



سیستم‌های دیجیتال ۲

اسلاید ۴

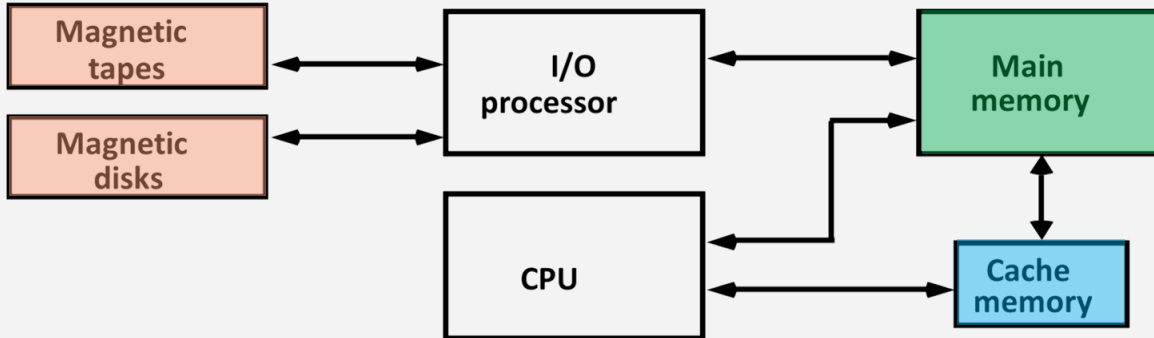
سازمان حافظه

محمد رضا رازقی زاده



سلسله مراتب حافظه

Auxiliary memory



حافظه اصلی (Main Memory):

واحد حافظه‌ای که مستقیماً با CPU ارتباط برقرار می‌کند (RAM).

حافظه کمکی (Auxiliary Memory):

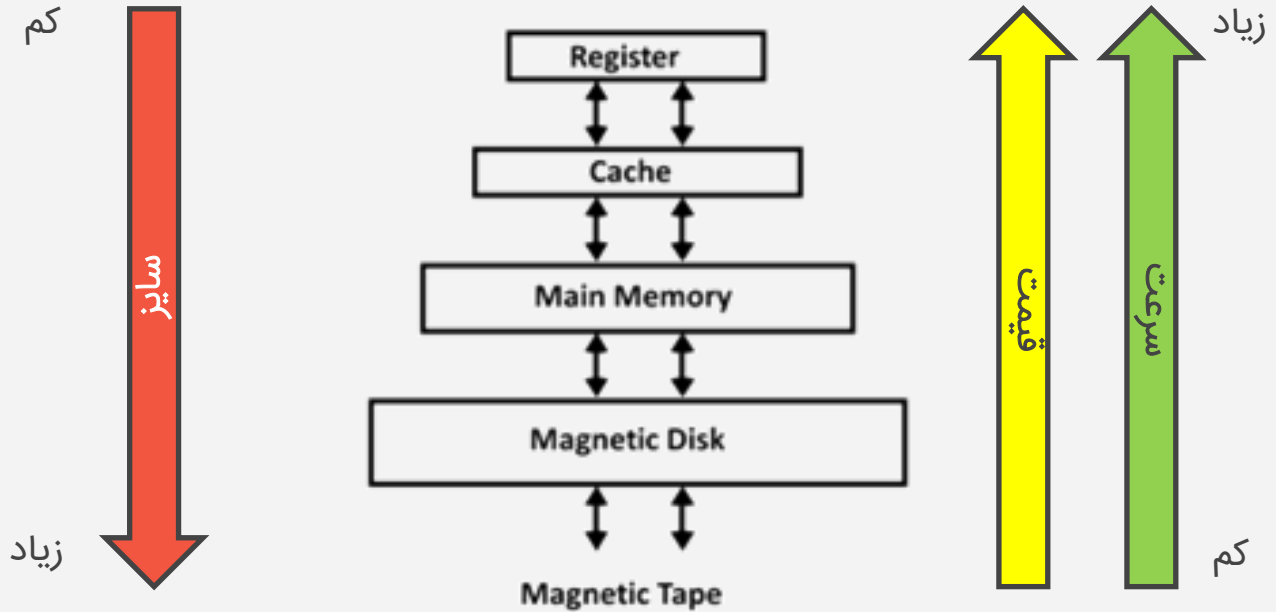
دستگاه‌هایی که ذخیره‌سازی پشتیبان را فراهم می‌کنند (Disk Drives).

حافظه کش (Cache Memory):

حافظه‌ای بسیار پرسرعت که برای افزایش سرعت پردازش به کار می‌رود (Cache RAM).



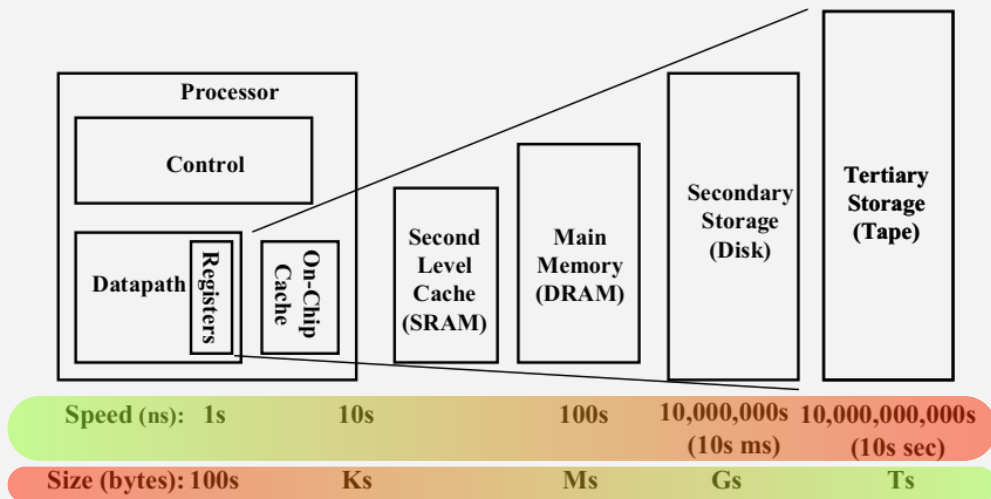
سلسله مراتب حافظه



➤ **سلسله مراتب حافظه** با هدف سرعت بخشیدن دسترسی به حافظه با حداقل هزینه بوجود آمده است.



