplikują odnowienie pamięci w środowisku wieloprocesorowym z pamięcią wspólną**

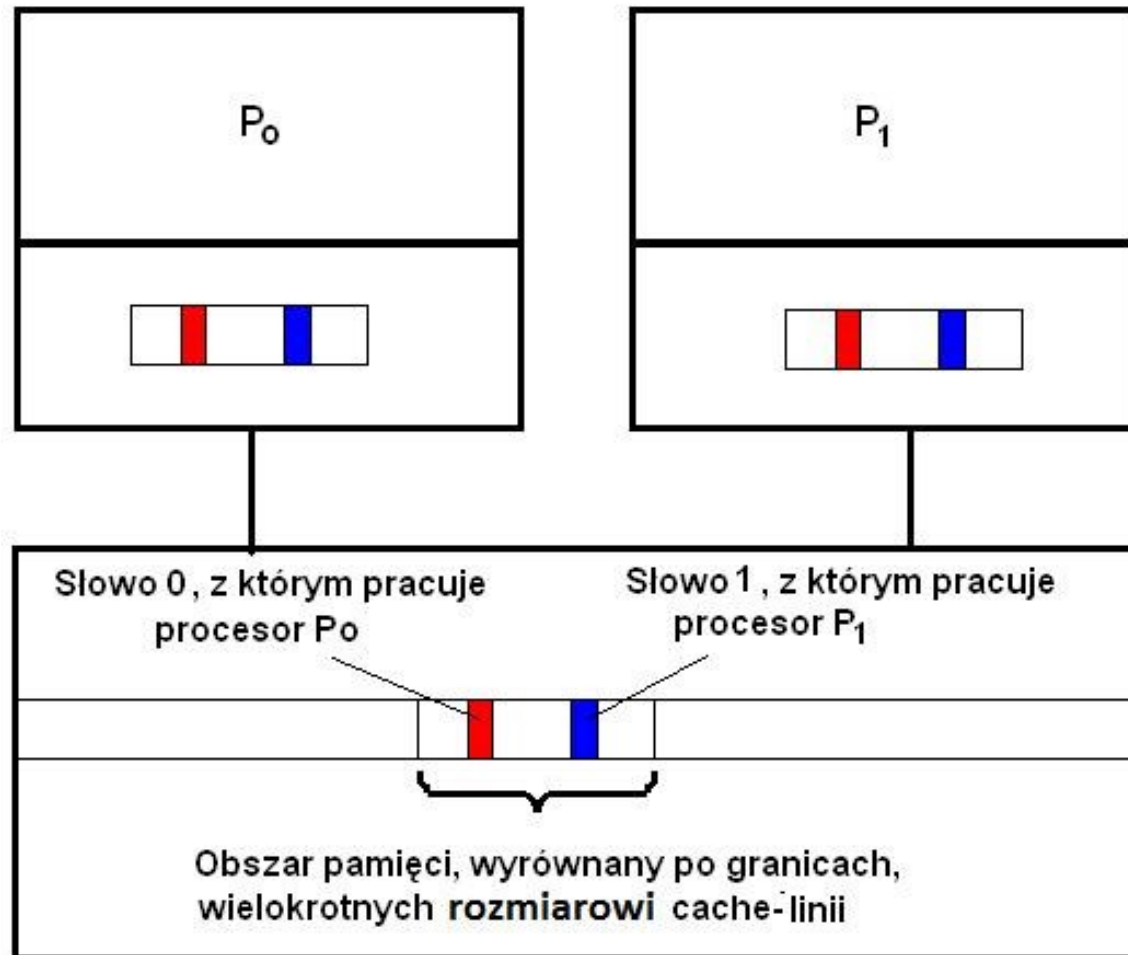
### **Przykład**

- 1. Procesor P0 odczytuje słowo 0, a również i słowa sąsiedni w swój cache**
- 2. Procesor P1 odczytuje słowo 1, a również i słowa sąsiedni w swój cache, przy czym słowo 0 procesora P0 w skutek wyrównania po granicach, wielokrotnych rozmiary cache-linii, znajduje się w cache-linii procesora P1 też.**
- 3. Procesor P0 modyfikuje słowo 0 i umieszcza ten wynik w swój cache. Wartość tego słowa w pamięci głównej na razie pozostaje bez zmian.**
- 4. Procesor P1 nic nie wie o tym, że zawartość słowa 1 podległa zmianie. Teraz wartość słowa 0 zależy od tego, cache którego procesora będzie wyładowany w pamięć główną później.**

