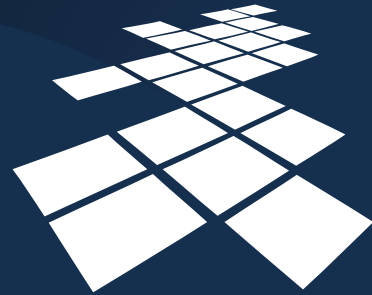


Pamięć wirtualna

**Wykład prowadzą:
Jerzy Brzeziński
Dariusz Wawrzyniak**



**UCZELNIA
ONLINE**

Pamięć wirtualna jest organizacją zasobów pamięci, zrealizowaną w oparciu o tzw. przestrzeń wymiany w pamięci drugiego rzędu (na dysku). Pamięć operacyjna (fizyczna) jest dla tych zasobów tylko pewnym oknem, przechowującym część zawartości na potrzeby bieżącego przetwarzania.

Stosowanie pamięci wirtualnej ma wiele zalet nie tylko związanych z możliwością powiększenia zasobów pamięci ponad dostępną pamięć fizyczną. Umożliwia bardziej racjonalne wykorzystanie pamięci operacyjnej, gdyż programy tworzone są często z nadmiarem w stosunku do typowych potrzeb. Na przykład rozmiary tablic statycznych ustala się z nadmiarem w stosunku do typowych potrzeb, w kodzie uwzględnia się obsługę sytuacji wyjątkowych do których może nigdy nie dojść. Ten nadmiar nie musi być ładowany do pamięci. Zastosowanie pamięci wirtualnej może też zmniejszyć czas odpowiedzi, gdyż skraca czas ładowania kodu, który często odwzorowywany jest w przestrzeń adresową procesu bezpośrednio z pliku i sprowadzany w niewielkich porcjach na żądanie.

Celem wykładu jest przedstawienie zasady działania i problemów realizacji pamięci wirtualnej oraz omówienie algorytmów wymiany stron pomiędzy pamięcią operacyjną a pamięcią drugiego rzędu (obszarem wymiany).

- Stronicowanie na żądanie
 - obsługa błędu strony
 - wymiana
- Problemy realizacji stronicowania na żądanie
 - problem wyboru ofiary
 - problem wznawiania rozkazów
- Algorytmy wymiany
 - algorytmy wymiany na żądanie
 - algorytmy wymiany ze sprowadzaniem na żądanie
 - algorytmy wstępnego stronicowania

Podstawą funkcjonowania pamięci wirtualnej jest mechanizm stronicowania na żądanie, od omówienia którego rozpoczyna się wykład. Działanie mechanizmu oparte jest na stronicowaniu i polega na sprowadzaniu do pamięci operacyjnej stron adresowanych przez procesor. Sprowadzenie strony jest zadaniem systemu operacyjnego i realizowane jest w przypadku wystąpienia błędu strony. Ponieważ pamięć operacyjna jest na ogół mniejsza od pojemności zasobów pamięci wirtualnej, sprowadzenie może wymagać usunięcia innej strony — dochodzi wówczas do wymiany. Realizacja pamięci wirtualnej oprócz wsparcia sprzętowego wymaga rozwiązania dwóch zasadniczych problemów na poziomie systemu operacyjnego:

- problemu wyboru ramki ofiary do wymiany strony, jeśli zajdzie potrzeba wymiany,
- problemu wznawiania rozkazów, którego rozwiązanie sprowadza się do zapewnienia dostępności odpowiednio dużej liczby ramek dla procesu.

Rozwiązanie problem wyboru ofiary bazuje na przesłankach o charakterze losowym i związane jest ściśle z algorytmem wymiany. Klasyfikacja i omówienie algorytmów wymiany stanowią ostatnią część wykładu.

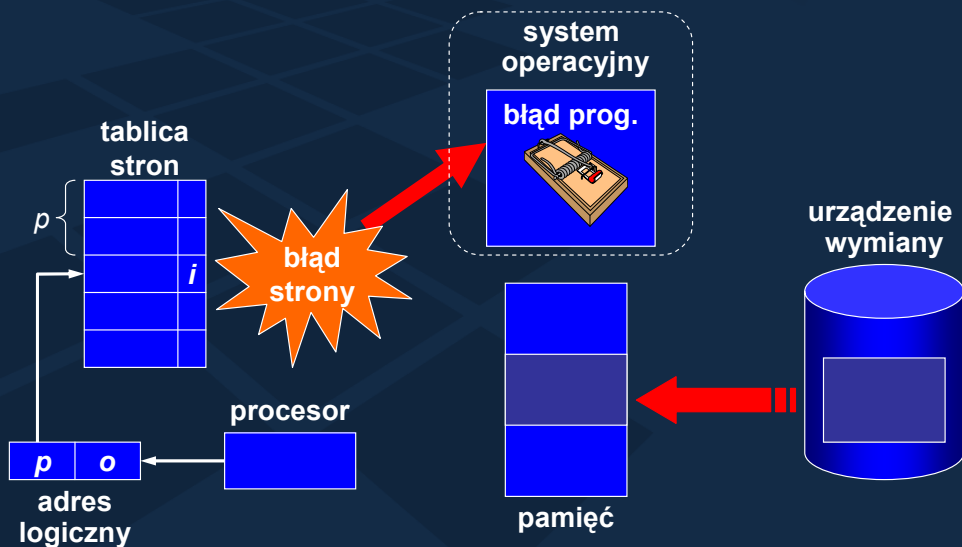
- Działanie mechanizmu: strony są sprowadzane do pamięci tylko wówczas, gdy jest to konieczne, czyli wówczas gdy następuje odniesienie do komórki o adresie, znajdującym się na tej stronie
- Wymaganie sprzętowe
 - tablica stron z bitem poprawności (ang. valid-invalid bit) dla każdej pozycji (dodatkowo z bitem modyfikacji i odniesienia),
 - mechanizm reakcji na odniesienie do strony niepoprawnej,
 - urządzenie wymiany (ang. swap device) — pamięć pomocnicza.

Stronicowanie na żądanie związane jest przede wszystkim z wymianą stron pomiędzy pamięcią pierwszego rzędu (operacyjną, fizyczną) i drugiego rzędu (masową, dyskową). Dzięki wykorzystaniu pamięci masowej można rozszerzyć wirtualną przestrzeń adresową i tym samym zwiększyć stopień wieloprogramowości lub uruchamiać zadania, których rozmiar wykracza poza dostępną pamięć operacyjną. Kosztem wprowadzenia takiego mechanizmu jest złożoność zarządzania pamięcią i narzut czasowy związany z dostępem, wynikający z wykorzystania stosunkowo wolnej pamięci dyskowej.

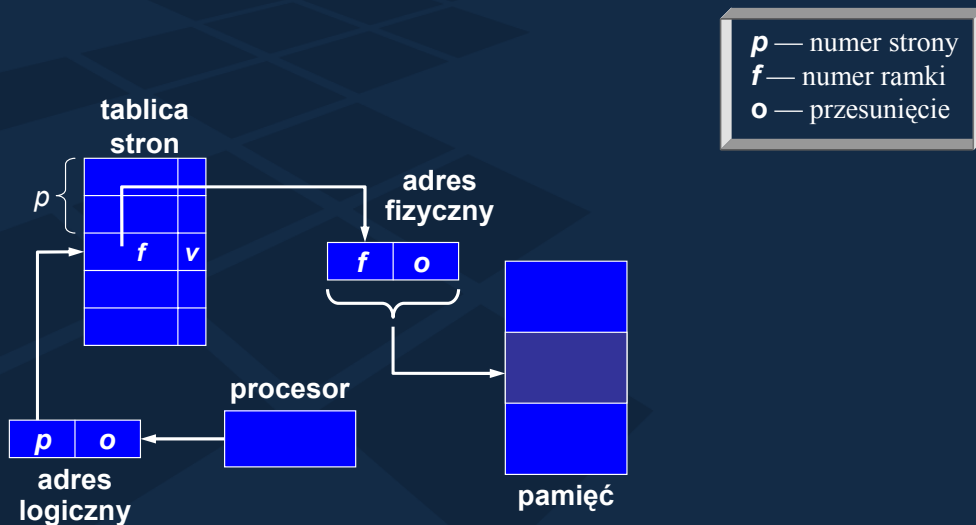
Adresowanie stron odbywa się tak samo, jak w prostym stronicowaniu, omówionym w poprzednim module. W tablicy stron przechowywany jest jednak bit poprawności, informujący o stanie strony. Strona poprawna (ang. valid) to taka, która zlokalizowana jest w pamięci operacyjnej. Odniesienie do takiej strony nie wymaga jej sprowadzania. Jeśli zgodnie z wartością bitu poprawności strona jest uznana za niepoprawną, występuje błąd braku strony i zgłaszane jest odpowiednie przerwanie diagnostyczne. W ramach obsługi błędu strony następuje sprowadzenie strony z obszaru wymiany (ang. swap space), umieszczenie w wolnej ramce i ponowne wykonanie rozkazu. Takie działanie nazywa się *leniwą wymianą* (ang. lazy swapping).

Urządzeniem wymiany jest najczęściej dysk, na którym znajduje się plik wymiany lub specjalnie wyodrębniona strefa (tzw. partycja wymiany).

W dalszej części wykładu zamiennie będzie używane sformułowanie: *adresowanie strony* i *odniesienie do strony*, oznaczające wystawienie przez procesor adresu logicznego komórki, zlokalizowanej na tej stronie.



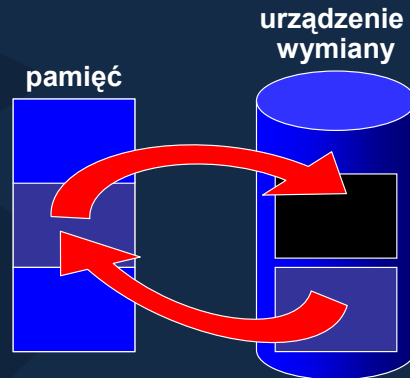
Zgodnie z zasadą transformacji adresu w pamięci stronicowanej, w celu zamiany numeru strony na numer ramki lokalizowany jest odpowiedni wpis w tablicy stron. Wpis może być jednak w danej chwili niepoprawny (ang. invalid), gdyż strona mogła zostać usunięta z pamięci — zastąpiona inną stroną. W przypadku odwołania się do strony, opisanej jako niepoprawna, następuje zgłoszenie błędu (przerwania diagnostycznego), które obsługuje system operacyjny. W ramach obsługi adresowana strona sprowadzana jest z obszaru wymiany do pamięci operacyjnej.