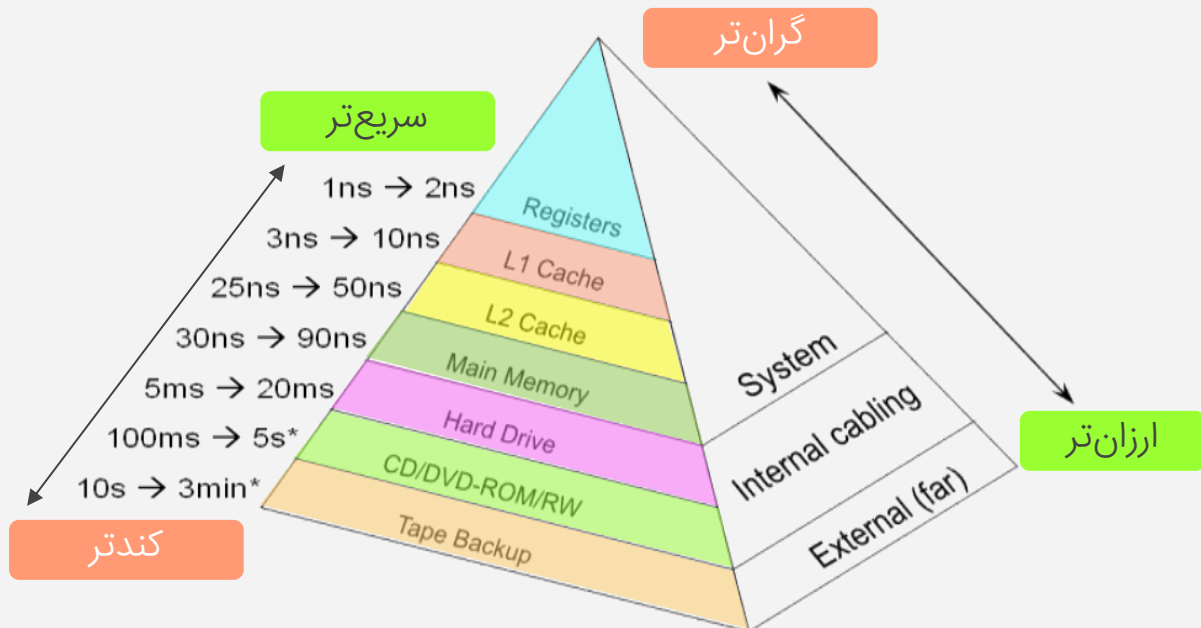
سلسله مراتب حافظه



- با حرکت از پردازنده به سمت حافظه‌های ثالثیه، **سرعت کاهش** و **سایز افزایش** می‌یابد.
- سیستم‌های مدرن از این سلسله‌مراتب برای بهینه‌سازی دسترسی به داده‌ها و کارایی کلی سیستم استفاده می‌کنند.



سلسله مراتب حافظه



• در سلسله مراتب حافظه اگر از پایین (کف هرم) به بالا حرکت کنیم:

✓ سرعت دسترسی به حافظه بیشتر می شود.

✓ هزینه سخت افزار افزایش پیدا می کند.

✓ حجم حافظه کاهش پیدا می کند.



اهمیت حافظه

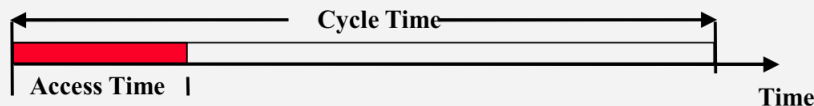
- حافظه از چه نظر اهمیت دارد؟
 - << سرعت دسترسی به حافظه يك گلوگاه اصلی برای افزایش کارایی کامپیوتر است.
- در حالی که در سالیان گذشته حجم حافظه به سرعت افزایش پیدا کرده است، کاهش در زمان دسترسی به محتوای حافظه پیشرفت بسیار کمتری داشته است.

Year	DRAM	Cycle Time
	Size	
1980	64 Kb	250 ns
1983	256 Kb	220 ns
1986	1 Mb	190 ns
1989	4 Mb	165 ns
1992	16 Mb	145 ns
1995	64 Mb	120 ns



زمان دسترسی به حافظه

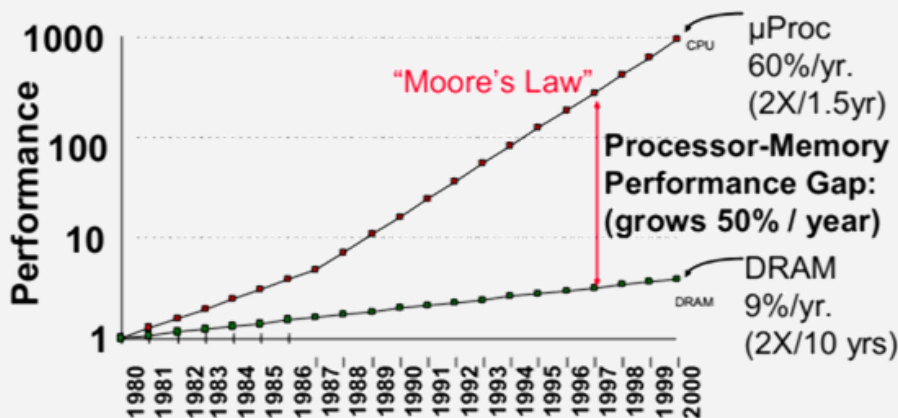
- زمان دسترسی به حافظه (Access Time) عبارت است از زمانی که بین تقاضای داده و آماده شدن آن طول می کشد.
- زمان سیکل (Cycle Time) عبارت است از زمان بین دو تکرار متوالی



- برای مثال در حالی که زمان دسترسی به يك حافظه DRAM ممکن است ۶۰ ns باشد با احتساب زمان های لازم برای انتقال آدرس از پردازنده به حافظه، تاخیر باس ها، و ... زمان سیکل آن ممکن است به ۱۸۰-۲۵۰ ns برسد.



بر اساس قانون مور که گردون مور از بنیانگذاران شرکت اینتل آن را ارایه کرد، تعداد ترانزیستورهای روی يك تراشه (با مساحت ثابت) هر دو سال يك بار تقريبا دو برابر خواهد شد.





➤ ویژگیهای حافظه:

✓ حافظه های بزرگ کند هستند

✓ حافظه های سریع کوچک هستند

➤ سؤال: چگونه می توان حافظه بزرگ سریع و ارزانی داشت؟

✓ استفاده از سلسله مراتب

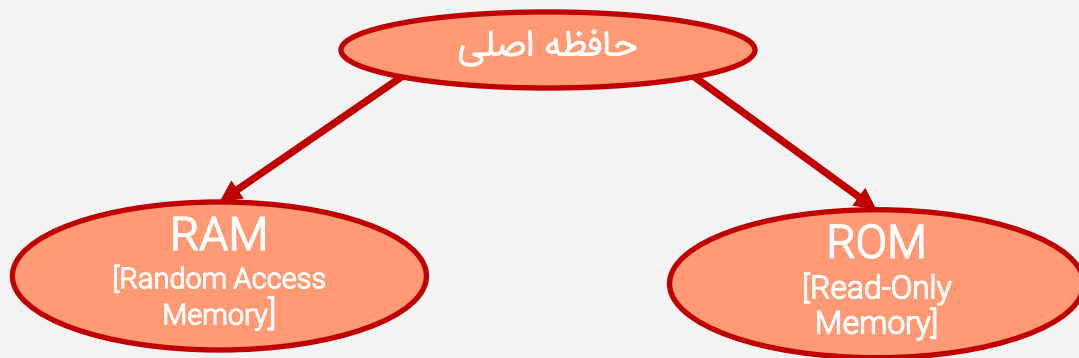
✓ دسترسی موازی





حافظه اصلی

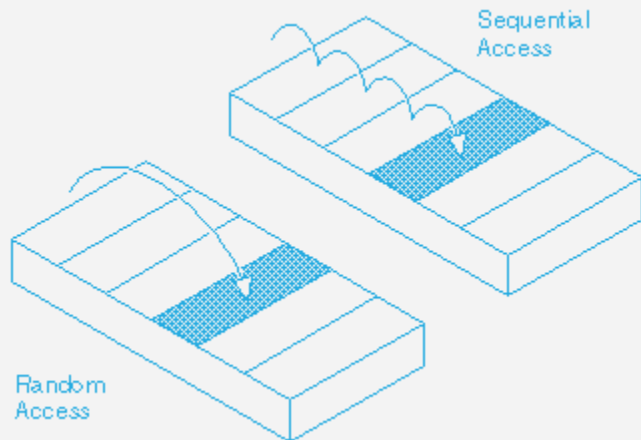
- حافظه اصلی از مدارات سریعی ساخته می شود که برنامه ها و داده های مورد نیاز را در هنگام اجرا نگهداری می نماید. حافظه اصلی بر دو نوع است:



- از نظر دسترسی پذیری تصادفی (Random Access)، هر دو حافظه ROM و RAM مشابه هستند، زیرا در هر دو نوع حافظه می توان به هر موقعیت حافظه به صورت مستقیم و تصادفی دسترسی پیدا کرد. به عبارت دیگر، در هر دو نوع حافظه، دسترسی به هر بخش از داده ها بدون نیاز به خواندن داده های پیشین انجام می شود.

- داده‌های با دسترسی تصادفی

اجازه می‌دهند که به هر بخش از داده‌ها به‌طور مستقیم و مستقل دسترسی پیدا کنید (مانند RAM و HDD/SSD).



- داده‌های با دسترسی ترتیبی

نیازمند دسترسی به داده‌های قبلی برای رسیدن به بخش خاصی از داده‌ها هستند (مانند نوارهای مغناطیسی و نوار کاست).

حافظه اصلی < حافظه ROM

- حافظه ROM حافظه‌ای است فقط خواندنی که محتوی آن يك بار نوشته شده و پس از نصب در کامپیوتر محتوای آن غیر قابل تغییر می‌باشد.
- معمولا از این حافظه برای ذخیره برنامه هایی نظیر bootstrap loader که برای راه اندازی اولیه کامپیوتر مورد نیاز هستند استفاده می شود.



- این حافظه انواع مختلفی دارد:

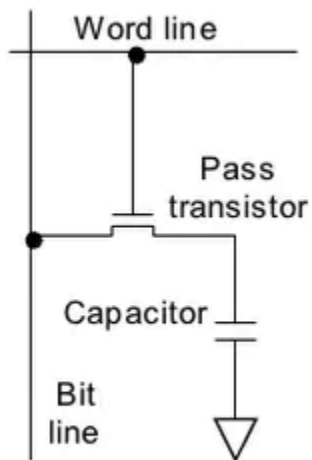




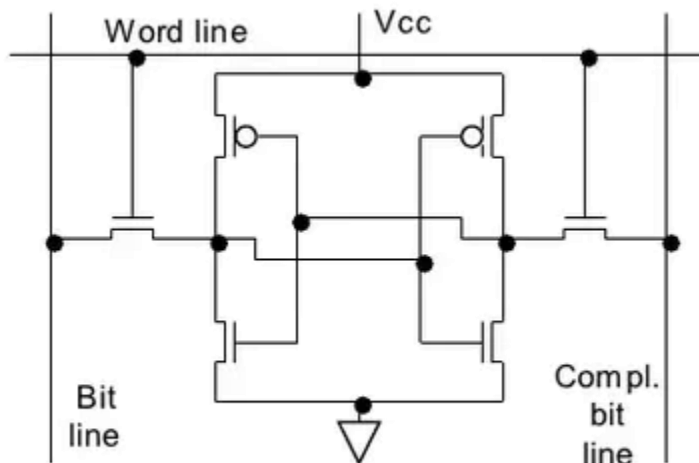
➤ عمده حافظه اصلی کامپیوتر از حافظه RAM ساخته می شود.

➤ معمولا در کامپیوترها دو نوع حافظه RAM مورد استفاده هستند:

- **DRAM: Dynamic Random Access Memory**
 - High density, low power, cheap, slow
 - Dynamic: need to be “refreshed” regularly
- **SRAM: Static Random Access Memory**
 - Low density, high power, expensive, fast
 - Static: content will last “forever” (until lose power)



ساختار حافظه DRAM



ساختار حافظه SRAM



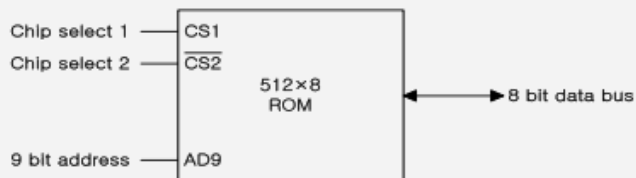
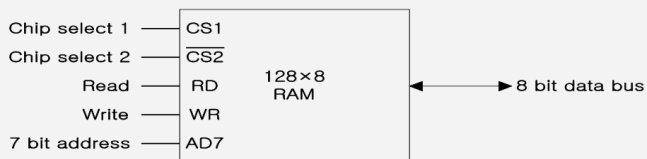
حافظه اصلی < حافظه ROM