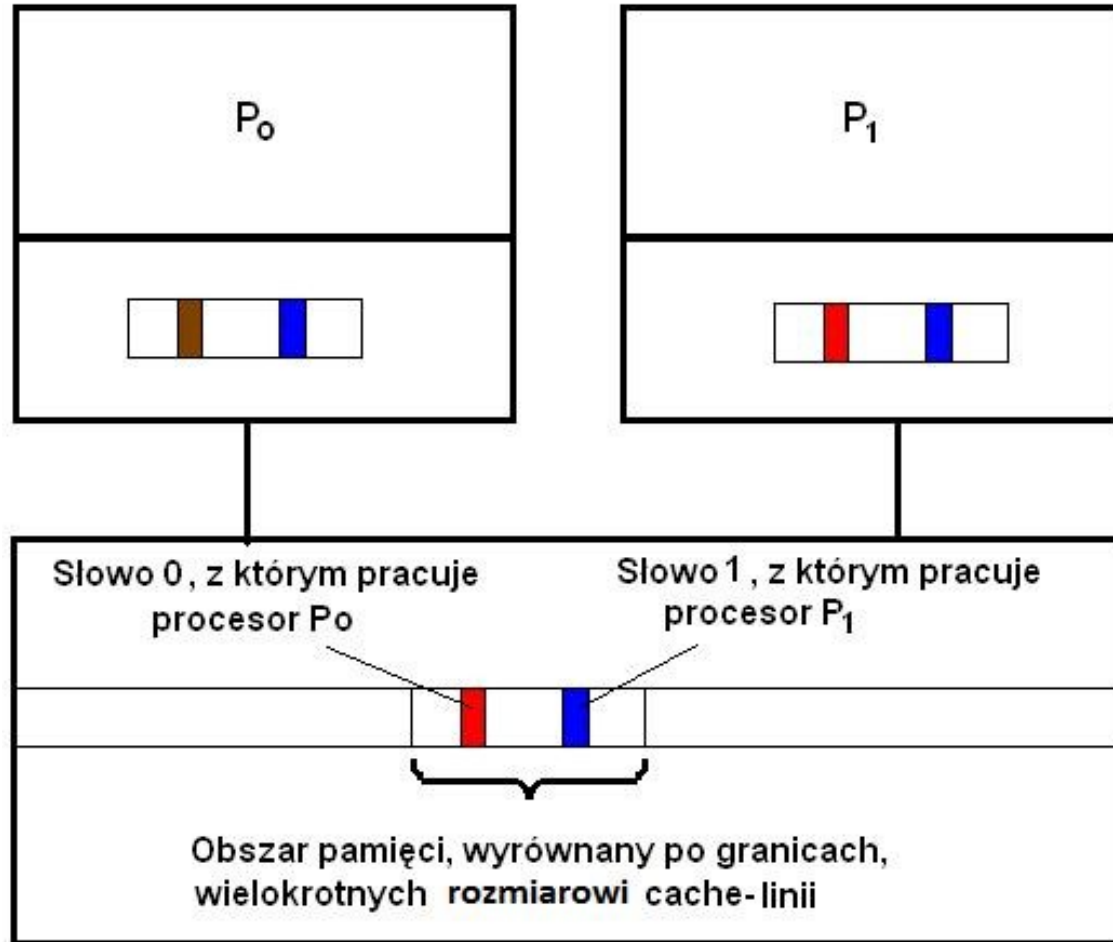
Procesor P<sub>0</sub> odczytuje słowo 0, a również i słowa sąsiedni w swój cache



