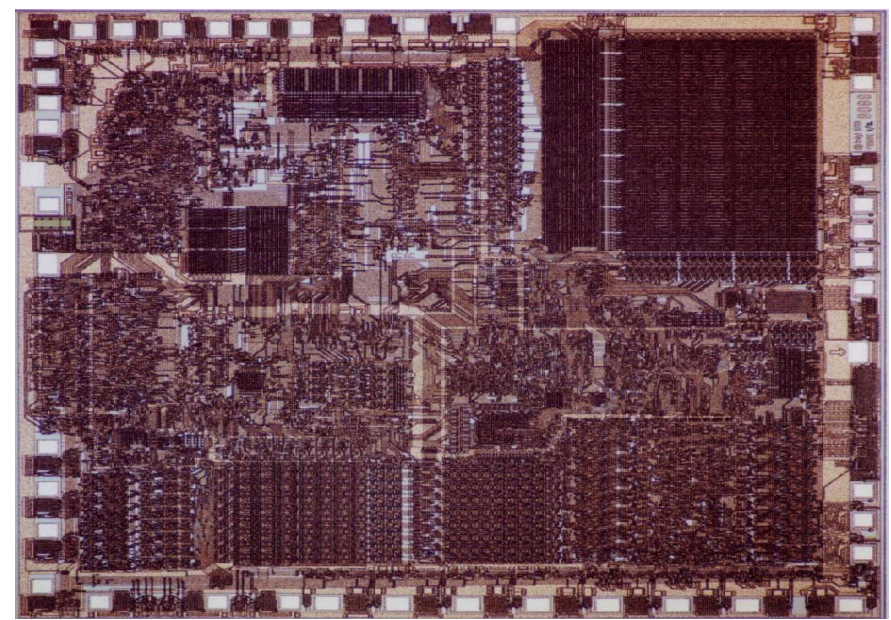
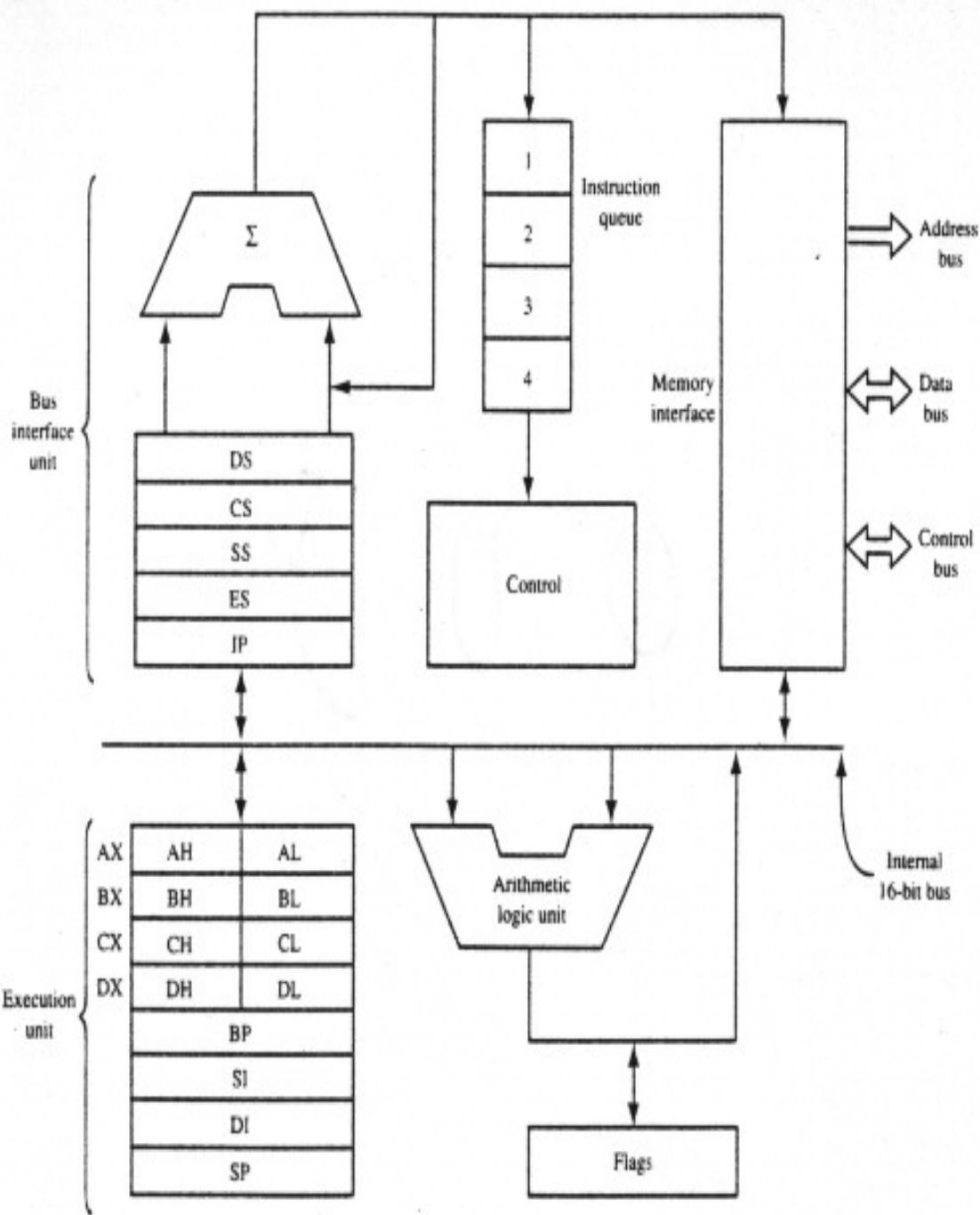