Po sprowadzeniu strony do pamięci operacyjnej odpowiedni bit poprawności w tablicy stron jest ustawiany. Przyjmijmy, że ustawienie bitu poprawności oznacza, że strona jest poprawna (ang. valid). Następnie rozkaz, który spowodował błąd strony, wykonywany jest ponownie przez procesor. To odniesienie do pamięci, które spowodowało wcześniej błąd strony powinno teraz wykonać się poprawnie.

Problem zastępowania (wymiany, ang. page replacement) stron pojawia się w, gdy w pamięci fizycznej brakuje wolnych ramek i konieczne jest zwolnienie jakieś ramki poprzez usunięcie z niej strony.

Jeśli strona była modyfikowana w pamięci, konieczne jest zapisanie jej w obszarze wymiany ⇒ konieczność wprowadzenia bitu modyfikacji (ang. modify bit), zwanego też bitem zabrudzenia (ang. dirty bit).



Pamięć wirtualna (6)

Ramkę, która jest użyta do wymiany określa się jako *ramkę ofiarę* (ang. victim frame), chociaż to raczej strona jest ofiarą. W dalszej części używane będą niezależnie od kontekstu terminy: *ramka ofiara*, *strona ofiara*, lub po prostu *ofiara*.

Brak bitu modyfikacji w tablicy stron oznacza konieczność zapisania strony na urządzeniu wymiany, jeśli strona to mogła być potencjalnie modyfikowana. Jeśli bit modyfikacji istnieje, ustawiany jest przed jednostką zarządzania pamięcią zawsze, gdy wystąpił cykl maszynowy zapisu na tej stronie, nawet jeśli zapis niczego w stanie strony nie zmienił (np. nastąpiło wpisanie do komórki pamięci takiej samej wartości, jaka była tam wcześniej).

Jeśli bit modyfikacji nie jest ustawiony, to znaczy, że w obszarze wymiany jest aktualna kopia strony, która znajduje się również w pamięci. Zastąpienie takiej strony sprowadza się do jej nadpisania w pamięci. Przy ustawionym bicie modyfikacji strona zastępowana musi być najpierw zapisana w obszarze wymiany, co istotnie zwiększa koszt operacji wymiany.

- Całkowity koszt sprowadzenia stron w ogólnym przypadku dany jest wzorem:

$$\sum_t h(k_t)$$

gdzie:

- t — dyskretna chwila czasu, w której następuje odniesienie do pamięci (do strony)
- k_t — liczba stron sprowadzanych w chwili t ,
- $h(k)$ — koszt jednorazowego sprowadzenia grupy k stron, przy czym $h(0) = 0$, $h(1) = 1$.

Z punktu widzenia efektywności przetwarzania koszt wymiany można utożsamiać z czasem realizacji wymiany. W celu uproszczenia analizy przyjmuje się, że koszt usunięcia strony stanowi stałą część kosztu jej sprowadzenia. Koszt wynika zatem z czasu sprowadzania, który w ogólności zależy od:

- ciągu odniesień,
- liczby dostępnych ramek,
- algorytmu wymiany.

Przyjmując funkcję kosztu sprowadzenia zbioru stron — h , koszt realizacji wymiany jest po prostu sumą kosztów w chwilach odniesienia do pamięci. Jeśli żadna strona nie jest sprowadzana, koszt wynosi 0. Jeśli sprowadzana jest jedna strona, koszt wynosi 1.

- W przypadku zastosowania dysku jako urządzenia wymiany, na czas sprowadzania wpływ mają:
 - T_w — czas oczekiwania (suma czasu oczekiwania w kolejce do urządzenia oraz czasu przygotowania urządzenia do transmisji)
 - T_{tr} — czas dostępu do danych
- Postać funkcji $h(k)$ dla dysku jest zatem następująca:

$$h(k) = T_w + k \cdot T_{tr}$$
- Własność funkcji $h(k)$ przy założeniu, że $T_w > 0$ i $k > 0$:

$$h(k) < k \cdot (T_w + T_{tr}) \Rightarrow 1 \leq h(k) \leq k \cdot h(1)$$

W przypadku dysku żądanie odczytu bloków czeka w kolejce, aż urządzenie dyskowe będzie wolne, czyli skończy obsługę wcześniej zgłoszonych żądań. Czas oczekiwania nie zależy od wielkości zgłaszanego żądania, ale wielkość ta ma wpływ na łączny czas dostępu, na który składa się czas wyszukiwania ścieżki, opóźnienie obrotowe, czas przesyłania. Czas dostępu dotyczy każdego sektora, można jedynie minimalizować czas wyszukiwania i opóźnienie obrotowe, gdyż sektory z zawartością strony są lokalizowane blisko siebie.

Jak wynika z własności funkcji kosztu dla dysku, sprowadzenie kilku stron w wyniku realizacji **jednego żądania** jest bardziej kosztowne niż sprowadzenie jednej strony, ale mniej kosztowne niż sprowadzanie poszczególnych stron w wyniku osobnych żądań.

- Koszt wymiany można aproksymować za pomocą parametrów FN (liczba wygenerowanych błędów strony) i TN (liczba transmisji stron).
- Całkowity czas oczekiwania na urządzenie (istotna składowa czasu wymiany) również jest zależny od liczby błędów strony.
- Wniosek: należy minimalizować liczbę błędów strony i tym samym redukować czas realizacji wymiany, czyli koszt.

Konieczność sprowadzenia strony pojawia się dopiero przy wystąpieniu błędu strony. Sprowadzanie stron w okolicznościach innych niż obsługa błędu strony nie ma sensu, ale jak wynika z zaprezentowanej analizy, w ramach obsługi błędu strony sensowne może być sprowadzenie kilku stron, przewidując przeszłe odniesienia. Koszt sprowadzenia jest więc ściśle uzależniony od liczby błędów strony i liczby transmisji stron.

Liczba błędów strony wpływa również na długość kolejki żądań dostępu do dysku, co wydłuża czas oczekiwania. Można zatem wyciągnąć ogólny wniosek, że w celu redukcji kosztów wymiany należy minimalizować liczbę błędów strony, gdyż wpływa ona pośrednio lub bezpośrednio na wszystkie składowe koszty.

- Problem wyboru ofiary — niewłaściwy wybór ramki-ofiary powoduje wzrost kosztu wymiany.
 - W skrajnym przypadku może dojść do zjawiska migotania, w przypadku którego często dochodzi do wystąpienia odniesienia do właśnie usuniętej strony.
- Problem wznawiania rozkazów — w przypadku wielokrotnego odniesienia do pamięci w jednym cyklu rozkazowym należy zapewnić, że wszystkie adresowane strony są jednocześnie dostępne w ramach w pamięci fizycznej.

Zasadnicze problemy wymiany stron dotyczą decyzji odnośnie usuwanej strony oraz zagwarantowania dostępności wszystkich stron wymaganych do realizacji cyklu rozkazowego.

Usunięcie z pamięci strony, która będzie potrzebna przyszłości, oznacza konieczność ponownego sprowadzenia jej do pamięci, a więc koszt, którego być może dałoby się uniknąć. Nasilenie tego zjawiska określane jest jako *migotanie* i oznacza drastyczny spadek efektywności działania systemu komputerowego, gdyż większość czasu cyklu przetwarzania proces spędza w stanie oczekiwania na gotowość urządzenia dyskowego. Rozwiązanie problemu wyboru ofiary oparte jest na przesłankach o charakterze losowym.

Problem wznawiania rozkazów wiąże się z koniecznością ponownego wykonania całego cyklu rozkazowego, w którym wystąpił błąd strony. Cykl rozkazowy może składać się z kilku cykli maszynowych, z których każdy może się wiązać z odniesieniem do innej strony pamięci. W niesprzyjających okolicznościach mogłoby dojść do sytuacji, w której w wyniku obsługi błędu strony następuje usunięcie z pamięci innej strony, adresowanej w tym samym cyklu rozkazowym. Wznowienie rozkazu ponownie zakończy się błędem strony w jednym z cykli maszynowych. W skrajnym przypadku rozkaz mógłby się nigdy nie wykonać. Problem wznawiania rozkazów można rozwiązać, zapewniając procesowi pewną minimalną liczbę ramek — nie mniejszą niż liczba różnych adresów, wystawianych na magistrali w jednym cyklu rozkazowym.

- Zakładając, że przyszły ciąg odniesień do pamięci nie jest znany, na podstawie historii odniesień należy wybrać taką ramkę, do której prawdopodobieństwo odniesienia w przyszłości jest małe.
- Podstawowa własność programów, na podstawie której można szacować takie prawdopodobieństwo, nazywana jest lokalnością.

Problem z wyborem ofiary wynika z faktu, że nie znamy przyszłego ciągu odniesień do stron. Przyszłe odniesienia możemy jedynie przewidywać z pewnym prawdopodobieństwem na podstawie różnych przesłanek z odniesień w przeszłości. Większość typowych programów charakteryzuje się własnością lokalności. Często również optymalizatory kodu dokonują takiego rozmieszczenia obiektów, aby w trakcie wykonywania ujawniała się taka własność.

- Lokalność czasowa — tendencja procesów do generowania w stosunkowo długich przedziałach czasu odniesień do niewielkiego podzbioru stron wirtualnych zwanego zbiorem stron aktywnych.
 - Formalnie jest to tendencja procesu do generowania z dużym prawdopodobieństwem w przedziale czasu $(t, t + \tau)$ odniesień do stron adresowanych w przedziale czasu $(t - \tau, t)$.
- Lokalność przestrzenna — tendencja procesu do generowania z dużym prawdopodobieństwem kolejnych odniesień do stron o zbliżonych numerach (stron sąsiednich) lub stron o numerach skojarzonych w trakcie przetwarzania.