**Procesor P1 odczytuje słowo 1, a również i słowa sąsiedni w swój cache, przy czym słowo 0 procesora P0**

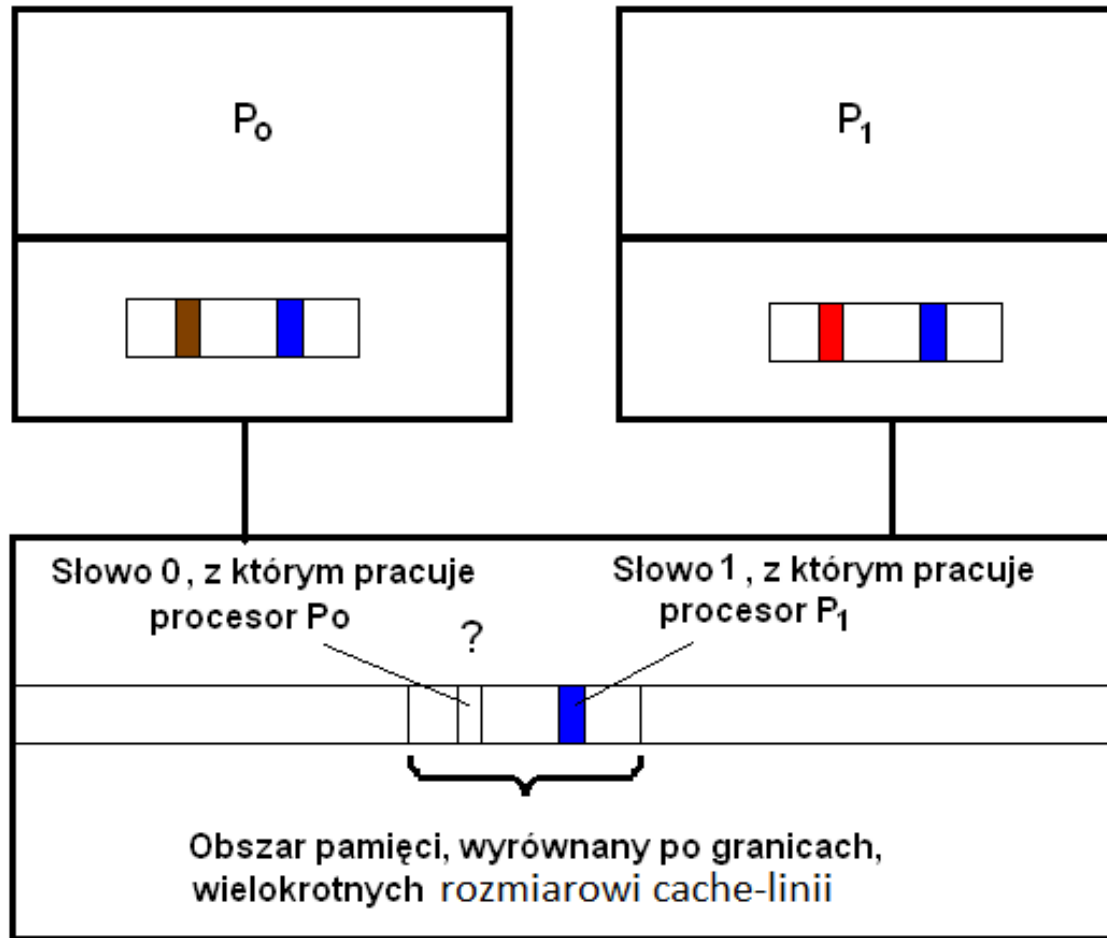


**Procesor P<sub>0</sub> modyfikuje słowo 0 i umieszcza ten wynik w swój cache. Wartość tego słowa w pamięci głównej na razie pozostaje bez zmian.**