Introduction to Intel IA-32 Architecture

Intel 8086 “Multi-cycle” Architecture (June 1979, 4.77 MHz)

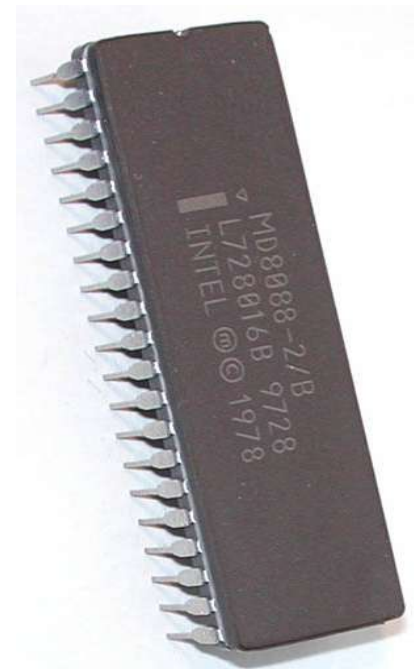


Vss (GND)	1	40	Vcc (+5V)
AD14	2	39	AD15
AD13	3	38	A16/S3
AD12	4	37	A17/S4
AD11	5	36	A18/S5
AD10	6	35	A19/S6
AD9	7	34	BHE/S7
AD8	8	33	MN/MX
AD7	9	32	RD
AD6	10	31	RQ/GT0
AD5	11	30	RQ/GT1
AD4	12	29	LOCK
AD3	13	28	S2
AD2	14	27	S1
AD1	15	26	S0
AD0	16	25	QS0
NMI	17	24	QS1
INTR	18	23	TEST
CLK	19	22	READY
Vss (GND)	20	21	RESET

8086

MIN	MODE
HOLD	WR
HLDA	M/I $\bar{O}$
DT/R	DEN
ALE	INTA

MAX MODE



IA-32 Intel Assembler Instructions

	<u>MIPS</u>	<u>Intel IA-32</u>	<u>C Language</u>
Register	<code>addi \$t0, \$t1, 0</code>	<code>mov eax, ebx</code>	<code>; eax=ebx</code>
Increment	<code>addi \$t0, \$t0, 1</code>	<code>inc eax</code>	<code>; ++eax</code>
Decrement	<code>addi \$t0, \$t0, -1</code>	<code>dec eax</code>	<code>; --eax</code>
Immediate	<code>addi \$t0, \$0, 10</code>	<code>mov eax, 10</code>	<code>; eax=10</code>
Direct addressing	<code>la \$a0, n</code>	<code>mov eax, n</code>	<code>; eax=n</code>
	<code>lw \$t0, 0(\$a0)</code>		
	<code>n .word 3</code>	<code>n dw 3</code>	<code>; int n</code>
Indirect addr.	<code>lw \$t0, 0(\$t1)</code>	<code>mov eax, [ebx]</code>	<code>; eax=*ebx</code>
..with offset	<code>lw \$t0, 8(\$t1)</code>	<code>mov eax, [ebx+8]</code>	<code>; eax=* (ebx+8)</code>
Load byte	<code>la \$a0, c</code>	<code>xor eax, eax</code>	<code>; eax=0</code>
	<code>lb \$t0, 0(\$a0)</code>	<code>mov al, c</code>	<code>; unsigned</code>
	<code>c .byte 4</code>	<code>c db 4</code>	<code>; char c</code>