Lokalność można rozważać w dwóch aspektach: czasowym i przestrzennym. Lokalność czasową można sprowadzić do wielokrotnego odwołania w krótkim czasie do tej samej strony. Przypadek taki ma miejsce przy wykonywaniu programu, chyba że występują częste i dość dalekie (wychodzące poza stronę) skoki. Przypadek taki dotyczy również danych, np. tablic, a nawet pojedynczych zmiennych, gdyż programiści definiują często kilka zmiennych na potrzeby obsługi jakiegoś krótkiego fragmentu programu, np. pętli. Z punktu widzenia stronicowania oznacza to, że jeśli strona zostanie sprowadzona do pamięci operacyjnej, przyda się wielokrotnie w czasie wykonywania określonego fragmentu programu. Własność lokalności czasowej wykorzystywana jest często przy podejmowaniu decyzji o usunięciu strony z pamięci.

Lokalność przestrzenna sprowadza się do wnioskowania na podstawie odniesień do pewnych stron o przyszłych odniesieniach do innych stron. Podstawą takiego skojarzenia może być numeracja stron. Jeśli program lub większa struktura danych (tablica) zajmuje kilka kolejnych stron, to prawdopodobne jest, że po odniesieniu do pierwszej z tych stron nastąpi odniesienie do drugiej, a później do trzeciej. Skojarzenie takie można również oprzeć o obserwację wcześniejszych odniesień. Jeśli strona p_2 często była adresowana zaraz po stronie p_1 , to być może na stronie p_1 jest skok do instrukcji, znajdującej się na stronie p_2 . Ten rodzaj lokalności można wykorzystać przy podejmowaniu decyzji o usuwaniu, ale również przy sprowadzaniu stron, kiedy stosowane jest tzw. sprowadzanie wstępne (sprowadzanie kilku stron w wyniku jednego błędu strony).

- Efektywność działania systemu pamięci wirtualnej zależy od precyzji identyfikacji zbioru stron aktywnych i możliwości utrzymania ich w pamięci fizycznej.
- Wobec braku a priori pełnego ciągu odniesień do stron wirtualnych identyfikacja takiego zbioru może wynikać z różnych przesłanek, czego skutkiem jest duża różnorodność algorytmów wymiany.

Podstawą efektywnego funkcjonowania pamięci wirtualnej jest oczywiście precyzyjna identyfikacja zbioru stron aktywnych. Wykorzystanie tej informacji polega na utrzymaniu stron aktywnych w pamięci. Jeśli ze względu na zbyt duże potrzeby procesów strony aktywne muszą być usuwane z pamięci, to i tak trzeba ponieść koszt ich ponownego sprowadzenia. W tego typu sytuacji dochodzi najczęściej do migotania stron (szamotania, ang. trashing), jednak identyfikacja stron aktywnych umożliwia wykrycie takiego ryzyka i podjęcie pewnych działań zapobiegawczych.

Sama identyfikacja wynikać może z różnych przesłanek, a precyzja tej identyfikacji uzależniona jest od zasadności przyjętych przesłanek.

- Algorytmy wymiany na żądanie
 - sprowadzenia odbywa się na żądanie — strona sprowadzana jest dopiero wówczas, gdy następuje odniesienie do niej i nie ma jej w pamięci,
 - usuwanie odbywa się na żądanie — strona usuwana jest wówczas, gdy konieczne jest sprowadzenie innej strony w wyniku odniesienia i nie ma wolnej ramki
- Algorytmy wymiany ze sprowadzaniem na żądanie
 - tylko sprowadzanie odbywa się na żądanie
- Algorytmy wstępnego sprowadzania
 - sprowadzana jest strona żądana, a wraz z nią inne strony

Klasyfikacja związana jest z okolicznościami, w jakich podejmowana jest decyzja o sprowadzeniu lub usunięciu strony. W algorytmach wymiany na żądanie zarówno sprowadzanie, jak i usuwanie wykonywane jest wówczas, gdy jest absolutna konieczność. W przypadku sprowadzania oznacza to, że jeśli strony nie ma w pamięci, to nie zostanie sprowadzona wcześniej, niż po wystąpieniu odniesienia do niej (zaadresowania jej przez procesor). W przypadku usuwania absolutna konieczność występuje dopiero wówczas, gdy brakuje miejsca w pamięci operacyjnej (nie ma wolnej ramki), a ze względu na występujące odniesienia do stron, znajdujących się poza pamięcią operacyjną, konieczne jest zwolnienie ramki na potrzeby sprowadzanej strony. Ponieważ strona do sprowadzenia jest jednoznacznie określona przez stan systemu (stan pamięci oraz adres wystawiony przez procesor), specyfika algorytmu tej klasy zależy od wyboru strony usuwanej z pamięci.

W klasie algorytmów wymiany ze sprowadzaniem na żądanie o specyfice algorytmu również decyduje sposób wyboru strony usuwanej, a często stron usuwanych. W tych algorytmach strona (lub kilka stron) może być usunięta w dowolnym momencie, niekoniecznie przy okazji obsługi błędu strony.

W klasie algorytmów ze wstępnym sprowadzaniem odniesienie do strony nieobecnej w pamięci operacyjnej skutkuje błędem strony, ale przy okazji jego obsługi może nastąpić sprowadzenie oprócz strony żądanej innych stron. Istotą algorytmów tej klasy jest wybór stron wstępnie sprowadzanych, można więc przyjąć dowolny sposób usuwania, łącząc w ten sposób ideę wstępnego sprowadzania z różnymi koncepcjami utrzymania stron w pamięci.

- Zastępowanie lokalne (ang. local replacement) — algorytm wymiany zastępuje tylko strony w ramach przydzielonych procesowi, który spowodował błąd strony.
- Zastępowanie globalne (ang. global replacement) — algorytm wymiany zastępuje strony znajdujące się w dostępnej puli ramek w całym systemie (w szczególności zatem usuwa strony innych procesów).

Klasyfikacja wiąże się z zakresem ramek uwzględnianych przy poszukiwaniu ofiary. W przypadku zastępowania lokalnego proces dysponuje pewną liczbą ramek, w których muszą się pomieścić jego strony. Błąd strony tego procesu skutkuje ewentualnym usunięciem innej jego strony. Można w ten sposób ograniczyć (ale nie zlikwidować) negatywny wpływ zbyt częstego generowania błędów strony przez proces na efektywność działania całego systemu.

W zastępowaniu globalnym poszukiwanie ramki ofiary dotyczy całej puli ramek niezależnie od bieżącego przydziału. Może więc dojść do podkradania ramek, kiedy jeden proces traci ramkę na rzecz innego. Takie podejście ułatwia adaptację liczby ramek używanych przez procesy do zróżnicowanych potrzeb, ale jest niebezpieczne dla efektywności całego systemu, gdy liczba ramek jest zbyt mała. Może wówczas dojść do zjawiska szamotania.

- Przydział statyczny — liczba ramek przydzielonych procesowi jest ustalona i nie ulega zmianie w trakcie przetwarzania.
- Przydział dynamiczny — liczba ramek przydzielonych procesowi może się zmienić w trakcie przetwarzania.

Klasyfikacja związana jest z adaptacją przydziału zasobów do zmieniających się potrzeb procesów. W przydziale statycznym proces otrzymuje pewną liczbę ramek na cały czas cyklu przetwarzania. W przydziale dynamicznym liczba ramek przydzielonych procesowi może się zmienić zależnie od intensywności generowania błędów strony. Jeśli liczba błędów strony, generowanych przez proces, jest stosunkowo duża to w zależności od możliwości, jakimi dysponuje system, można przydzielić dodatkowe ramki. Dodatkowe ramki mogą pochodzić z puli ramek odebranych procesom, które generowały mało błędów strony.

Zastępowanie globalne w naturalny sposób prowadzi do przydziału dynamicznego, nie ma natomiast sensu w przypadku przydziału statycznego. Przydział statyczny wymaga zastępowania lokalnego, ale w przypadku zastępowania lokalnego również można zastosować przydział dynamiczny. W ramach obsługi błędu strony wymiana jest lokalna, ale okresowo liczba ramek może się np. zwiększyć, co wyeliminuje na jakiś czas problem zastępowania w przypadku kolejnych błędów strony. Liczba ramek przydzielonych procesowi może też oczywiście się zmniejszyć.

- Minimalna liczba ramek — zdefiniowana przez architekturę komputera (zależna od maksymalnej liczby komórek adresowanych przez jeden rozkaz).
- Liczba ramek przydzielona dla procesu
 - podział równomierny (ang. equal allocation)
 - podział proporcjonalny (ang. proportional allocation)
 - przydział zależny od priorytetu procesu

Z przydziałem ramek wiąże się ustalanie ich liczby. Jest to szczególnie istotne w przydziale statycznym, ale dotyczy również początkowej liczby ramek w przydziale dynamicznym. Niezależnie od sposobu przydziału proces, który jest w pamięci powinien mieć do dyspozycji pewną minimalną liczbę ramek, gwarantującą uniknięcie problemu wznawiania rozkazów. Liczba ramek powinna być nie mniejsza, niż maksymalna liczba różnych adresów, jakie mogą być wystawione na magistrali w jednym cyklu rozkazowym.

Zakładając, że spełnione jest minimum dla każdego procesu przebywającego w pamięci, liczba ramek może być:

- w przypadku przydziału równomiernego równa dla wszystkich procesów,
- proporcjonalna do rozmiaru obrazu procesu,
- uzależniona od priorytetu procesu — proces o wyższym priorytecie będzie miał więcej ramek, żeby jego przetwarzania odbywało się sprawniej.