ویژگی	SRAM (Static RAM)	DRAM (Dynamic RAM)
مخفف	Static Random Access Memory	Dynamic Random Access Memory
ساختار ذخیره‌سازی	از فلیپ-فلاپ‌ها و ترانزیستور برای ذخیره هر بیت استفاده می‌کند.	از خازن و ترانزیستور برای ذخیره هر بیت استفاده می‌کند.
نیاز به تازه‌سازی	نیاز به تازه‌سازی (Refresh) ندارد.	نیاز به تازه‌سازی مداوم دارد تا داده‌ها حفظ شوند.
سرعت	سریع‌تر	کندتر در مقایسه با SRAM
تراکم حافظه	تراکم کمتر (نیاز به فضای بیشتری برای هر بیت).	تراکم بالاتر (نیاز به فضای کمتر برای هر بیت)
مصرف توان	مصرف توان کمتر در حالت ثابت، اما در حالت فعال بیشتر است.	مصرف توان کمتر در حالت فعال، اما به دلیل تازه‌سازی مداوم، مصرف توان کلی بیشتر است.
پیچیدگی طراحی	طراحی پیچیده‌تر به دلیل تعداد بیشتر ترانزیستورها.	طراحی ساده‌تر به دلیل استفاده از خازن‌ها.
ظرفیت حافظه	معمولاً ظرفیت‌های کوچک‌تر (چند مگابایت)	معمولاً ظرفیت‌های بزرگ‌تر (چند گیگابایت)
هزینه	گران‌تر به دلیل پیچیدگی و تعداد ترانزیستورها.	ارزان‌تر به دلیل تراکم بالا و طراحی ساده‌تر.
کاربردها	کش CPU (L1، L2، L3) - حافظه‌های با سرعت بالا	-حافظه اصلی (RAM) در کامپیوترها و سرورها - کارت‌های گرافیکی
زمان دسترسی	کوتاه‌تر (چند نانوثانیه)	بلندتر (چند ده نانوثانیه)
پایداری داده	تا زمانی که تغذیه برقرار باشد، پایدار است.	نیاز به شارژ مداوم خازن‌ها برای حفظ داده‌ها.



بلوك دياگرام حافظه

- ✓ Typical RAM chip
 - > 128 X 8 RAM : 27 = 128 (7 bit address lines)
- ✓ Typical ROM chip
 - > 512 X 8 ROM : 29 = 512 (9 bit address lines)



CS ۱	CS ۲	R D	W R	عمل حافظه	وضعیت گذرگاه داده
۰	۰	X	X	مسدود	امپدانس بالا
۰	۱	X	X	مسدود	امپدانس بالا
۰	۰	۰	۱	مسدود	امپدانس بالا
۱	۰	۰	۱	نوشتن	ورودی داده به RAM
۱	۰	۱	۰	خواندن	خروجی داده به RAM
۱	۱	X	X	مسدود	امپدانس بالا



Memory Configuration : **512** bytes **RAM** + **512** bytes **ROM**

[illegible]



ارتباط حافظه با پردازنده

Memory Configuration : **512 bytes RAM** + **512 bytes ROM**

» **4 x 128 bytes RAM** + **1 x 512 byte ROM**

Memory Address Map

» Address line **9 8**

- RAM 1 0 0 : 0000 - 007F
- RAM 2 0 1 : 0080 - 00FF
- RAM 3 1 0 : 0100 - 017F
- RAM 4 1 1 : 0180 - 01FF

» Address line **10**

- ROM 1 : 0200 - 03FF

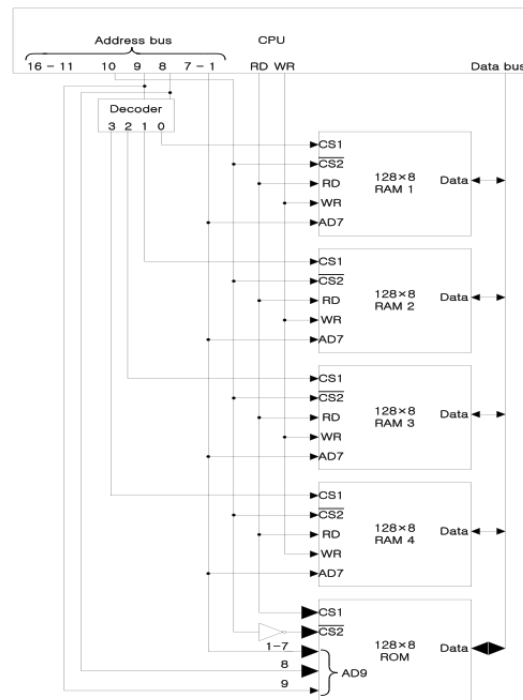
Memory Connection to CPU

» **2 x 4 Decoder** : RAM select (**CS1**)

» Address line **10**

- RAM select : **CS2**
- ROM select : **CS2 Invert**

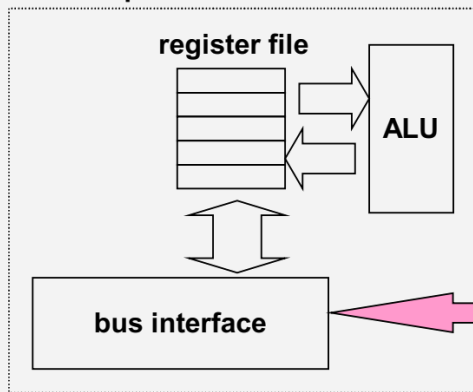
Component	Address	$A_{10} A_9 A_8 A_7 A_6 A_5 A_4 A_3 A_2 A_1$
RAM1	0000-007F	0 0 0 x x x x x x x
RAM2	0080-00FF	0 0 1 x x x x x x x
RAM3	0100-017F	0 1 0 x x x x x x x
RAM4	0180-01FF	0 1 1 x x x x x x x
ROM	0200-03FF	1 x x x x x x x x x





ارتباط حافظه اصلی با پردازنده

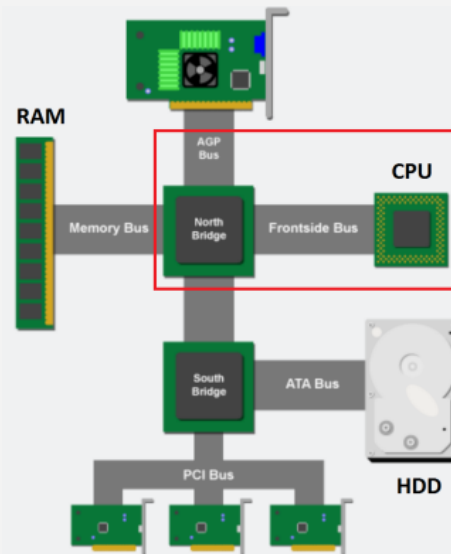
CPU chip



memory bus

Memory Interface

main memory Modules



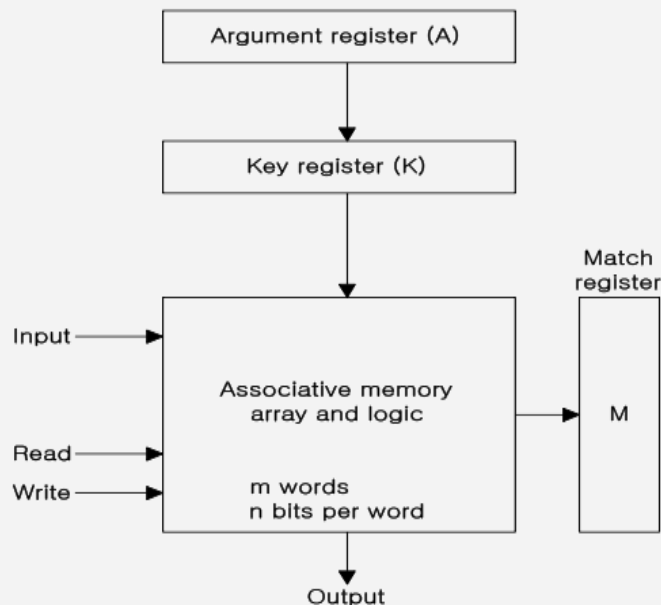


حافظه شرکت پذیر یا هم پیوند (Associative Memory (AM)