W1

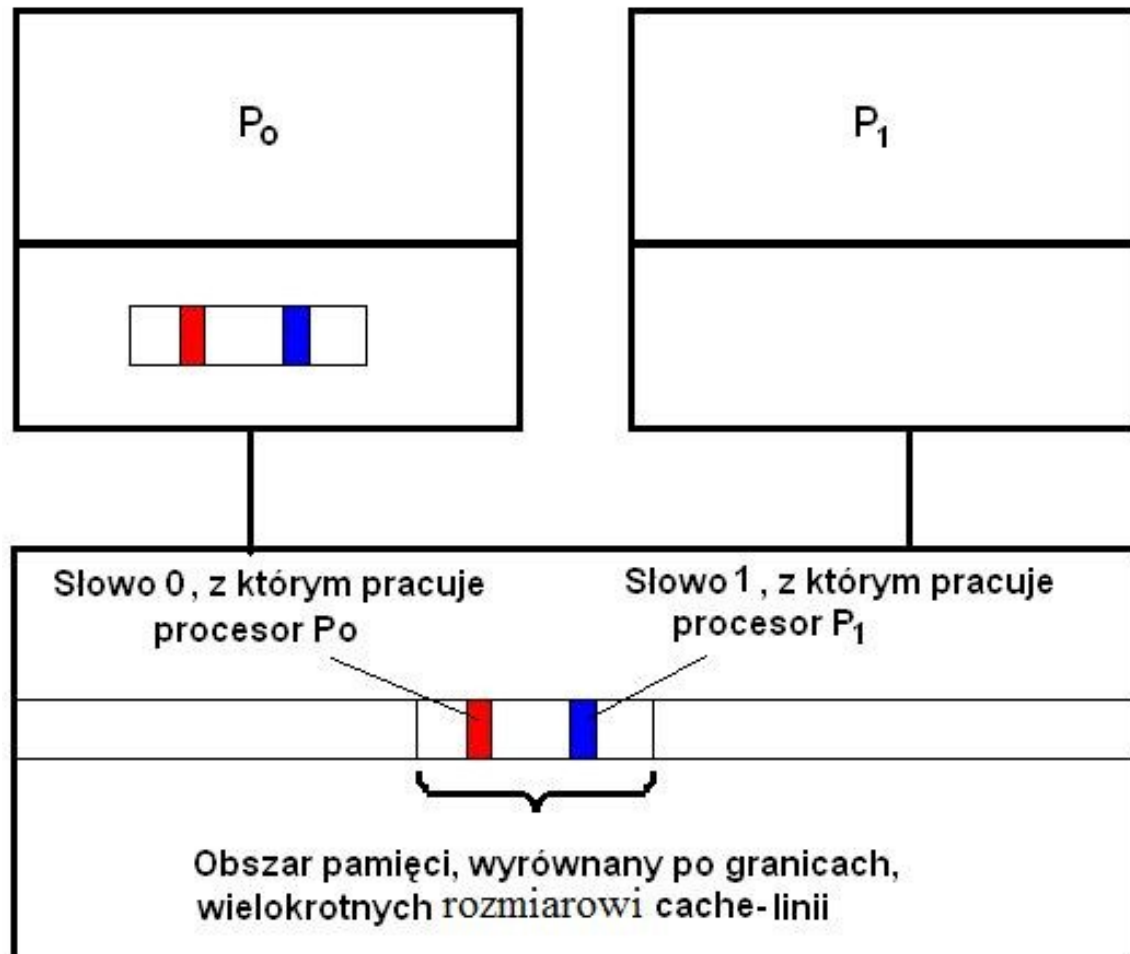
Procesor P<sub>1</sub> nic nie wie o tym, że zawartość słowa 1 podległa zmianie. Teraz wartość słowa 0 zależy od tego, cache którego procesora będzie wyładowany w pamięć główną później.