Intel machine language example

Note: variable size instructions

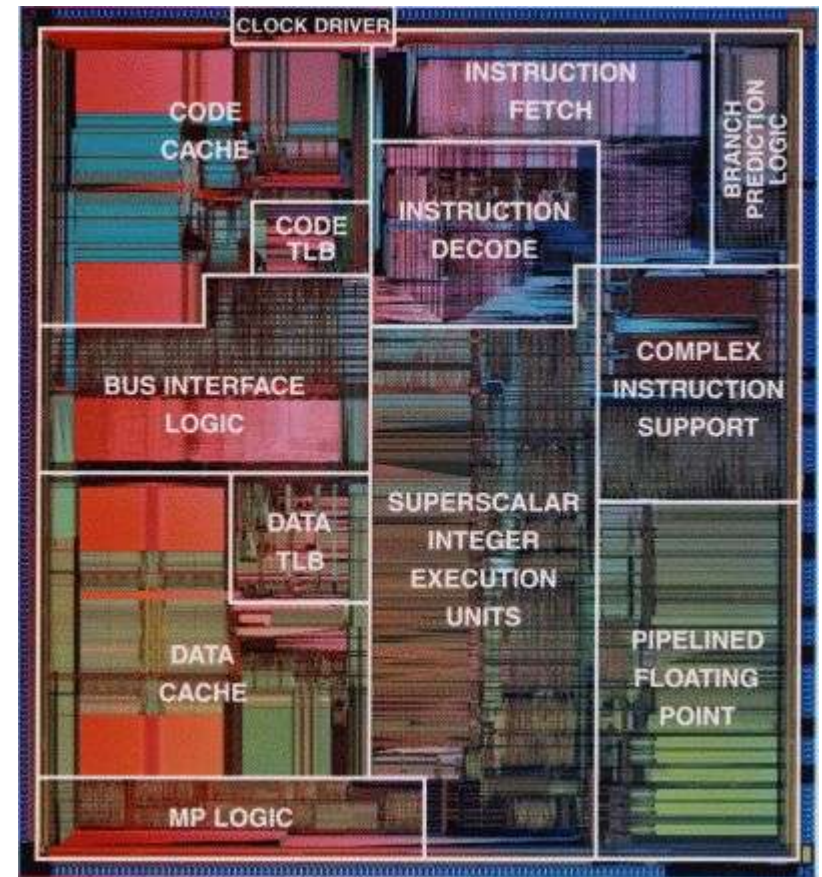
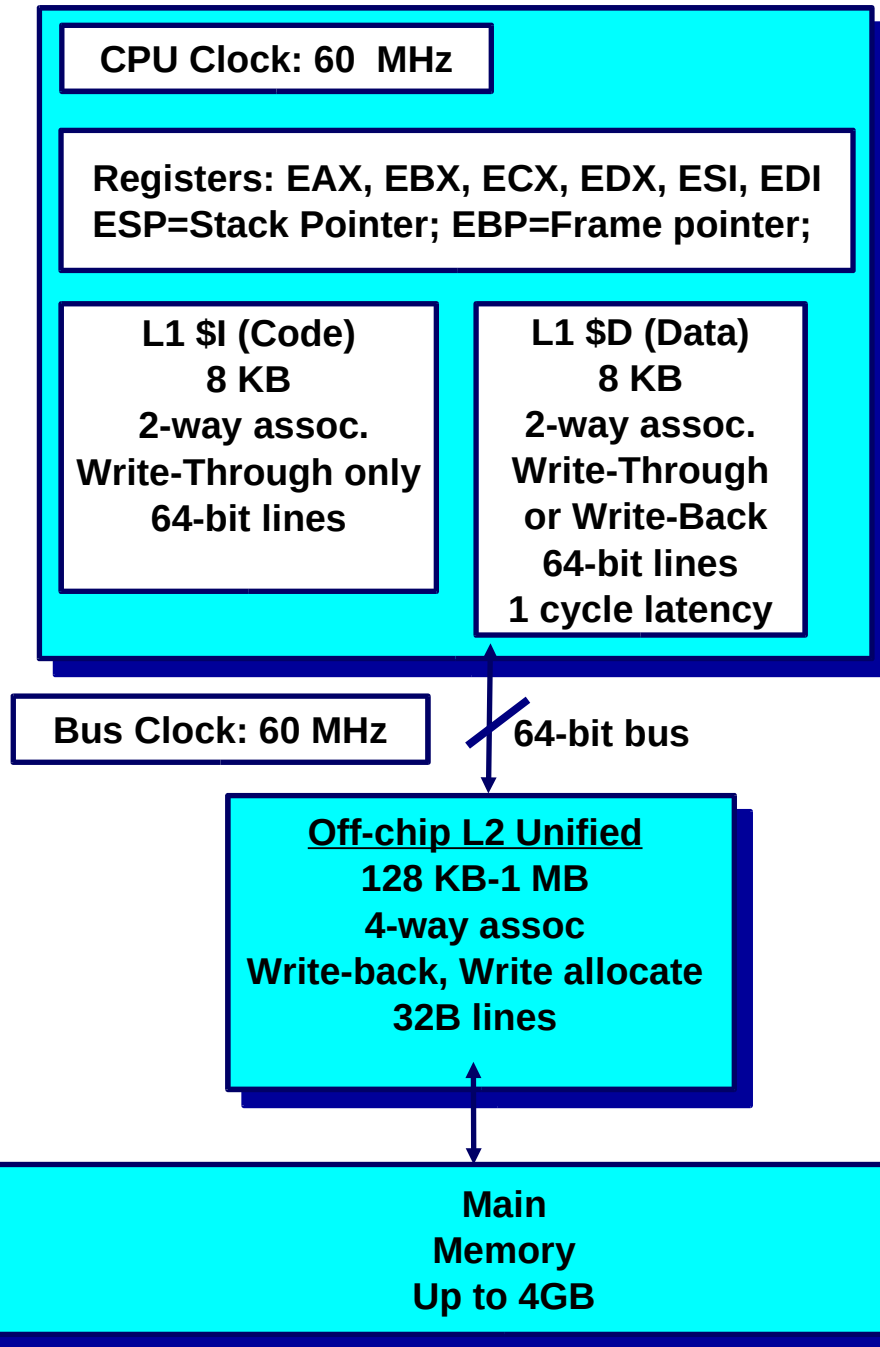
<u>Address</u>	<u>machine code</u>	<u>8086 assembler</u>
0000	000A	a dw 10
0002	00	b db ?
		.code
0000	8B DF	mov bx,di
0002	8A F9	mov bh,cl
0004	8B 1E 0000 R	mov bx,a
0008	8A 26 0002 R	mov ah,b
000C	8B 12	mov dx,[si][bp]
000E	A0 0002 R	mov al,b
0011	8A 26 0002 R	mov ah,b
0015	BB 0003	mov bx,3
0018	B1 03	mov cl,3
001A	C7 06 0000 R 0064	mov a,100
0020	C6 06 0002 R FF	mov b,255