- از این حافظه برای تسریع عمل جستجو در بین داده های ذخیره شده استفاده می شود.
- برای کاهش زمان جستجو به جای استفاده از آدرس، از محتوی خود داده استفاده می شود.
- نام دیگر این روش Content Addressable Memory (CAM) می باشد.
- هنگام نوشتن داده در این حافظه نیازی به آدرس نیست!
- این حافظه قادر است تا جای خالی را پیدا نموده و داده را در آنجا ذخیره می نماید.
- هنگام خواندن داده از حافظه، مقدار داده و یا بخشی از آن به حافظه ارائه شده و حافظه تمامی کلمات ذخیره شده ای را که با داده مورد نظر یکسان هستند را علامت گذاری کرده و آماده خواندن می نماید.



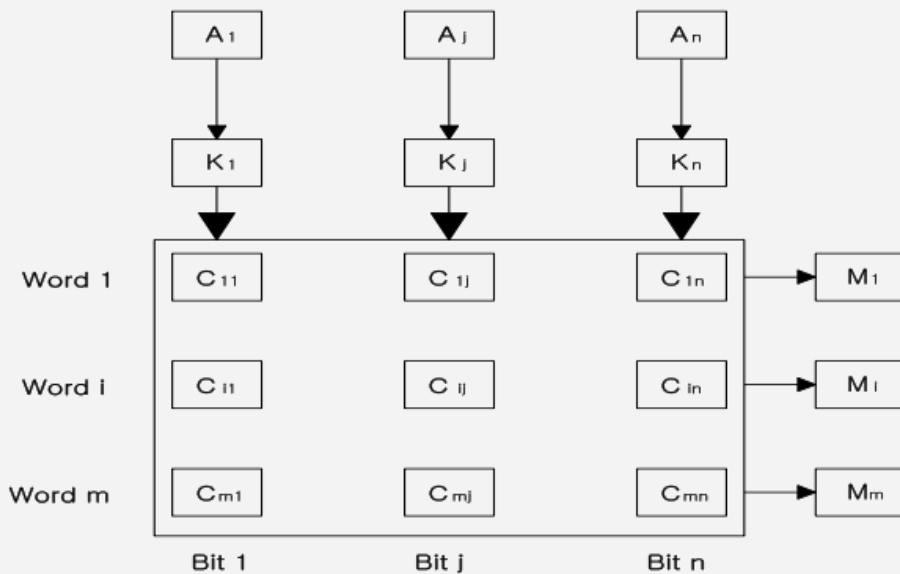
بلوک دیاگرام حافظه شرکت پذیر



- هر کلمه در حافظه به طور موازي با محتویات **ثبات** (Argument Register) مقایسه می شود.
- به کمک **ثبات کلید** (Key Register) می توان محدوده جستجو را کاهش داد.
- اگر در حافظه، $\text{word}[i] = A$ آنگاه $M[i] = 1$. در غیر این صورت $M[i] = 0$.
- تمام کلماتی که بیت متناظر آنها در M یک شده باشد، از حافظه خوانده می شوند.



عملکرد حافظه شرکت پذیر





	A	K	M
Word 0	1	1	0
Word 1	1	0	0
Word 2	1	0	0
Word 3	0	1	0
Word 4	1	1	0
Word 5	1	1	1
Word 6	0	0	0
Word 7	1	1	0



	A	K	M
Word 0	1	1	0
Word 1	1	0	1
Word 2	1	0	0
Word 3	0	1	0
Word 4	1	1	0
Word 5	1	0	1
Word 6	0	0	0
Word 7	1	1	0



حافظه نهان (Cache)

- اگر قسمت فعال برنامه و داده ها در حافظه سریع و کوچکی قرار داده شود، می توان با کم کردن میانگین زمان دسترسی به حافظه، زمان اجرای برنامه را کاهش داد.
- این حافظه سریع و کوچک را حافظه **cache** می نامند.
- معمولا حافظه cache زمان دسترسی بسیار کمتری نسبت به حافظه اصلی دارد (۵ تا ۱۰ برابر کمتر)
- دریک کامپیوتر ممکن است عمل cache تا چند سطح تکرار گردد.

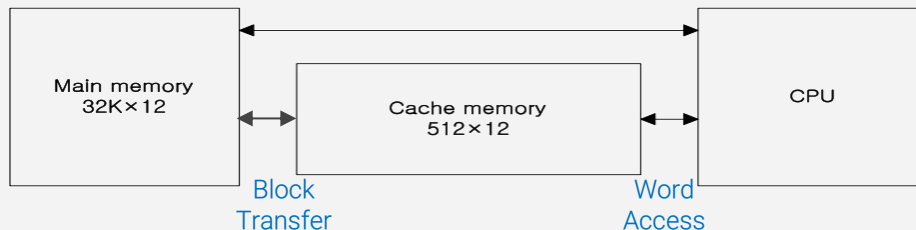


حافظه نهان (Cache)

- وقتي که CPU نیاز به دسترسی به حافظه دارد ابتدا حافظه cache را جستجو می نماید؛ اگر داده در این حافظه سریع موجود بود (hit) از آن استفاده می شود در غیر اینصورت (miss) با رجوع به حافظه، بلوکی از داده ها که شامل داده مورد نیاز CPU می شود از حافظه اصلی به حافظه cache منتقل می گردد.



- ممکن است CPU منتظر نوشته شدن داده و کلمات مجار آن (یک بلوک) در حافظه نهان نماند و داده را مستقیماً از حافظه اصلی دریافت کند.