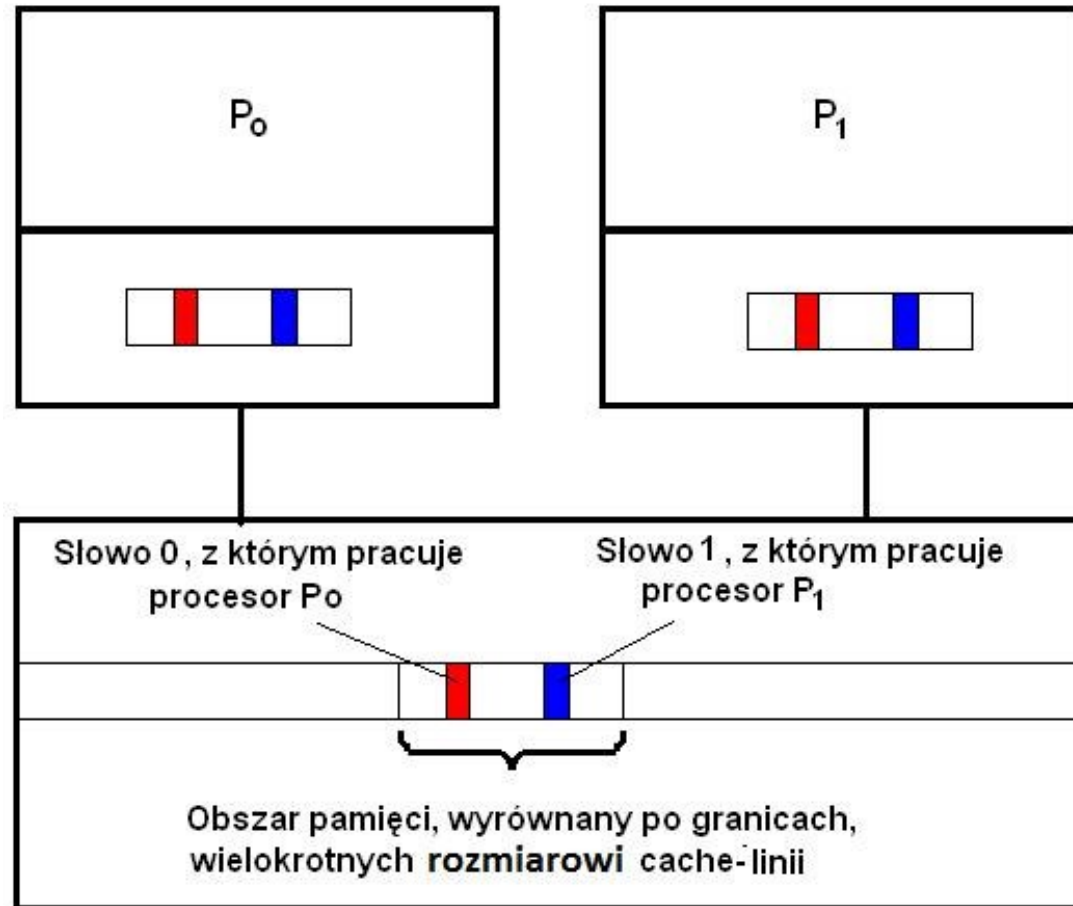
**Taki scenariusz był by katastrofą. Producenci sprzętu komputerowego bardzo dobrze wiedzą o tym, i dla tego rzeczywisty komputer nigdy tak nie działa.**

- 1. Procesor P0 odczytuje słowo 0, a również i słowa sąsiedni w swój cache**
- 2. Procesor P1 odczytuje słowo 1, a również i słowa sąsiedni w swój cache, przy czym słowo 0 procesora P0 w skutek wyrównania po granicach, wielokrotnych rozmiary cache-linii, znajduje się w cache-linii procesora P1 też.**
- 3. Procesor P0 modyfikuje słowo 0 i umieszcza ten wynik w swój cache. Wartość tego słowa w pamięci głównej na razie pozostaje bez zmian, ale procesor P1 w skutek snooping'u zgłasza zawartość swojego cache unieważnioną.**
- 4. Procesor P0 wyładowuje zawartość swojego cache w pamięć główną, a procesor P1 ponownie załadowuje swoją cache-linie. Teraz dane w pamięci głównej i w cache-liniach procesorów P0, P1 są spójne (koherentne). Przecież płacą za tą spójność – nadliczbowe przeładowania cache-linij procesorów P0, P1.**