- MIN — zastępowana jest strona, która najdłużej nie będzie używana (optymalny w tej klasie)
- FIFO (ang. First In First Out) — zastępowana jest strona najstarsza (najwcześniej sprowadzona)
- LIFO (ang. Last In First Out) — zastępowana jest strona najmłodsza (najpóźniej sprowadzona)
- LRU (ang. Least Recently Used) — zastępowana jest najdawniej użyta strona (najdłużej nie używana)
- LFU (ang. Least Frequently Used) — zastępowana jest najrzadziej używana strona
- MFU (ang. Most Frequently Used) — zastępowana jest najczęściej używana strona

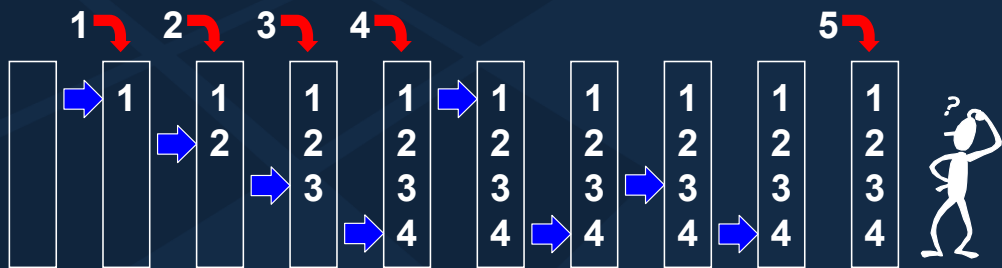
Algorytmy wymiany na żądanie stosowane są w systemach jednozadaniowych lub w systemach wielozadaniowych ze statycznym przydziałem ramek.

Algorytm MIN oparty jest na nierealnej przesłance, wymagającej znajomości **przyszłego** ciągu odniesień. Z drugiej strony jest to algorytm optymalny w tej klasie, dlatego wykorzystywany jest dla celów porównawczych. Można w ten sposób sprawdzać, ile tracimy, opierając się na przesłankach z historii dotychczasowych odniesień. Jest to zatem swego rodzaju miara zasadności tych przesłanek.

Algorytmy FIFO i LIFO oparte są na kolejności sprowadzania stron do pamięci. FIFO usuwa strony w kolejności ich sprowadzania LIFO w kolejności odwrotnej. FIFO sprawdza się dobrze w przypadku programów, w których jest prosty, sekwencyjny przepływ sterowania od początku programu do końca, z małą liczbą pętli, czy wywołań podprogramów. LIFO natomiast właściwy jest dla przypadków pętli, gdyż ponowne przejście sterowania do tej samej instrukcji nastąpi dopiero w następnej iteracji. Dla pętli istnieje jeszcze inny algorytm — LD (ang. loop detection), którego prezentację tu pominięto.

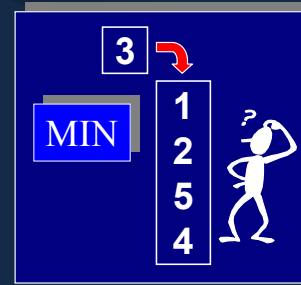
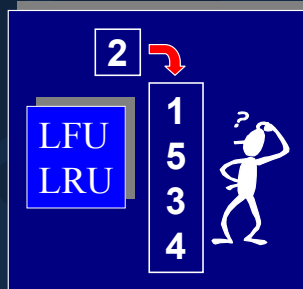
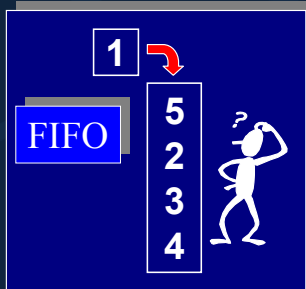
Algorytmy LRU, LFU i MFU oparte są na przesłankach, wymagających monitorowania odniesień do pamięci. Dla LRU istotny jest czas ostatniego odniesienia, a dla LFU i MFU liczba odniesień w przeszłości. LRU jest typowym algorytmem dla programów, charakteryzujących się lokalnością czasową odniesień do pamięci. Algorytmy LFU i MFU (tzw. algorytmy licznikowe) oparte są na zupełnie przeciwnych przesłankach: LFU usuwa stronę, do której było najmniej odniesień do początku przetwarzania lub od momentu sprowadzenia do pamięci (są to dwa warianty algorytmów licznikowych), a MFU usuwa stronę, do której było najwięcej odniesień.

- W systemie pamięci wirtualnej są 4 ramki.
- Wszystkie ramki są początkowo puste
- W systemie pojawiają się następujące odniesień (odwołań) do stron: 1, 2, 3, 4, 1, 4, 3, 4, 5, 2, 1, 4, 3, 4



Przedstawiony przykład pokazuje przebieg odniesień do stron o numerach od 1 do 5 w obrębie 4 początkowo pustych ramek. Pierwsze 4 odniesienia powodują błąd strony, ale nie występuje problem zastępowania, gdyż dostępne są wolne ramki. Dopiero odniesienie do strony nr 5 wymaga zwolnienia jakiejś ramki. Wybór strony do usunięcia zależy od algorytmu wymiany.

Dalszy ciąg odniesień:
2, 1, 4, 3, 4



Pamięć wirtualna (20)

Na kolejnych rysunkach pokazany jest stan poszczególnych ramek, zależnie od zastosowanego algorytmu. Przy okazji pokazane jest odniesienie do pamięci, które spowoduje kolejny błąd błęd strony i tym samym problem zastępowania.

Usunięcie tej samej strony w przypadku algorytmów LFU i LRU oraz MFU i LIFO jest przypadkowe i kolejna wymiana w przypadku tych algorytmów może wykazać różnice.

Warto też przy tej okazji zwrócić uwagę, że dla algorytmów licznikowych wybór może być niejednoznaczny, gdyż liczba odniesień do różnych stron może być taka sama. Jako kryterium rozstrzygnięcia można przyjąć czas sprowadzenia, czyli strony z taką samą wartością licznika usuwane są w kolejności FIFO.

Przy założeniu, że strony sprowadzane są pojedynczo nie ma tego problemu w algorytmach opartych na kolejności zdarzeń w czasie, czyli MIN, FIFO, LIFO i LRU.

Następujący ciąg odniesień obsługiwany jest zgodnie z algorytmem FIFO w systemie z 3, a następnie 4 ramkami:
1, 2, 3, 4, 1, 2, 5, 1, 2, 3, 4, 5



Pamięć wirtualna (21)

W przypadku 3 ramek pierwsze 7 odniesień do stron kończy się błędem, ale kolejna 2 odniesienia przebiegają bez błędów, a później pojawiają się jeszcze 2 błędy. W sumie jest ich 9.

W przypadku 4 ramek początek jest bardziej optymistyczny, gdyż po 4 pierwszych odniesieniach z błędem, wypełnione ramki zdają się gwarantować stabilność przetwarzania. Jednak sprowadzenie strony nr 5 powoduje usunięcie strony numer 1, która sprowadzana jest jako następna, usuwając z kolei potrzebną później stronę nr 2 itd. Ostatecznie mamy 10 błędów strony pomimo **zwiększenia liczby ramek**. Paradoks ten, nazywany anomalią Belady'ego, dotyczy wybranych algorytmów, między innymi algorytmu FIFO. Anomalia nie ujawnia się być może zbyt często, ale może wzbudzać obiekcje co do efektów przydziału dodatkowych zasobów.

- Implementacja algorytmu FIFO
- Implementacja algorytmu LRU
- Algorytmy przybliżające metodę LRU
 - algorytm dodatkowych bitów odwołań
 - algorytm drugiej szansy (FINUFO)
 - ulepszony algorytm drugiej szansy

Jak już wcześniej wspomniano, kluczowa dla efektywności funkcjonowania systemu z pamięcią wirtualną jest minimalizacja liczby błędów strony. W celu ograniczenia kosztów czasowych dostępu do pamięci stosowane są odpowiednie algorytmy wymiany. Na przykładzie algorytmów FIFO i LRU wyjaśnione zostaną natomiast zagadnienia związane z **kosztem działania samego algorytmu**.

- Utrzymywanie listy numerów stron w kolejności ich sprowadzania do pamięci
- Umieszczanie numeru sprowadzanej strony na końcu listy
- Usuwanie z pamięci (i z listy) strony, której numer znajduje się na początku listy

Implementacja algorytmu FIFO bazuje na kolejce FIFO, którą można łatwo zaimplementować za pomocą listy jednokierunkowej.

Warto zwrócić uwagę, że aktualizacja listy następuje w ramach obsługi błędu strony, gdyż stronę umieszczamy na liście po sprowadzeniu. W przypadku usuwania na żądanie, również usuwanie strony z listy wykonywane jest przy obsłudze błędu strony.

- Licznik — przy każdym odniesieniu do pamięci zwiększana jest wartość pewnego licznika i wpisywana do odpowiedniej pozycji opisującej stronę w tablicy stron (lub w innej specjalnej strukturze systemu operacyjnego). Z pamięci usuwana jest wówczas strona z najmniejszą wartością tego licznika, co wymaga przejrzania całej tablicy stron.
- Stos — numery stron, do których następuje odwołanie, odkładane są na szczycie stosu. Przed odłożeniem na szczycie numer strony musi być wydobyty ze środka stosu, czyli z miejsca, gdzie był ostatnio odłożony. W tej implementacji z pamięci usuwana jest strona, która jest na dnie stosu.

Rozwiązanie z licznikiem wymaga zwiększenia licznika odniesień do pamięci i stosownej aktualizacji opisu strony (lub ramki). Jest to operacja stosunkowo szybka w porównaniu z wyszukiwaniem strony do usunięcia w przypadku konieczności zwolnienia ramki. Pewnym problemem jest ryzyko przepełnienia licznika.

W przypadku rozwiązania opartego na stosie jest odwrotnie. Aktualizacja stosu zbudowanego na liście wymaga zmiany zaledwie kilku wskaźników, ale wyszukanie strony w głębi stosu jest bardziej czasochłonne. Wybór strony do usunięcia w przypadku konieczności zwolnienia ramki jest za to natychmiastowy.

Porównując oba podejścia należałoby zauważyć, że błąd strony jest zjawiskiem nieporównywalnie rzadszym niż odniesienie do strony. W przeciwnym przypadku stosowanie pamięci wirtualnej nie miałoby żadnego sensu. Preferowane byłoby zatem rozwiązanie licznikowe. Jednak implementacja tego rozwiązania na poziomie systemu operacyjnego wymaga przekazania sterowania do programu jądra **przy każdym odniesieniu do strony**. Nawet jeśli aktualizacja licznika i tablicy stron wymaga wykonania zaledwie kilku rozkazów, to i tak powoduje to kilkunastokrotne spowolnienie **każdego dostępu** do pamięci. Programowe rozwiązanie oparte na stosie byłoby jeszcze bardziej czasochłonne.

Mogłoby się zatem okazać, że koszt algorytmu LRU pochłania zysk, wynikający z trafności decyzji co do stron wymienianych. Konieczne jest zatem odpowiednie wsparcie na poziomie architektury komputera.