درصد موفقیت

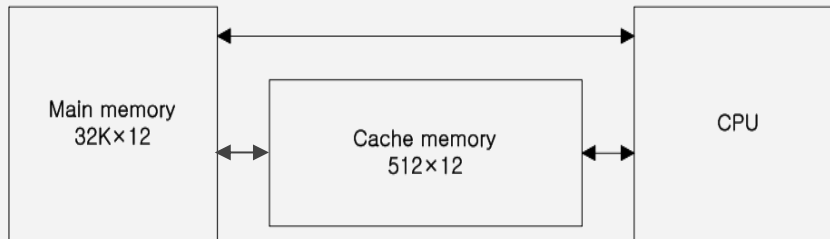
- کارایی حافظه cache با ضریبی با نام **درصد موفقیت (نرخ دسترسی) hit ratio** اندازه‌گیری می‌شود.
- این درصد عبارت است از نسبت تعداد دفعاتی که داده در cache یافت شده به تعداد کل دفعات رجوع به حافظه.
- هر چه مقدار این درصد بیشتر باشد سرعت دسترسی به حافظه به سرعت cache نزدیکتر می‌شود.

$$\text{درصد موفقیت (h)} = \frac{\text{تعداد hit ها}}{\text{تعداد کل مراجعات}} = \frac{\text{تعداد hit ها}}{\text{تعداد hit ها} + \text{تعداد miss ها}}$$



زمان متوسط دسترسي

فرض: حافظه اصلي مي تواند مستقيما داده در اختيار CPU قرار دهد.



t_{eff} : مدت زمان دستیابی متوسط (موثر) cache

t_{cache} : مدت زمان دستیابی cache

t_{main} : مدت زمان دستیابی حافظه اصلي

h : درصد موفقیت

$$t_{eff} = h \cdot t_{cache} + (1-h) \cdot t_{main}$$



زمان متوسط دسترسی

مثال: فرض کنید یک کامپیوتر با مشخصات زیر در اختیار داریم:

- ✓ زمان دسترسی ns ۱۰۰ برای حافظه cache
 - ✓ زمان دسترسی ns ۱۰۰۰ برای حافظه اصلی
 - ✓ درصد موفقیت ۰.۹
 - ✓ حافظه اصلی می تواند مستقیماً به CPU داده بدهد
- زمان متوسط دسترسی را محاسبه نمایید؟

$$t_{\text{cache}} = 100 \text{ ns}$$

$$t_{\text{eff}} = h \cdot t_{\text{cache}} + (1-h) \cdot t_{\text{main}}$$

$$t_{\text{main}} = 1000 \text{ ns}$$

$$t_{\text{eff}} = (0.9) \cdot (100 \text{ ns}) + (0.1) \cdot (1000 \text{ ns}) = 190 \text{ ns}$$

$$h = 0.9$$



زمان متوسط دسترسي

فرض: حافظه اصلي نمي تواند مستقيماً داده در اختيار CPU قرار دهد.



t_{eff} : مدت زمان دستیابی متوسط (موثر) cache

t_{cache} : مدت زمان دستیابی cache

t_{main} : مدت زمان دستیابی حافظه اصلي

h : درصد موفقیت

$$t_{eff} = h \cdot t_{cache} + (1-h) \cdot (t_{main} + t_{cache}) = h \cdot t_{cache} + (1-h) \cdot t_{main}$$



زمان متوسط دسترسی

مثال: فرض کنید يك کامپیوتر با مشخصات زیر در اختیار داریم:

- ✓ زمان دسترسی ns ۱۰۰ برای حافظه cache
 - ✓ زمان دسترسی ns ۱۰۰۰ برای حافظه اصلی
 - ✓ درصد موفقیت ۰.۹
 - ✓ حافظه اصلی می تواند مستقیماً به CPU داده ندهد
- زمان متوسط دسترسی را محاسبه نمایید؟

$$t_{\text{cache}} = 100 \text{ ns}$$

$$t_{\text{eff}} = t_{\text{cache}} + (1-h) \cdot t_{\text{main}}$$

$$t_{\text{main}} = 1000 \text{ ns}$$

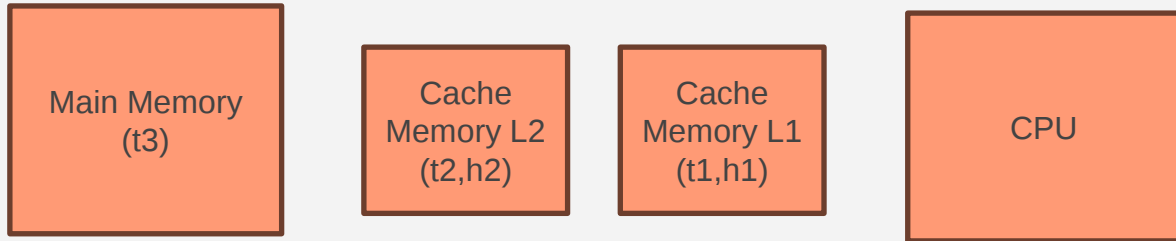
$$t_{\text{eff}} = (100 \text{ ns}) + (0.1) \cdot (1000 \text{ ns}) = 200 \text{ ns}$$

$$h = 0.9$$



زمان متوسط دسترسی

زمان متوسط دسترسی برای یک سیستم با چندین سطح cache به صورت زیر محاسبه می گردد:



اگر CPU مستقیماً به حافظه متصل **نباشد**:

$$t_{\text{eff}} = t_1 + (1-h_1) \cdot t_2 + (1-h_1) \cdot (1-h_2) \cdot t_3$$

اگر CPU مستقیماً به حافظه متصل **باشد**:

$$t_{\text{eff}} = h_1 \cdot t_1 + (1-h_1)h_2 \cdot t_2 + (1-h_1) \cdot (1-h_2) \cdot t_3$$



تأثیر اندازه بلوک های cache

در حافظه کش، اندازه بلوک‌ها نقش مهمی در عملکرد آن دارد. اگر اندازه بلوک‌ها بزرگ‌تر باشد:

بهره‌مندی از محلی‌سازی فضایی Spatial Locality:

داده‌هایی که نزدیک به هم در حافظه اصلی قرار دارند، احتمالاً با هم نیاز خواهند شد. بلوک‌های بزرگ‌تر می‌توانند داده‌های بیشتری را یکجا وارد کش کنند و این باعث کاهش نیاز به دسترسی مکرر به حافظه اصلی می‌شود.

افزایش جریمه خطا Miss Penalty:

زمانی که داده مورد نظر در کش وجود ندارد، نیاز به انتقال یک بلوک کامل از حافظه اصلی به کش است. هرچه بلوک بزرگ‌تر باشد، زمان بیشتری برای این انتقال نیاز خواهد بود و جریمه خطا افزایش می‌یابد.