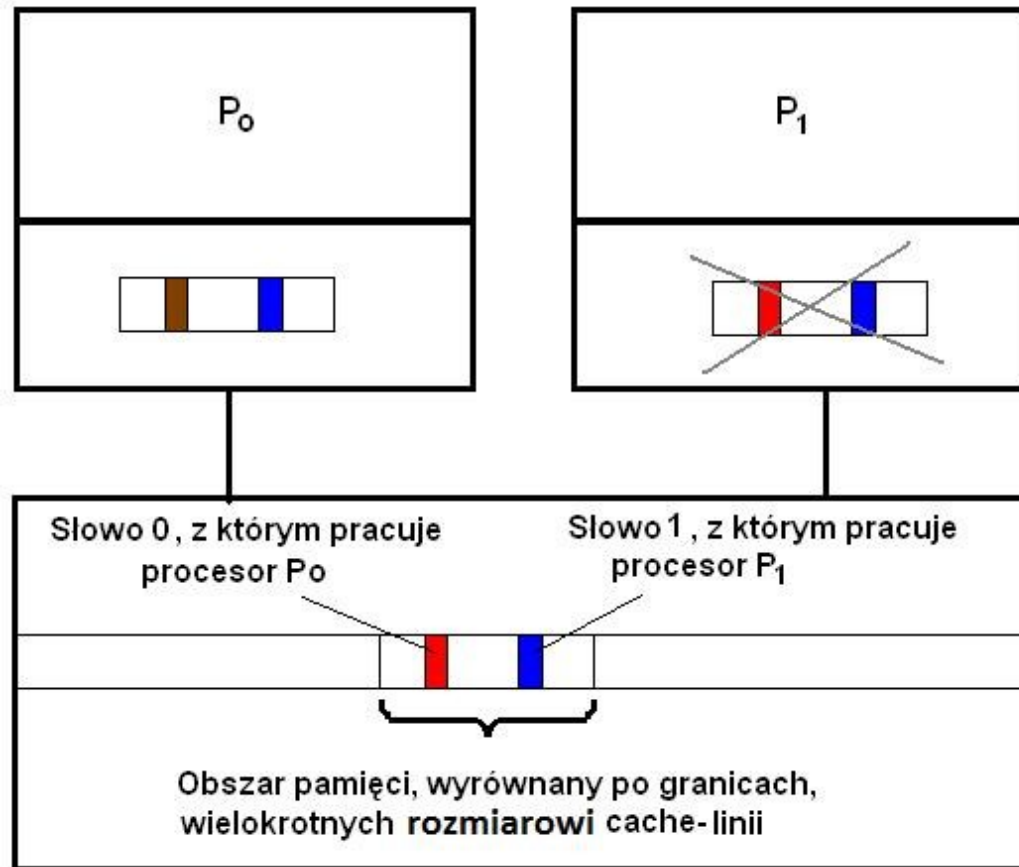
Procesor P<sub>0</sub> odczytuje słowo 0, a również i słowa sąsiedni w swój cache