IA-32 For Loop Example

```
; for (i=0; i!=10; i++) { a[i]++; }
```

```
                MOV EBX,A    ; EBX = &A = address of A
                MOV EAX,0    ; EAX = i = 0
FOR:            CMP EAX,10   ; if (i==10)
                JNE EXIT    ; then { goto FEND; }
                INC [EBX]    ; Mem[EBX] (=a[i])++
                ADD EBX,4     ; EBX = &a[i+1]
                INC EAX      ; EAX++
                JMP FOR      ; goto FOR
FEND:          ...
```


Intel Pentium Cache Hierarchy



Introduced March 1993, P60 & P66
Pentium (aka. 80586, P5)
100 MIPS peaks, 50% better than 486
273 pin package, 3.1 Million Transistors
BiCMOS 0.8 micron Process
Production ended in January 1998
<http://www.pcguide.com/ref/cpu/fam/g5P54-c.html>

Intel Pentium P5 “pipeline” Architecture

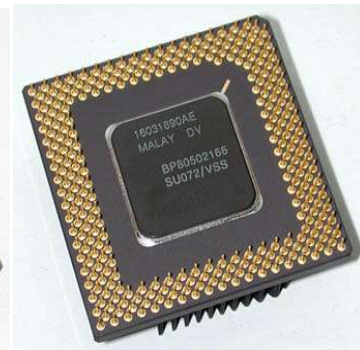
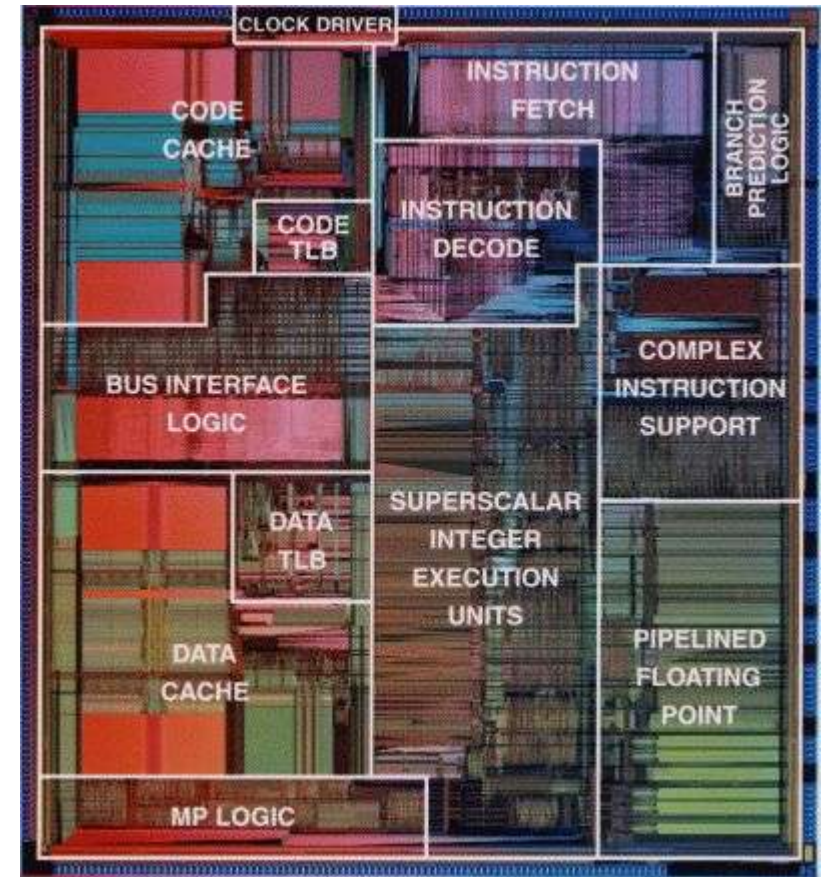
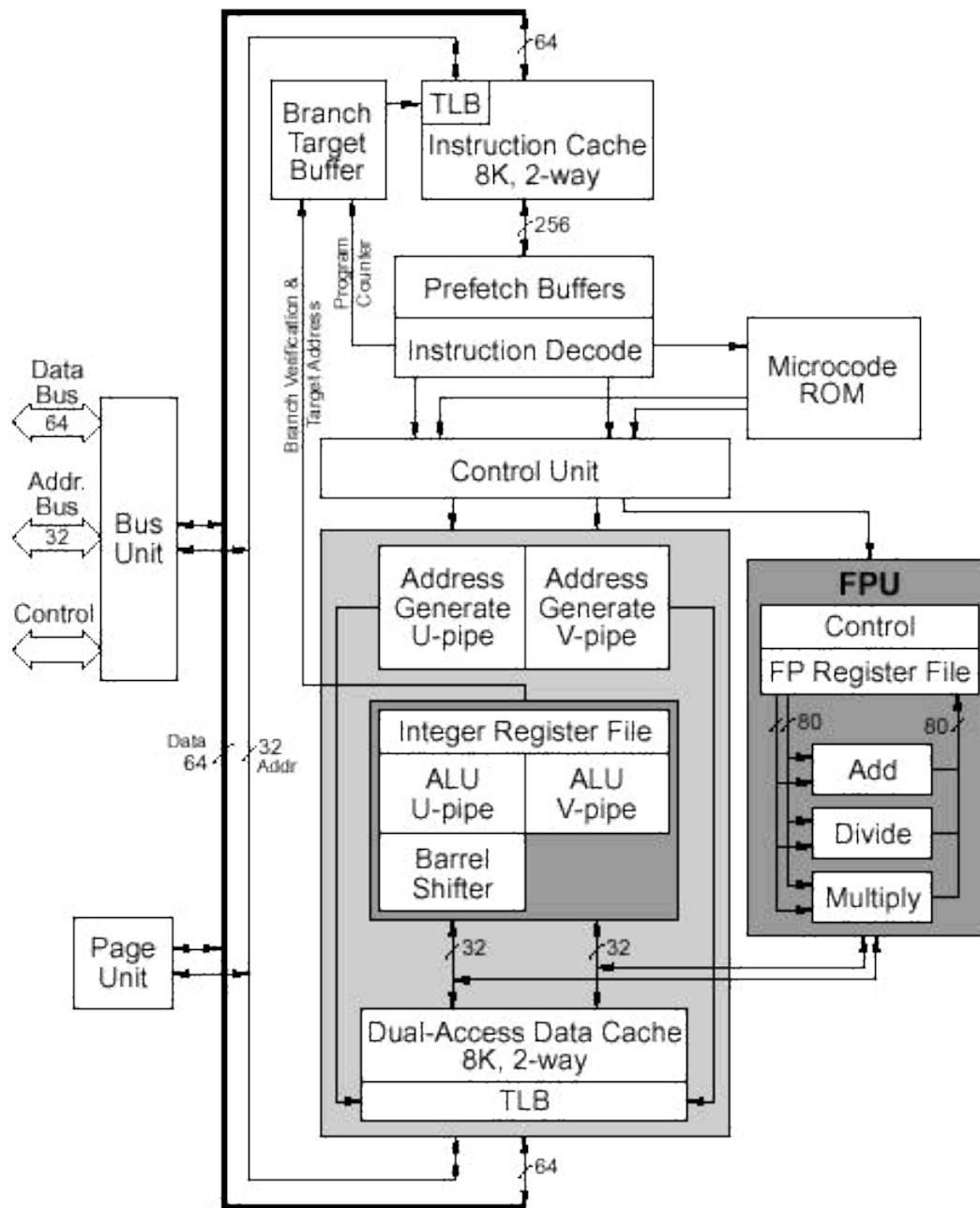