- Niezbędne wspomaganie sprzętowe:
 - bit odniesienia (ang. reference bit) — ustawiany, gdy następuje odniesienie do strony,
 - bit modyfikacji (ang. modify bit) — ustawiany, gdy następuje zapis na stronie.
- Algorytmy korzystające ze wspomagania sprzętowego:
 - algorytm dodatkowych bitów odniesienia — wykorzystuje bit odniesienia,
 - algorytm drugiej szansy — wykorzystuje bit odniesienia,
 - ulepszony algorytm drugiej szansy — wykorzystuje bit odniesienia i bit modyfikacji.

Implementacja na poziomie architektury komputera nie może być zbyt kosztowana, gdyż nie wiadomo, jaki algorytm wymiany będzie implementowany w systemie operacyjnym i czy w ogóle realizowana będzie pamięć wirtualna z wymianą stron. Jeśli zatem dostarczane jest wsparcie dla implementacji wymiany, powinno ono być na tyle uniwersalne, żeby można była je zaadaptować do różnych algorytmów.

W ramach wsparcia na poziomie architektury jednostka zarządzania pamięcią ustawia odpowiednie bity na właściwej pozycji w tablicy stron w przypadku wykrycia odniesienia do strony:

- bit odniesienia (ang. reference bit) — ustawiany dla danej strony zawsze, gdy następuje zapis lub odczyt jakiejś komórki na tej stronie,
- bit modyfikacji (ang. modify bit) — ustawiany dla danej strony zawsze, gdy następuje zapis na tej stronie. (Wspomniany przy omawianiu problemu wymiany).

W dalszej części omówione zostaną 3 algorytmy, w których wykorzystywane jest omówione wspomaganie sprzętowe:

- algorytm dodatkowych bitów odniesienia (wykorzystanie bitu odniesienia),
- algorytm drugiej szansy (wykorzystanie bitu odniesienia),
- ulepszony algorytm drugiej szansy (wykorzystanie bitu odniesienia i bitu modyfikacji).

Poza tym, bity te wykorzystywane są w niektórych algorytmach wymiany ze sprowadzaniem na żądanie.



W algorytmie dodatkowych bitów odniesienia okresowo sprawdzane są bity odniesienia każdej ze stron i kopiowane do dodatkowej struktury, zwanej tablicą dodatkowych bitów odniesienia. Bity kopiowane są na najbardziej znaczącą pozycję, ale przed skopiowaniem następuje logiczne przesunięcie bitów w prawo na każdej pozycji. Bit najmniej znaczący jest tracony. Tablica przechowuje zatem informacje o wystąpieniu błędu strony w kilku ostatnich okresach kontrolnych. Ciąg samych zer oznacza, że nie było żadnego odniesienia w ostatnim czasie. Ogólnie im mniejsza wartość tym bardziej odległe w czasie jest ostatnie odniesienie lub mniejsza jest liczba odniesień. Horyzont czasowy tej informacji zależy od liczby bitów na każdej pozycji oraz długości okresu kontrolnego. Oczywiście im dłuższy okres kontrolny, tym mniej precyzyjna jest informacja, gdyż reprezentowana jest tylko przez 1 bit.

W przypadku konieczności wymiany poszukiwana jest strona, dla której wartość w tablicy jest najmniejsza, a ramka tej strony jest użyta do wymiany. Precyzja działania algorytmu wymaga uwzględnienia również bieżącej wartości bitu odniesienia na najbardziej znaczącej pozycji wektora. Można przyjąć, że tuż przed wybraniem strony ofiary następuje przesunięcie bitów w tablicy i skopiowanie bitów odniesienia.



Algorytm drugiej szansy (FINUFO, ang. First In Not Used First Out) jest jednym z tzw. algorytmów zegarowych (wskazówkowych), w których numery stron (lub ramek, w tym przypadku ramek danego procesu) tworzą listę cykliczną (tarczę zegara), a wskazówka wskazuje kandydatkę do usunięcia z pamięci w przypadku błędu strony. Jeśli jednak strona, wskazywana do usunięcia była ostatnio adresowana (ma ustawiony bit odniesienia), otrzymuje drugą szansę. Bit odniesienia dla tej strony jest kasowany, a wskazówka przesuwa się na stronę następną. Sprowadzana strona zastępuje więc stronę na wskazanej ostatecznie pozycji, po czym wskazówka przesuwa się do następnej strony.

W przedstawionym przykładzie wskazywana jest strona nr 5, ale bit odniesienia do tej strony jest ustawiony, więc po jego skasowaniu wskazówka przesuwa się na stronę nr 2. Dla tej strony bit odniesienia jest skasowany, więc nastąpi jej usunięcie, konkretnie ramka tej strony zostanie użyta do wymiany, a wskazówka przesunie się na stronę nr 4.

W niektórych publikacjach odróżnia się algorytm zegarowy od algorytmu drugiej szansy (FINUFO), przedstawiając ten drugi, jako algorytm FIFO, w którym strona na czole kolejki jest:

- usuwana, jeśli nie ma ustawionego bitu odniesienia, albo
- przesuwana jest na koniec, jeśli jest ustawiony bit odniesienia, który jest przy tym kasowany.

Algorytm zegarowy natomiast przedstawia się jako listę cykliczną ramek, w których wymieniane są strony zgodnie z położeniem wskazówki.



Ulepszenie algorytmu drugiej szansy w tym przypadku polega na uwzględnieniu bitu modyfikacji. W celu redukcji kosztów unika się po prostu wymiany stron modyfikowanych, które muszą być dodatkowo zapisane na urządzeniu wymiany. Wyszukiwanie strony do wymiany polega więc na analizie wartości 2-bitowej, złożonej z bitu odniesienia na bardziej znaczącej pozycji i bitu modyfikacji na mniej znaczącej pozycji. W przedstawionym przykładzie „najlepsza” do usunięcia zgodnie z obiegiem wskazówki jest strona nr 6, gdyż ma wyzerowane oba bity.

- VMIN — usuwane są strony, których koszt utrzymania w pamięci jest większy od kosztu ponownego sprowadzenia
- WS — usuwane są strony, do których nie było odniesień przez określony czas
- WSClock — przybliżona wersja algorytmu WS, oparta na bicie odniesienia
- PFF, VSWS — przydział i zwalnianie ramek procesów realizowane jest na podstawie częstości zgłaszania błędów strony

Istotą algorytmów ze sprowadzaniem na żądanie jest zwalnianie ramek, uznanych za nie wykorzystywane, wcześniej niż to wynika z potrzeb realizacji samej wymiany. Celem wcześniejszego zwalniania jest uzyskanie ramek na potrzeby innych procesów, którym tych ramek brakuje. W skrajnych przypadkach niewystarczająca liczba ramek może być przyczyną zawieszenie procesu i przeniesienia całego obrazu w obszar wymiany. Uzyskanie wolnych ramek umożliwia z kolei aktywację procesu — wprowadzenie do pamięci operacyjnej niezbędnych jego stron.

Algorytm VMIN (ang. Variable-space MIN) jest optymalny w klasie algorytmów ze sprowadzaniem na żądanie, ale oparty jest na znajomości **przyszłego** ciągu odniesień do stron. Decyzja o usunięciu strony z pamięci wynika z porównania kosztu ponownego sprowadzenia, oraz kosztu utrzymania do czasu następnego odniesienia. Jeśli koszt ponownego sprowadzania jest mniejszy, strona jest usuwana.

Algorytm WS (ang. Working Set) oparty jest na koncepcji *zbioru roboczego*, czyli zbioru stron, do których było odniesienie w ostatnim okresie czasu w przeszłości. Długość tego okresu mierzona liczbą odniesień, jest parametrem algorytmu. Algorytm WS wymaga monitorowania odniesień do pamięci, więc wymaga wsparcia sprzętowego. Wykorzystywany jest tutaj bit odniesienia.

Przybliżoną realizacją idei zbioru roboczego, wykorzystującą bit odniesienia, jest również algorytm WSClock.

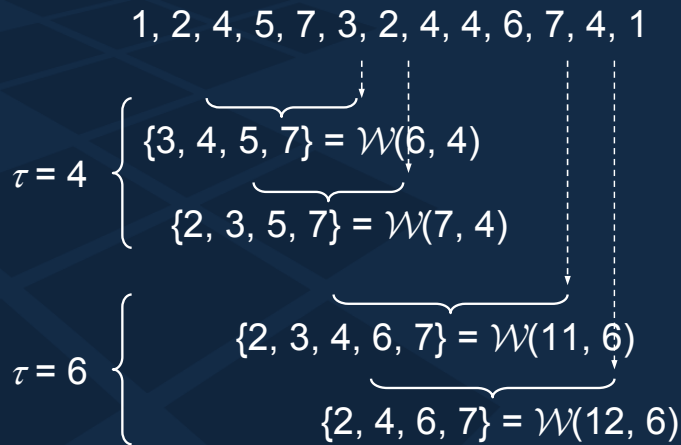
Algorytmy PFF (ang. Page Fault Frequency) i VSWS (ang. Variable-Interval Sampled Working Set) dają efekt podobny do zbioru roboczego, przy czym oparte są na częstość generowania błędów strony.

- Zbiór roboczy procesu (ang. working-set) — zbiór stron, które zostały zaadresowane w ciągu ostatnich τ odniesień do pamięci (w tzw. oknie zbioru roboczego)
- Okno zbioru roboczego (ang. working-set window) — zakres odniesień do pamięci, które adresują strony należące do zbioru roboczego
- Zbiór roboczy w chwili t przy rozmiarze okna τ oznaczony będzie jako $\mathcal{W}(t, \tau)$
- Algorytm zbioru roboczego usuwa z pamięci wszystkie strony, które nie należą do zbioru roboczego.

Przedstawione definicje — zbioru roboczego oraz okna zbioru roboczego — oznaczają właściwie to samo. Pierwsza definiuje zbiór roboczy w oparciu o pojęcie okna, a druga definiuje okno w oparciu o pojęcie zbioru roboczego. Z punktu widzenia strony, pozostaje ona w zbiorze roboczym tak długo, jak długo ostatnie odniesienie do niej pozostaje w bieżącym oknie. Oznacza to, że od momentu ostatniego zaadresowania strony, pozostaje ona w zbiorze roboczym przez τ odniesień.