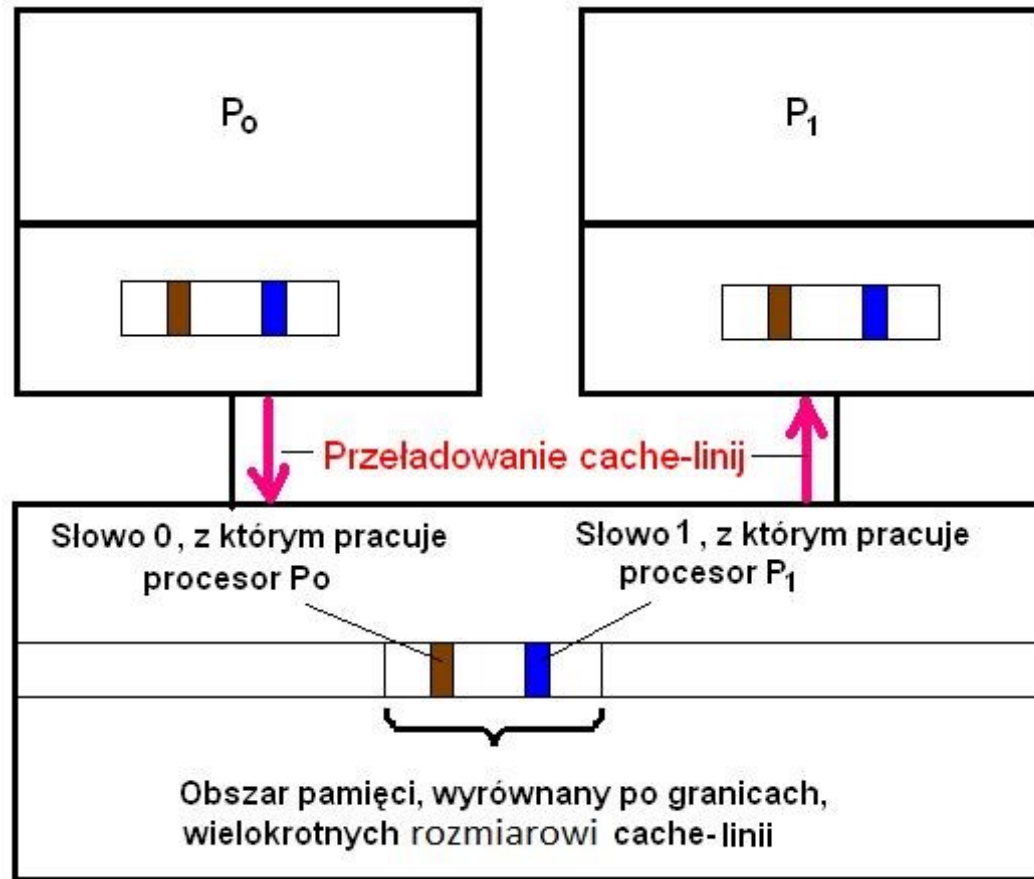
**Procesor P1 odczytuje słowo 1, a również i słowa sąsiedni w swój cache, przy czym słowo 0 procesora P0**



Procesor P<sub>0</sub> modyfikuje słowo 0 i umieszcza ten wynik w swój cache. Wartość tego słowa w pamięci głównej na razie pozostaje bez zmian, **ale procesor P<sub>1</sub> w skutek snooping'u zgłasza zawartość swojego cache unieważnioną.**



**Procesor P0 wyładowuje zawartość swojego cache w pamięć główną, a procesor P1 ponownie załadowuje swoją cache-linie. Teraz dane w pamięci głównej i w cache-liniach procesorów P0, P1 są spójne (koherentne).**



**W1** To się okazuje, że obecność pamięci podręcznej, przeznaczonej dla przyspieszenia obliczeń w skutek zmniejszenia ilości odczytów-zapisów bezpośrednich z/w pamięć(i) głównej, może stać przyczyną istotnego zmniejszenia wydajności obliczeń, jeśli dane nie będą wyrównane z uwzględnieniem podanej powyżej specyfiku pracy z cache-liniami.

To leży na programiście.

Środki, które zapobiegają wyniknięciu cache-konfliktów:

- **wyrównanie rozmieszczenia danych w pamięci głównej tak, żeby różne procesory mieli dostęp do różnych adresów pamięci, rozdzielonej granicą przynajmniej jednej cache-linii.**
- Rozdzielenie danych tylko dla odczytu od danych dla odczytu i zapisu. **Dane dla odczytu nie wymagają wyrównania na granicy cache-linii.**
- Użycie pamięci lokalnej dla każdego wątku. Wtedy wyrównanie na granicy cache linii odbywa się automatycznie.

**UWAGA!** Cache-konflikty powstają dla architektury ze wspólną pamięcią. **Dla architektury z pamięcią rozproszoną każdy procesor odczytuje-zapisuje dane, pobierane z pamięci lokalnej – cache konflikty nie występują.**