Figure 1. Pentium block diagram.

CPU Pipeline comparisons

<u>CPU</u>	<u># pipeline stages</u>	<u>Max. clock frequency</u>
Pentium	5	300 MHz
Motorola G4	4	500 MHz
Motorola G4e	7	1 GHz
Pentium II and III	10	1.4 GHz
Athlon XP	10/15	2.5 Ghz
Athlon 64	12/17	>3.0 GHz
Pentium 4	20	>3.0 GHz
Pentium 4 "Prescott"	31	>5.0 GHz

Pipeline Differences



The translation of x86 instructions into micro-ops is made in Willamette outside the pipeline

Intel Pentium Architecture Comparison

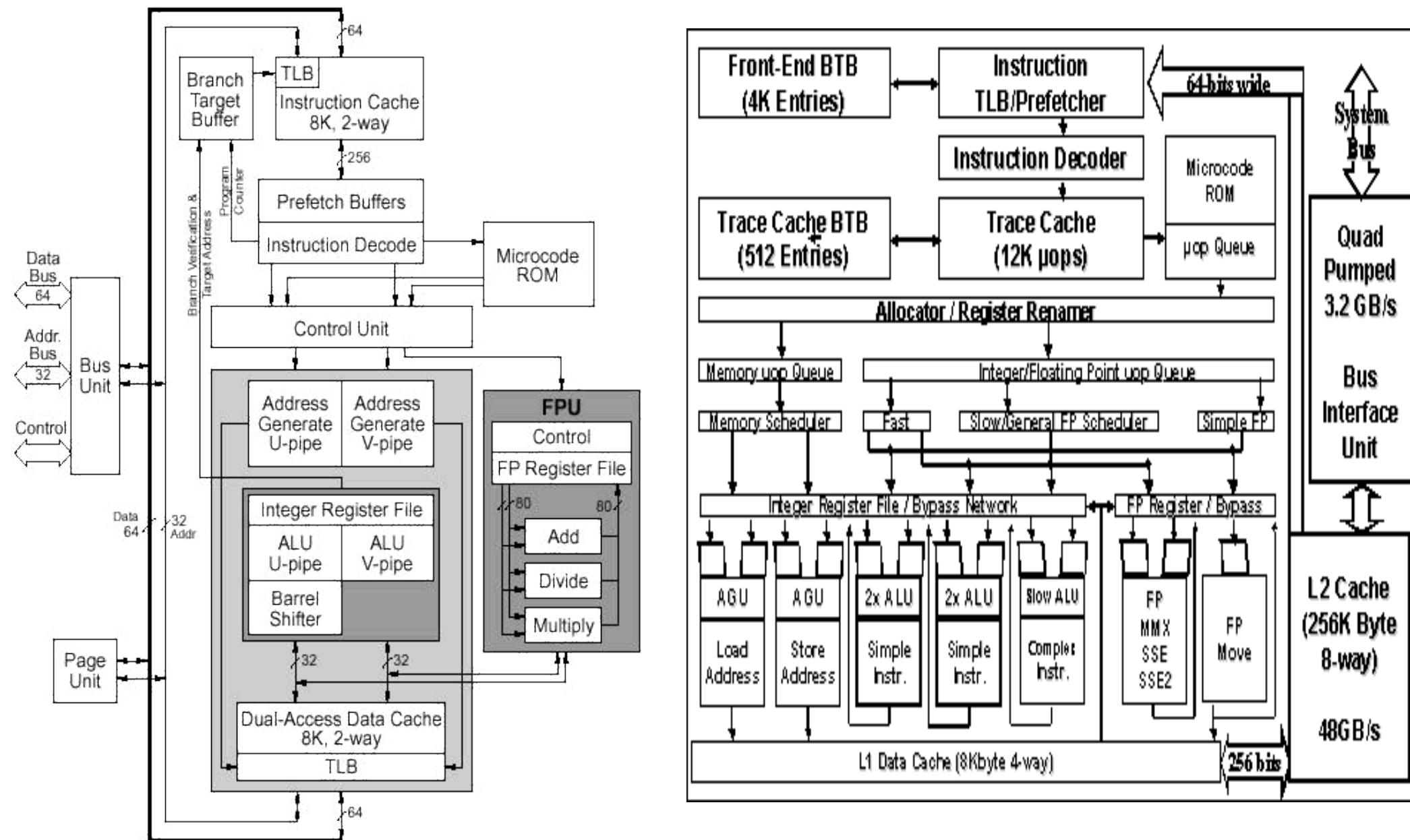
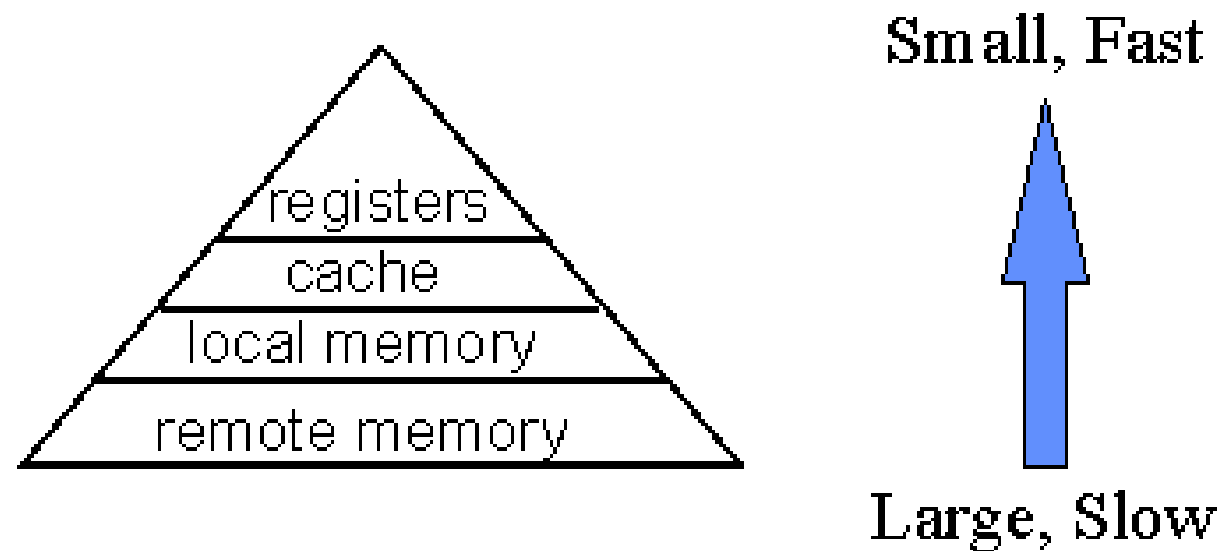


Figure 1. Pentium block diagram.

Memory Hierarchy



- ♦ **Can only do useful work at the top**

- 90-10 rule: 90% of time is spent on 10% of program

- Take advantage of locality

 - temporal locality

 - keep recently accessed memory locations in cache

 - spatial locality

 - keep memory locations nearby accessed memory locations in cache

Matrix Storage

4 by 4 matrix stored in
local memory
by row and by column

A_{11}	A_{12}	A_{13}	A_{14}
A_{21}	A_{22}	A_{23}	A_{24}
A_{31}	A_{32}	A_{33}	A_{34}
A_{41}	A_{42}	A_{43}	A_{44}

Memory locations 

by row
(C)

A_{11} A_{12} A_{13} A_{14}	A_{21} A_{22} A_{23} A_{24}	A_{31} A_{32} A_{33} A_{34}	A_{41} A_{42} A_{43} A_{44}
-------------------------------------	-------------------------------------	-------------------------------------	-------------------------------------