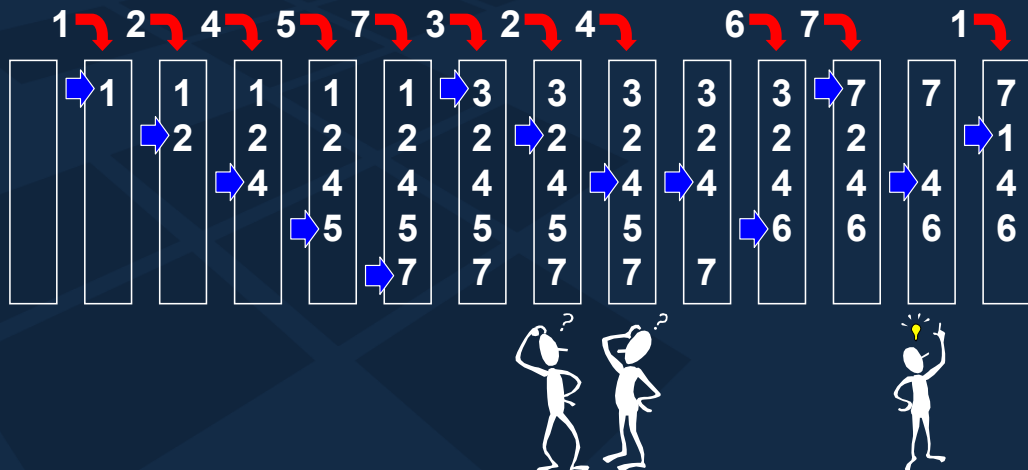
Ogólna (teoretyczna) koncepcja zbioru roboczego polega utrzymywaniu takiego zbioru dla każdego procesu. Strony, należące do zbioru roboczego, pozostawiane są w pamięci, natomiast ramki stron spoza zbioru roboczego są zwalniane. Zapotrzebowanie na ramki w systemie jest sumą mocy zbiorów roboczych wszystkich procesów. Jeśli liczba ta jest większa niż dostępna liczba ramek, może dojść do szamotania. Zadaniem zarządcy pamięci jest wówczas przeniesienie całego obrazu któregoś z procesów w obszar wymiany i odzyskanie jego ramek. Po uzyskaniu odpowiednio dużej liczby wolnych ramek (w przypadku zmniejszenia się zbiorów roboczych) można aktywować któryś z procesów zawieszonych. Wybór procesu do usunięcia lub aktywowania jest zadaniem planisty średnioterminowego.

Osobnym problemem jest dobór wielkości okna. Okno powinno być na tyle duże, żeby objąć tzw. strefę, czyli zbiór stron niezbędnych do wykonania określonego fragmentu programu (np. procedury). W praktyce zależy to od częstości generowania błędów strony przez proces.



Dla przykładowego ciągu odniesień przedstawiono 2 zbiory robocze przy rozmiarze okna 4 oraz 2 zbiory robocze przy rozmiarze okna 6. Ze względu na powtarzające się odniesienia moc zbioru roboczego (liczba elementów) może być mniejsza niż rozmiar okna, co widać w przypadku okna o rozmiarze 6.

Ciąg odniesień: 1, 2, 4, 5, 7, 3, 2, 4, 4, 6, 7, 4, 1
Rozmiar okna: 5



Pamięć wirtualna (32)

Przykład pokazuje realizację ciągu odniesień z poprzedniego slajdu przy rozmiarze okna 5. Proces ma więc do dyspozycji 5 ramek. Początkowo wymiana przebiega tak samo, jak w przypadku FIFO lub LRU. Ciekawym momentem jest drugie odniesienie do strony nr 2. Formalnie strona nr 2 została właśnie usunięta ze zbioru roboczego, ale nastąpiło do niej odniesienie. Jest więc błąd strony, chociaż jego obsługa nie wymaga sprowadzania strony ponownie z pliku wymiany. Podobna sytuacja ma miejsce przy następnym odniesieniu do pamięci — do strony nr 4. Po kolejnym odniesieniu — ponownie do strony nr 4 — poza bieżącym oknem znajduje się ostatnie odniesienie do strony nr 5, więc strona ta jest usuwana ze zbioru roboczego i tym samym z pamięci operacyjnej. Po ostatnim odniesieniu do strony numer 4 zbiór roboczy redukuje się do 3 stron.

- Dla każdej ramki utrzymywany jest wirtualny (mierzony odniesieniami do pamięci) czas ostatniego odniesienia do niej.
- Po każdym odniesieniu do strony zwiększana jest wartość licznika odniesień i wpisywana do odpowiedniej tablicy na pozycji odpowiadającej ramce, w której znajduje się adresowana strona.
- Do zbioru roboczego należą te strony, dla których różnica pomiędzy bieżącą wartością licznika odniesień, a wartością wpisaną w tablicy jest mniejsza lub równa rozmiarowi okna zbioru roboczego.

Precyzyjna identyfikacja zbioru roboczego mogłaby być zaimplementowana podobnie jak algorytm LRU. Zaprezentowana koncepcja jest adaptacją licznika odniesień do stron. Problemem podobnie, jak w przypadku LRU, jest monitorowanie odniesień do stron oraz złożoność struktur danych na potrzeby pamiętania licznika. Ze względu na koszt takiej realizacji algorytm WS implementowany jest w sposób przybliżony.

- Wspomaganie sprzętowe — bit odniesienia.
- Okresowe (wyznaczone przez czasomierz lub przez wystąpienie błędu strony) zwiększanie licznika reprezentującego upływ czasu wirtualnego oraz sprawdzanie bitu odniesienia dla każdej z ramek.
- Jeśli bit odniesienia jest ustawiony, to skasowanie bitu odniesienia i wpisanie bieżącej wartości licznika do odpowiedniej tablicy na pozycji odpowiadającej ramce.
- Jeśli bit odniesienia jest skasowany, to sprawdzenie różnicy pomiędzy bieżącą wartością licznika a wartością wpisaną na odpowiedniej pozycji w tablicy i usunięcie strony, gdy różnica jest większa niż rozmiar okna.

Czas wirtualny i tym samym czas odniesienia do strony wyznaczany jest w sposób przybliżony — z dokładnością do okresu kontrolnego, wyznaczanego przez czasomierz lub błąd strony. Wystąpienie odniesienia do strony sprawdzane jest na podstawie bitu odniesienia w tablicy stron. Bieżąca wartość licznika staje się wirtualnym czasem odniesienia dla tych stron, dla których bit odniesienia jest ustawiony. Brak odniesienia do strony może oznaczać, że ostatnie odniesienie jest już poza oknem i stronę należy usunąć. W tym celu sprawdzana jest różnica pomiędzy bieżącą wartością licznika (bieżącym czasem wirtualnym), a czasem ostatniego odniesienia, odczytanym dla danej strony/ramki z przeznaczonej na to tablicy.

Upływ czasu wirtualnego musi być rejestrowany oddzielnie dla każdego procesu i tylko wówczas, gdy jest on wykonywany. Przerwanie zegarowe powinno więc zwiększać licznik w tym procesie, w kontekście którego jest obsługiwane.

Pomimo uniknięcia konieczności monitorowania wszystkich odniesień do pamięci, implementacja tego algorytmu i tak jest kosztowana, gdyż po każdym przerwaniu zegarowym wymagane jest testowanie bitu odniesienia **każdej** ze stron procesu.

- Dla każdej ramki (strony w ramce) utrzymywany jest wirtualny czas (przybliżony) ostatniego odniesienia.
- Wszystkie ramki pamięci (niezależnie od przynależności do procesu) powiązane są w cykl.
- W wyniku **wystąpienia błędu strony** sprawdzany jest bit odniesienia do strony wskazywanej jako kolejna do usunięcia.
- Jeśli bit odniesienia jest ustawiony, zostaje on skasowany, po czym następuje wskazanie następnej ramki w cyklu i sprawdzenie bitu odniesienia.

Jak sama nazwa wskazuje, algorytm łączy podejście oparte na zbiorze roboczym z koncepcją algorytmu zegarowego. Podobnie jak w przybliżonej realizacji algorytmu zegarowego, utrzymywany jest wirtualny czas ostatniego odniesienia do strony, mierzony osobno dla każdego procesu.

Warto zwrócić uwagę na dwie istotne cechy algorytmu:

- lista cykliczna (tarcza zegara) obejmuje wszystkie ramki w systemie, a nie ramki procesu, jak w przypadku klasycznego algorytmu zegarowego,
- próba usunięcia strony podejmowana jest w reakcji na błąd strony, co oznacza, że jest to właściwie algorytm wymiany na żądanie.

Zasadniczą różnicą w stosunku do algorytmów wymiany na żądanie jest zastępowanie globalne. Podobieństwo do koncepcji zbioru roboczego przejawia się natomiast w usuwaniu pierwszej napotkanej strony, nie należącej do zbioru roboczego swojego procesu.

- Jeśli bit odniesienia jest skasowany, sprawdzana jest różnica pomiędzy wirtualnym czasem bieżącym, a czasem ostatniego odniesienia do wskazywanej strony:
 - jeśli różnica jest większa od rozmiaru okna zbioru roboczego, następuje wymiana strony w ramce,
 - w przeciwnym razie strona pozostaje i następuje wskazanie i sprawdzenie ramki następnej.
- W przypadku wykonania pełnego obiegu przez wskazówkę i stwierdzenia braku stron do wymiany następuje zawieszenie jakiegoś procesu.

Postępowanie w drugiej części algorytmu jest podobne, jak w przypadku właściwego przybliżenia algorytmu WS, czyli strona, do której nie było odniesienia, weryfikowana jest pod kątem przynależności do zbioru roboczego. Wyszukiwanie kończy się jednak na znalezieniu pierwszej ramki ofiary. Brak takiej ramki (brak strony ofiary) zgodnie z ideą zbioru roboczego oznacza, że suma rozmiarów zbiorów roboczych wyczerpuje limit dostępnych ramek pamięci fizycznej i należy zwiesić jakiś proces.