تأثیر بر زمان و نرخ دسترسی:

بلوک‌های بزرگ‌تر ممکن است زمان دسترسی به کش Hit Time را افزایش دهند و همچنین به دلیل پر شدن سریع‌تر کش، نرخ خطا Miss Rate نیز بالا برود. بنابراین، زمان دسترسی میانگین AMAT با توجه به این فاکتورها محاسبه می‌شود:

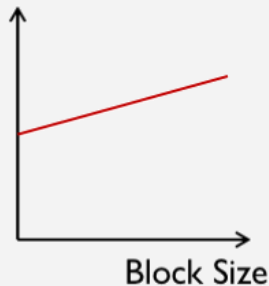
$$AMAT = Hit Time \times Hit Ratio + Miss Penalty \times (1 - Hit Ratio)$$



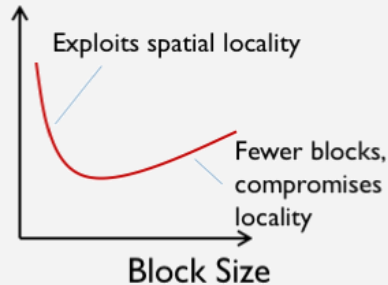
تأثیر اندازه بلوک های cache

به بیان ساده، بلوک های بزرگ تر می توانند عملکرد را بهبود دهند اما اگر بیش از حد بزرگ باشند، ممکن است اثر معکوس بگذارند و کش کندتر عمل کند.

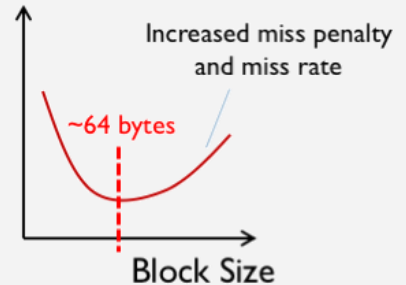
Miss Penalty



Miss Ratio

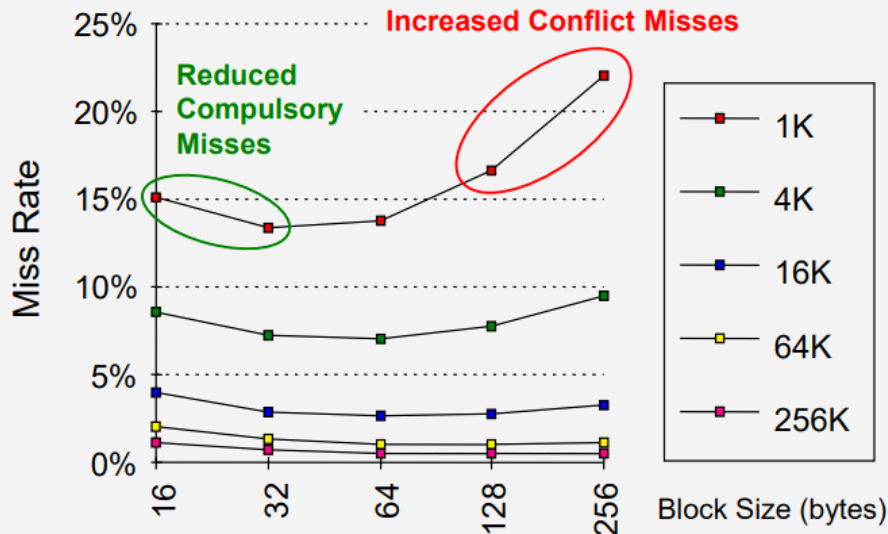


AMAT



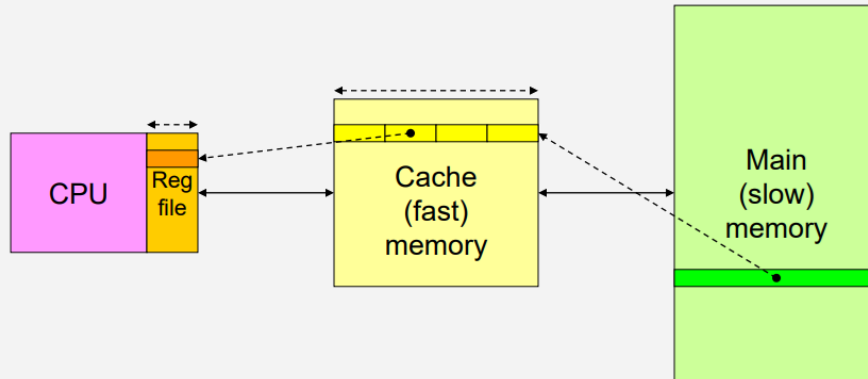


تأثير اندازه بلوك هاي cache



64-byte blocks
are common in
L1 caches

128-byte block
are common in
L2 caches



عمل انتقال داده از حافظه اصلی به حافظه cache نگاشت نامیده می‌شود. این کار به سه طریق قابل انجام است:

نگاشت مستقیم (Direct Mapping)

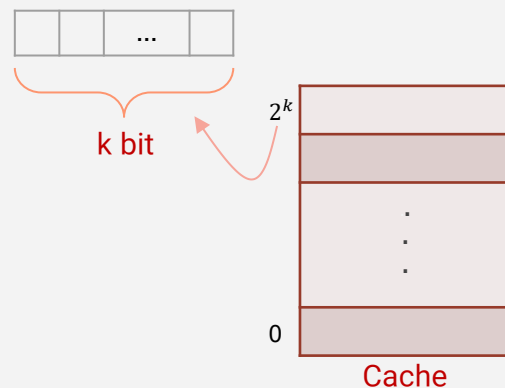
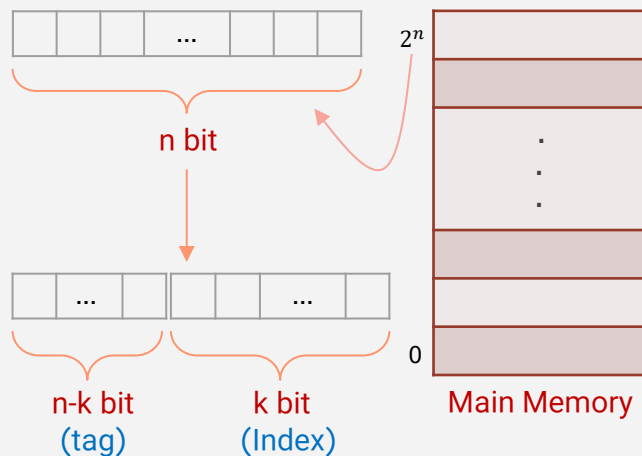
نگاشت شرکت پذیر (Associate Mapping)

نگاشت مجموعه های هم پیوند (Set-associative mapping)



نگاشت مستقیم (Direct Mapping)