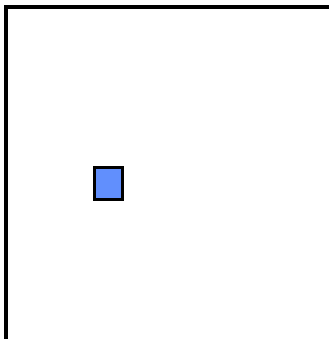
by column
(FORTRAN)

A_{11} A_{21} A_{31} A_{41}	A_{12} A_{22} A_{32} A_{42}	A_{13} A_{23} A_{33} A_{43}	A_{14} A_{24} A_{34} A_{44}
-------------------------------------	-------------------------------------	-------------------------------------	-------------------------------------

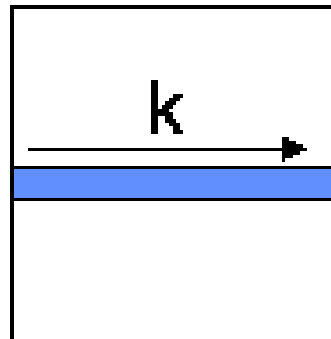
Matrix multiplication: IJK

```
for i = 1, n
  for j = 1, n
    for k = 1, n
      C(i,j) = C(i,j) + A(i,k) * B(k,j)
    end
  end
end
```

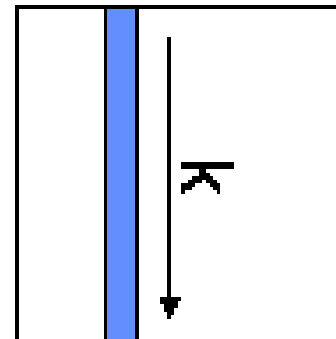
Inner loop computes $C(i, j) = \text{dotproduct}(A(i, :), B(:, j))$



C



A



B

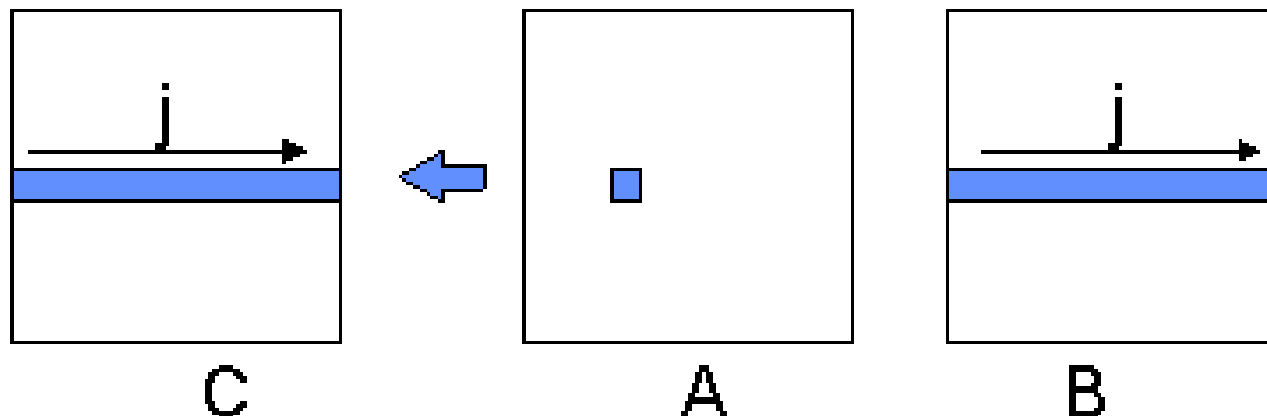
Matrix multiplication: IKJ

Improving Spatial Locality

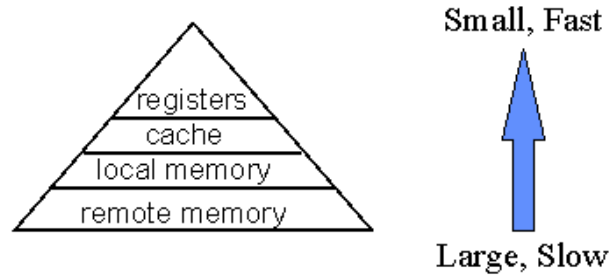
```
for i = 1, n
  for k = 1, n
    for j = 1, n
      C(i,j) = C(i,j) + A(i,k) * B(k,j)
    end
  end
end
```

Stride 1

All inner loop memory references are Stride 1; (refer to consecutive memory locations)



Memory Hierarchy



Can only do useful work at the top

•90-10 rule: 90% of time is spent of 10% of program

•Take advantage of locality

–temporal locality

•keep recently accessed memory locations in cache

–spatial locality

•keep memory locations nearby accessed memory locations in cache

K. Dackland, HPC2N

Matrix Storage

4 by 4 matrix stored in local memory by row and by column

A_{11}	A_{12}	A_{13}	A_{14}
A_{21}	A_{22}	A_{23}	A_{24}
A_{31}	A_{32}	A_{33}	A_{34}
A_{41}	A_{42}	A_{43}	A_{44}

Memory locations

by row (C)

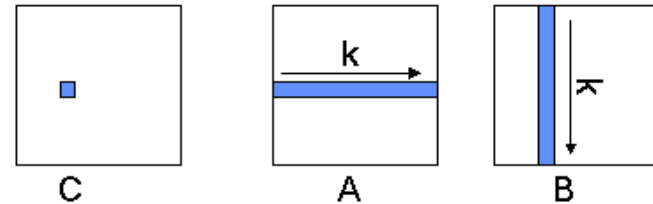
by column (FORTRAN)