W przykładzie wskazana do usunięcia jest strona nr 5, więc od ramki tej strony rozpocznie się poszukiwanie ofiary. Strona 5 ma ustawiony bit odniesienia więc z założenia jest w zbiorze roboczym i nie zostanie usunięta z pamięci. Wyzerowany zostanie jedynie bit odniesienia i ustawiony wirtualny czas odniesienia do strony na wartość aktualną dla bieżącego procesu, czyli 51. Dla uproszczenia przyjmijmy, że strony 2, 4 i 5 należą do tego samego procesu.

Kolejną kandydatką na liście jest strona nr 4. Bit odniesienia jest dla tej strony skasowany, ale ostatni czas użycia nie jest jeszcze poza oknem ($45 > 51 - 10$), więc strona pozostaje w pamięci, a czas jej użycia się nie zmienia. Następna na liście jest strona nr 2, której ostatnie użycie, zgodnie z wirtualnym czasem, jest poza oknem zbioru roboczego. Strona nr 2 we wskazanej ramce zostanie zastąpiona.

- Wszystkie strony pamięci (niezależnie procesu) powiązane są w listę cykliczną.
- Lista przeglądana jest okresowo przez 2 wskazówki:
 - wiodąca (przednia) zeruje bit odniesienia,
 - zamykająca (tylna) wskazuje stronę do usunięcia.
- Jeśli bit odniesienia przed nadejściem tylnej wskazówki zostanie ponownie ustawiony, strona pozostaje w pamięci, w przeciwnym przypadku jest usuwana.
- Algorytm sterowany jest następującymi parametrami:
 - tempo przeglądania,
 - rozstaw wskazówek.

Algorytm dwuwskazówkowy stosowany jest w niektórych systemach rodziny UNIX, np. SVR4. Realizuje on strategię, określaną jako NRU (Not Recently Used), która w tym przypadku oznacza, że jeśli w czasie pomiędzy przejściem wskazówki wiodącej a nadejściem zamykającej nie został ustawiony bit odniesienia, to strona jest usuwana. Czas pomiędzy przejściem tych wskazówek zależy od tempa przeglądania oraz rozstawu wskazówek. Parametry te są natomiast ustawiane zależnie od liczby wolnych ramek w systemie.

Jeśli uruchomienie algorytmu nie doprowadzi do uzyskania wystarczająco dużej liczby wolnych ramek, następuje zawieszenie jakiegoś procesu i zwolnienie jego ramek.



Wskazówka wiodąca w przykładzie ustawiana jest na stronie nr 5, której bit jest zerowany. Wskazówka zamykająca ustawiona jest na stronie nr 3, dla której bit odniesienia nie został ustawiony, więc nastąpi jej usunięcie. Warto zwrócić uwagę, że bit odniesienia został ustawiony dla strony nr 7, więc po przejściu wskazówki zamykającej strona ta pozostanie w pamięci. W międzyczasie może zostać ustawiony bit odniesienia do strony nr 8. Jeśli jednak nie nastąpi to przed nadejściem zamykającej wskazówki, strona ta zostanie usunięta.

- Jeśli częstotliwość błędów strony generowanych przez proces przekroczy ustalony poziom f_H , to dla stron tego procesu sprowadzanych do pamięci przydzielana są nowe ramki.
- Jeśli częstotliwość błędów strony spadnie poniżej ustalonego poziomu f_L , to
 - zwalniana jest jedna ramka procesu lub
 - zwalniane są wszystkie ramki ze stronami, do których nie było odniesienia od chwili wystąpienia ostatniego błędu strony w danym procesie.
- W szczególności $f_H = f_L$

Algorytmy bazujące na zbiorze roboczym są trudne w implementacji, gdyż ze względu na brak możliwości monitorowania odniesień do stron muszą być realizowane w sposób przybliżony.

Podobne efekty przy mniejszym koszcie implementacji można uzyskać stosując podejście oparte na kontroli *częstości błędów strony*. Podejście jest łatwiejsze w implementacji, gdyż wymaga aktualizacji struktur danych **tylko w przypadku wystąpienia błędu strony**, którego obsługa jest i tak dość czasochłonna.

Ogólna idea polega na tym, żeby zabierać ramki procesom zgłaszającym mało błędów strony, a przydzielać procesom, które często generują błędy strony. W przypadku przekroczenia dolnego progu można rozważać dwie strategie postępowania: żwawą i opieszłą. W strategii żwawej zwalniane są wszystkie ramki, dla których bit odniesienie jest skasowany, a w strategii opieszalej zwalniana jest tylko jedna ramka.

Implementacja algorytmu może być oparta na :

- liczeniu błędów strony, czyli rzeczywistym wyznaczaniu częstotliwości błędów,
- mierzeniu okresu pomiędzy błędami strony.

- Przy każdym błędzie strony generowanym przez proces zwiększany jest licznik błędów strony danego procesu oraz zerowane są bity odniesienia do jego stron.
- W określonych interwałach czasu, wyznaczanych przez czasomierz, sprawdzane są (a następnie zerowane) liczniki poszczególnych procesów.
- Jeśli wartość licznika jest większa od górnej granicy, to procesowi przydzielana jest dodatkowa ramka.
- Jeśli wartość licznika jest mniejsza od ustalonej dolnej granicy, to zwalniana jest:
 - jedna ramka lub
 - wszystkie ramki z wykasowanym bitem odniesienia.

Kontrola częstości błędów strony wymaga wykonania dodatkowego zadania okresowego, wyznaczanego przez czasomierz, niezależnego od błędu strony. Umożliwia to oparcie decyzji o przydziale lub odebraniu ramek na średniej z dłuższego okresu, w przeciwieństwie do implementacji opartej na sprawdzaniu okresu między kolejnymi błędami strony.

- Przy każdym błędzie strony sprawdzany jest czas jaki upłynął od poprzedniego błędu strony.
- Jeśli czas jest mniejszy od ustalonej wielkości T_{min} , to na potrzeby sprowadzanej strony przydzielana jest dodatkowa ramka.
- Jeśli czas jest większy od ustalonej wielkości T_{max} , to
 - zwalniana jest jedna ramka lub
 - zwalniane są wszystkie ramki z wykasowanym bitem odniesienia.
- Bit odniesienia dla stron pozostających w pamięci jest kasowany.

Sprawdzanie okresu między kolejnymi błędami strony nie wymaga dodatkowej kontroli poza obsługą błędu strony, gdyż czas odmierzany jest przez jądro na podstawie taktów zegara, niezależnie od zastosowanego podejścia do wymiany stron. Problemem może być jednak pewna destabilizacja przydziału ramek, wynikająca ze zbyt szybkiej reakcji na przypadkowe błędy strony występujące w krótkim czasie. W teorii automatycznej regulacji można by to podsumować jako dominację członu różniczkującego (duży czas wyprzedzenia).

- Przy każdym błędzie strony zwiększany jest licznik błędów strony danego procesu — n .
- Co pewien okres czasu T_t , wynikający ze stanu wymiany, następuje kontrola odniesień do stron.
- W ramach kontroli wykonywane są następujące czynności:
 - strony, do których nie było odniesienia są usuwane (ich ramki są zwalniane),
 - bity odniesienia stron pozostających na następny okres (do których było odniesienie) są kasowane,
 - wartość licznika n jest zerowana.

W zależności od implementacji, algorytm PFF albo słabo (wolno) dostosowuje się do zmieniającej się strefy w procesie, albo reaguje zbyt gwałtownie. Wolne dostosowanie polega na tym, że po wejściu w nową strefę zbyt długo utrzymywane są strony ze starej strefy. Może to powodować chwilowy wzrost zapotrzebowania na ramki z wszystkimi negatywnymi skutkami, np. zawieszaniem procesów. Zbyt gwałtowna reakcja może spowodować usunięcie stron potrzebnych do dalszego przetwarzania.

Algorytm VSWS (*zbioru roboczego ze zmiennym okresem próbkowania*) przez odpowiednią parametryzację umożliwia lepsze dostosowanie do tego typu sytuacji. Jego istotą jest wcześniejsze lub późniejsze reagowanie na zmieniającą się strefę, zależnie od liczby błędów strony. Czas reakcji jest jednak obustronnie ograniczony, więc reakcja nie będzie ani zbyt szybka (gwałtowna), ani znacząco spóźniona.

- Algorytm sterowany jest następującymi parametrami:
 - T_{max} — maksymalna wielkość interwału czasu,
 - T_{min} — minimalna wielkość interwału czasu,
 - n_{max} — maksymalna liczba błędów strony.
- Kontrola następuje w chwili t , po czasie T_t od momentu poprzedniej kontroli, ustalany następująco:

$$n_t < n_{max} \Rightarrow T_t = T_{max}$$

$$n_t \geq n_{max} \Rightarrow T_{min} \leq T_t \leq T_{max}$$

Interwał czasu pomiędzy kolejnymi kontrolami należy do przedziału obustronnie domkniętego od T_{min} do T_{max} , przy czym konkretna wartość zależy od momentu osiągnięcia liczby n_{max} błędów strony. Jeśli liczba błędów strony jest duża i osiągnie wartość n_{max} przed upływem czasu T_{min} , to kontrola następuje po czasie T_{min} . Jeśli upłynął czas T_{min} , to kontrola nastąpi zaraz po osiągnięciu n_{max} błędów strony, nie później jednak niż po czasie T_{max} . A więc najpóźniej kontrola wystąpi po czasie T_{max} , niezależnie do liczby błędów strony.

- Wymiana ramek zamiast wymiany stron — przyspieszenie wymiany brudnej strony
- Buforowanie usuniętych stron zamiast usuwania — mniej kosztowna obsługa błędu strony
- Czyszczenie stron — ułatwienie ewentualnej wymiany
- Utrzymywanie listy brudnych stron do usunięcia — opóźnienie zapisu stron w obszar wymiany

Przedstawione techniki służą do zmniejszenia czasu obsługi błędu strony. Są one szczególnie przydatne w algorytmach ze sprowadzaniem na żądanie.

Jeśli na potrzeby wymiany została znaleziona ramka, zawierająca zmodyfikowaną stronę, a w systemie dostępne są inne wolne ramki, to zamiast wymiany stron, wymagającej czasochłonnego zapisu, dla sprowadzanej strony przydzielana jest nowa ramka. Ramka z brudną stroną dołączana jest natomiast do puli ramek zwolnionych. W ten sposób następuje wymiana ramek zamiast wymiany stron.