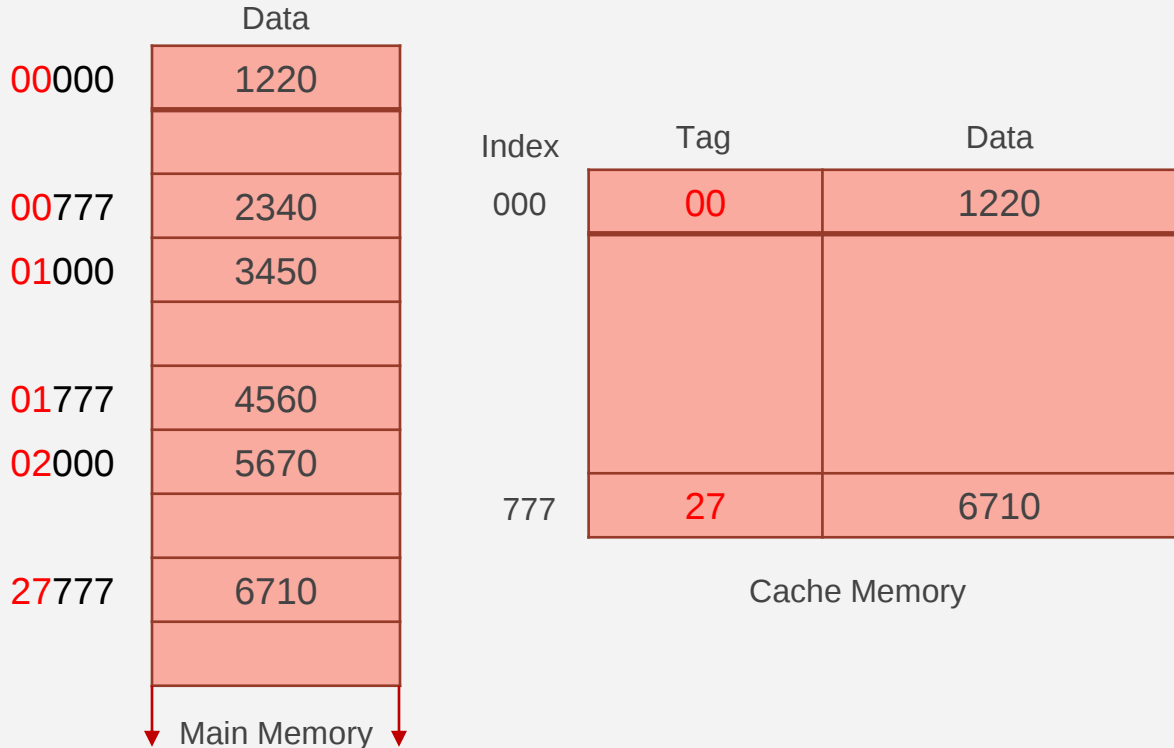
- در این روش برای پیاده سازی حافظه کش از یک حافظه RAM سریع استفاده می‌شود:
- اگر حافظه اصلی دارای 2^n کلمه و حافظه کش دارای 2^k کلمه باشند، در نتیجه به ترتیب برای هر یک به n و k بیت آدرس نیاز خواهد بود.
- در این صورت آدرس n بیتی حافظه اصلی به صورت زیر تقسیم می‌شود:





نگاشت مستقیم (Direct Mapping)

- در روش نگاشت مستقیم، اطلاعات و Tag هر دو در حافظه کش ذخیره می شود.





نگاشت مستقیم (Direct Mapping)



✓ وقتی که آدرسی توسط CPU تولید می شود با استفاده از قسمت Index آن یک کلمه از حافظه cache خوانده شده و مقدار Tag ذخیره شده در آن با مقدار Tag آدرس مقایسه می گردد.

✓ اگر دو Tag یکسان بودند، داده cache مورد استفاده قرار می گیرد (hit) در غیر اینصورت (miss) داده از حافظه خوانده شده و همراه با Tag جدید در cache نوشته می شود.

✓ در این روش اگر رجوع به آدرسهایی با index یکسان زیاد اتفاق بیافتد، درصد موفقیت پائین می آید.



نگاشت مستقیم (Direct Mapping)

✓ معمولا اندازه بلوک بزرگتر از يك در نظر گرفته مي شود:

	Index	Tag	Data
Block 0	000	0 1	3 4 5 0
	007	0 1	6 5 7 8
Block 1	010		
	017		
Block 63	770	0 2	
	777	0 2	6 7 1 0

6	6	3
Tag	Block	Word



نگاشت شرکت پذیر (Associate mapping)

✓ در این روش که سریعترین راه پیاده سازی يك cache است از يك حافظه associative استفاده می شود.

✓ در این حافظه هم آدرس (به طور کامل) و هم محتوای يك کلمه ذخیره می شوند.

✓ هنگام جستجو برای يك داده، آدرس آن به حافظه

associative عرضه شده و در صورتی که ورودی متناظری

در حافظه باشد، داده مربوطه در خروجی ظاهر می گردد. در

غیر اینصورت هر دوی آدرس و داده در حافظه

associative ذخیره خواهند شد.



✓ معمولاً از الگوریتم FIFO

برای جایگزینی در cache

استفاده می شود.



نگاشت مجموعه هاي هم پيوند (Set-associative mapping)

- ✓ در اين روش حافظه کش با ساختاری مشابه روش نگاشت مستقيم يعنی شامل Tag و Data در هر محل از حافظه cache مي تواند پر شود. با اين تفاوت که در هر سطر حافظه کش چندین کلمه با آدرس Index يکسان ذخيره می شود.
- ✓ تعداد data-tag هاي ذخيره شده در حافظه يك set خوانده مي شود.
- ✓ براي مقايسه tag آدرس توليد شده با مقادير ذخيره شده از حافظه associative استفاده مي شود.
- ✓ با بزرگ شدن set ها درصد موفقيت cache افزايش مي يابد.

Index	Tag	Data	Tag	Data
000	0 1	3 4 5 0	0 2	5 6 7 0
777	0 2	6 7 1 0	0 0	2 3 4 0



نگاشت مجموعه هاي هم پيوند (Set-associative mapping)

✓ وقتي كه يك miss اتفاق مي افتد، در صورتيكه يك set پر باشد لازم مي شود تا يكي از داده هاي آن خالي شده و داده جديد وارد cache شود.

✓ روشهاي مختلفي براي جاگزینی داده جديد از حافظه اصلی در حافظه کش:

- first-in, first-out
- Random replacement
- Least recently used (LRU)



نوشتن در حافظه cache

- ✓ یکی از مسایل مهم در ارتباط با حافظه cache نحوه عمل در هنگام نوشتن اطلاعات در حافظه است.
- ✓ هنگام خواندن داده ها وقتی داده در cache وجود داشته باشد، نیازی به رجوع به حافظه اصلی نیست اما وقتی داده را در حافظه می نویسیم ممکن است به دو طریق عمل شود:

۱. write through

در این روش در هر بار نوشتن در حافظه داده هم در cache و هم در حافظه اصلی نوشته می شود.

۲. write back

در این روش داده فقط در cache نوشته شده و با یک پرچم set می شود. تا زمانیکه این داده در cache قرار دارد از این داده استفاده خواهد شد، اما در صورت انتقال داده از cache مقدار آن در حافظه اصلی نیز update می گردد.



razeghizade@gmail.com

Razeghizade.pudica.ir

CREADITS: This presentation was created by M.Razeghizadeh
Please keep this slide for attribution