K. Dackland, HPC2N

Matrix multiplication: IJK

```
for i = 1, n
  for j = 1, n
    for k = 1, n
      C(i,j) = C(i,j) + A(i,k) * B(k,j)
    end
  end
end
```

Inner loop computes $C(i,j) = \text{dotproduct}(A(i,:), B(:,j))$



K. Dackland, HPC2N

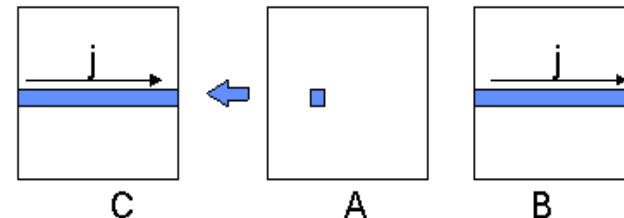
Matrix multiplication: IKJ

Improving Spatial Locality

```
for i = 1, n
  for k = 1, n
    for j = 1, n
      C(i,j) = C(i,j) + A(i,k) * B(k,j)
    end
  end
end
```

Stride 1

All inner loop memory references are Stride 1; (refer to consecutive memory locations)



K. Dackland, HPC2N

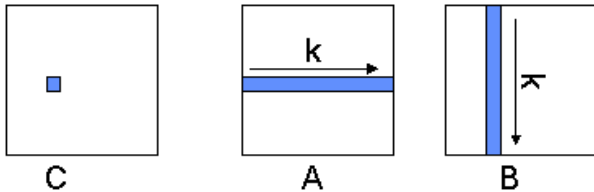
Matrix multiplication: IJK

```

for i = 1, n
  for j = 1, n
    for k = 1, n
      C(i,j) = C(i,j) + A(i,k) * B(k,j)
    end
  end
end

```

Inner loop computes $C(i,j) = \text{dotproduct}(A(i,:), B(:,j))$



K. Dackland, HPC2N

Improving Temporal Locality Blocking for better cache behavior

- Transform multiplication into multiplications of submatrices

$$\begin{array}{ccc}
 C_{11} & C_{12} & C_{13} \\
 C_{21} & C_{22} & C_{23} \\
 C_{31} & C_{32} & C_{33}
 \end{array}
 =
 \begin{array}{ccc}
 A_{11} & A_{12} & A_{13} \\
 A_{21} & A_{22} & A_{23} \\
 A_{31} & A_{32} & A_{33}
 \end{array}
 *
 \begin{array}{ccc}
 B_{11} & B_{12} & B_{13} \\
 B_{21} & B_{22} & B_{23} \\
 B_{31} & B_{32} & B_{33}
 \end{array}$$

$$C_{11} = A_{11} * B_{11} + A_{12} * B_{21} + A_{13} * B_{31}$$

- Size of submatrix multiplication is chosen to fit into cache
 - $m \times m$ submatrices
 - $\sim 2m^3$ flops for $3m^2$ data in cache
- Each word fetched into cache is used m times

K. Dackland, HPC2N

Matrix multiplication: IKJ Improving Spatial Locality

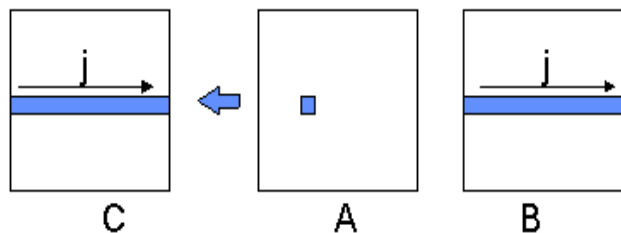
```

for i = 1, n
  for k = 1, n
    for j = 1, n
      C(i,j) = C(i,j) + A(i,k) * B(k,j)
    end
  end
end

```

Stride 1

All inner loop memory references are Stride 1; (refer to consecutive memory locations)



K. Dackland, HPC2N

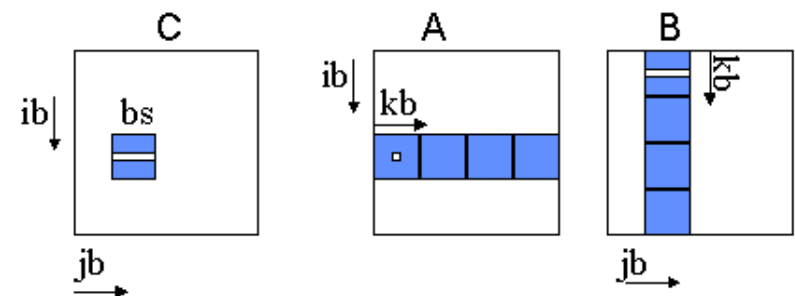
Matrix multiplication: blocked

```

do ib = 1, n, bs
  do jb = 1, n, bs
    do kb = 1, n, bs
      do i = ib, min(ib+bs, n)
        do k = kb, min(kb+bs, n)
          do j = jb, min(jb+bs, n)
            C(i,j) = C(i,j) + A(i,k) * B(k,j)
          end
        end
      end
    end
  end
end

```

Want C to stay in cache as long as possible



K. Dackland, HPC2N