Ramki zwolnione w przypadku stosowania algorytmów ze sprowadzaniem na żądanie nie zawsze są natychmiast wykorzystywane na potrzeby innych procesów. Można więc pozostawić zawartość strony bez zmian do czasu przydziału tej ramki. Jeśli jednak przed ponownym przydziałem wystąpi żądanie dostępu do usuniętej strony, to błąd strony można obsłużyć bez wykonywania czasochłonnej operacji wejścia-wyjścia, przydzielając ramkę ze stroną ponownie.

Czasami system jest niedociążony przez bieżące przetwarzanie. Taki chwilowy okres przestoju można wykorzystać do zapisu zmodyfikowanych stron. Kasowany jest wówczas bit modyfikacji i stronę łatwiej jest wymienić.

W przypadku zwalniania ramek stron modyfikowanych operację zapisu w obszarze wymiany można odłożyć w czasie, dołączając ramkę do specjalnej listy. W miarę dostępności zasobów system stopniowo będzie zapisywał strony w ramkach z tej listy w obszarze wymiany i dołączał ramki do puli oczyszczonych wolnych ramek, gotowych w każdej chwili do użycia

- DPMIN — sprowadzane są wszystkie strony potrzebne w najbliższej przyszłości, które mieszczą się dostępnych ramach
- OBL — sprowadzana jest strona żądana oraz strona następna (wg. numeracji w tablicy stron)
- SL — na podstawie wcześniejszych odniesień i błędów strony budowana jest tablica skojarzeń stron sprowadzanych po sobie i informacja taka wykorzystywana jest przy następnym sprowadzaniu
- FDPA — podejście uwzględniające sugestie programisty co do stron potrzebnych w przyszłości

Sens wstępnego sprowadzania wynika z analizy efektywności wymiany, przedstawionej wcześniej. W praktyce, skuteczność tego podejścia zależy od trafności decyzji o wstępnym sprowadzeniu. Na koszt wymiany wpływa również trafność decyzji o usuwaniu stron, ale sprowadzanie wstępne nie jest obligatoryjne, a więc daje potencjalny zysk, jednak obciążony ryzykiem. Przed omówieniem algorytmów warto przypomnieć, że sprowadzanie (również wstępne) przeprowadza się w wyniku wystąpienia błędu strony.

Algorytm DPMIN (ang. Demand Prepaging MIN) — jest optymalny w klasie algorytmów ze sprowadzaniem na żądanie, ale oparty jest (podobnie jak wcześniej przedstawione algorytmy optymalne) na znajomości **przyszłego** ciągu odniesień do stron. W wyniku wystąpienia błędu strony wszystkie dostępne ramki są wypełniane stronami, do których nastąpi odniesienie w najbliższej przyszłości.

Algorytm OBL (ang. One Block Lookahead) wykorzystuje założenie o lokalności przestrzennej odniesień do pamięci i wraz ze stroną, której zaadresowanie spowodowało błąd, sprowadza stronę następną (jeśli nie ma jej jeszcze w pamięci). Skojarzenie stron ma tu charakter statyczny.

Algorytm SL (ang. Spatial Lookahead) również wykorzystuje założenie o lokalności przestrzennej, ale informacja o powiązaniu stron budowana jest w czasie działania algorytmu, ma więc charakter dynamiczny.

Algorytm FDPA (ang. Fixed-space Demand Prepaging) opiera się na wskazówkach programisty lub wnioskach ze statycznej analizy kodu, przekazywanych do systemu poprzez wykonania specjalnych instrukcji. Instrukcje to dotyczą zarówno wskazania stron potrzebnych w przyszłości, jak i zbędnych. Algorytm uwzględnia więc również wskazówki odnośnie usuwania stron.

- System utrzymuje dodatkową tablicę *pred*, gdzie dla każdej strony pamiętana jest strona, sprowadzona po niej do pamięci. Początkowo $pred[i] = i + 1$.
- *u* jest zmienną przechowującą nr strony ostatnio sprowadzanej do pamięci w wyniku odniesienia do niej.
- W chwili wystąpienia błędu strony następuje sprowadzenie żądanej strony — *p*. Jeśli przy poprzedniej wymianie nie dokonano wstępnego sprowadzenia lub nie było odniesienia do strony wstępnie sprowadzonej do chwili obecnej $pred[u] := p$.
- Jeśli strona $q = pred[p]$ nie znajduje się w pamięci, to następuje również jej sprowadzenia.
- $u := p$

Krótką charakterystyka algorytmów DPMIN i OBL przy opisie wcześniejszego slajdu jest wystarczająca, natomiast omówienia wymaga algorytm SL.

Celem tablicy *pred* jest przechowywanie informacja o numerze strony, którą należy sprowadzić wstępnie w przypadku błędu braku strony. Jeśli więc wystąpi błąd strony przy odniesieniu do strony *p*, to wstępnie sprowadzona zostanie również strona, której numer jest na pozycji *p* w tablicy *pred* — $pred[p]$. Początkowo przyjmujemy, że będzie to strona następna, tak, jak w przypadku OBL.

Zmienna *u* przechowuje numer strony, do której odniesienie spowodowało ostatnio obsługony błąd strony. Jeśli zatem następny błąd strony występuje przy adresowaniu strony *p* (*p* jest kolejną sprowadzaną stroną po stronie *u*), to być może jest jakiś związek pomiędzy tymi stronami i w przyszłości również warto po wystąpieniu błędu strony przy dostępie do strony *u* wstępnie sprowadzić stronę *p*. Numer strony — *p* — powinien zostać zapisany w tablicy *pred* na pozycji odpowiadającej stronie *u*. Należy się jednak upewnić, że dotychczasowe skojarzenie dla strony *u* jest nieodpowiednie. Jeśli zatem przy sprowadzaniu strony *u* nie było wstępnego sprowadzenia innej strony lub strona wstępnie sprowadzona nie została użyta, skojarzenie dla *u* lepiej zmienić.

Na końcu następuje wstępne sprowadzenie strony, wynikające z bieżącego skojarzenia dla *p*. To, czy skojarzenie jest właściwe, okaże się przy wystąpieniu następnego błędu strony. W tym celu wartość *p* podstawiana jest do zmiennej *u*.

- W programie występują dodatkowe instrukcje dla systemu operacyjnego:
 - $\text{FREE}(p)$ — strona p nie będzie wykorzystywana w niedalekiej przyszłości,
 - $\text{PRE}(p)$ — strona p będzie wkrótce wykorzystywana.
- Wykonanie takiej instrukcji skutkuje ustawieniem odpowiednich bitów dla strony:
 - bitu usunięcia — D ,
 - bitu wstępnego sprowadzenia — P .

W wyniku analizy kodu w odpowiednich miejscach w programie pojawiają się instrukcje dla systemu operacyjnego, informujące o przyszłych potrzebach (lub ich braku) w odniesieniu do pewnych stron. Zgodnie z tymi instrukcjami w tablicy stron lub innej strukturze ustawiane są odpowiednie bity. W toku realizacji wymiany bity te są kasowane.

Bit P oznacza, że strona będzie w najbliższej przyszłości potrzebna, a jego wykasowanie nastąpi przy odwołaniu do strony po jej sprowadzaniu do pamięci.

Bit D oznacza, że strona nie będzie na razie potrzebna. Gdyby jednak strona została usunięta z pamięci i ponownie sprowadzana w wyniku błędu strony, bit ten zostanie skasowany.

- W przypadku wystąpienia błędu strony sprowadzana jest strona żądana, ewentualnie wymieniana z jakąś stroną z ustawionym bitem D . Jeśli nie ma takiej strony stosowany jest inny algorytm usuwania np. LRU, z pominięciem stron, dla których ustawiony jest bit P .
- Jeśli w systemie są wolne ramki lub mogą zostać zwolnione przez usunięcie stron z ustawionym bitem D , to sprowadza się tyle stron z ustawionym bitem P , na ile pozwala potencjalna liczba wolnych ramek.
- W przypadku odniesienia do strony z ustawionym bitem P , bit ten jest kasowany.

System sprowadza oczywiście strony na żądanie i dokonuje wymiany w przypadku braku wolnych ramek. Usuwane są strony z ustawionym bitem D , a gdy takich nie ma, zgodnie z innym algorytmem, np. LRU lub w praktyce z którymś z jego przybliżeń. Algorytm usuwania nie jest tu najbardziej istotny poza preferencją dla stron „wskazanych” przez bit D .

Strona z ustawionym bitem P może być zarówno w pamięci, jak i poza nią. Będąc poza pamięcią, strona taka czeka na sprowadzenie. Bit P ustawiony dla strony przebywającej w pamięci operacyjnej oznacza, że stronę tę wstępnie sprowadzono, ale nie było do niej jeszcze odniesienia, dlatego nie powinna być usunięta przy wymianie.

System stara się sprowadzać do pamięci wszystkie strony z ustawionym bitem P w miarę dostępnych ramek. W bilansie ramek uwzględnia się zwolnienie ramek przez strony z ustawionym bitem D . Jak już wcześniej wspomniano, gdyby bit D ustawiony był dla strony sprowadzanej, to zostanie wykasowany.

- Segmentacja na żądanie — wymiana segmentów pomiędzy pamięcią pierwszego i drugiego rzędu
- Skomplikowana wymiana — zróżnicowany rozmiar segmentów
- Mniejsza przewidywalność czasu dostępu — zróżnicowany czas przesyłania zawartości segmentu

W niektórych systemach (między innymi IBM OS/2) podjęto próbę realizacji segmentacji na żądanie. Zarządzanie wymianą segmentów jest jednak znacznie bardziej skomplikowane, gdyż sprowadzenie segmentu do pamięci może wymagać upakowania i/lub przesunięcia w obszar wymiany jednego lub kilku segmentów znajdujących się w pamięci fizycznej. Optymalizacja decyzji w takim przypadku jest więc procedurą czasochłonną. W przypadku stronicowania czas sprowadzenia stron do pamięci oraz zapisu na dysku jest zmienną losową, w przypadku segmentacji, ze względu na zróżnicowany rozmiar segmentów, sam czas przesyłania danych i tym samym czas oczekiwania w kolejce do dysku jest trudniej przewidzieć.