

Arquitectura de Computadores



TEMA 4

Paralelismo a nivel de datos: arquitectura vectorial,
instrucciones SIMD para multimedia, GPUs

DEPARTAMENTO DE
ARQUITECTURA DE **C**OMPUTADORES
Y **A**UTOMÁTICA

Curso 2015-2016

- ❑ Introducción
 - ❑ Arquitectura vectorial
 - o Repertorio de instrucciones vectoriales
 - o Tiempo de ejecución, medidas de rendimiento
 - o Procesamiento selectivo de elementos
 - o Ejemplos
 - ❑ Instrucciones SIMD para procesamiento multimedia
 - ❑ Unidades para procesamiento gráfico (GPUs)
 - o Modelo de programación a alto nivel: CUDA
 - o Instrucciones PTX
 - o Arquitectura del elemento de proceso
 - ❑ Paralelismo a nivel bucle: vectorización
 - o Detección de dependencias
 - ❑ Bibliografía
 - o Cap 4 y Apéndice G de [HePa12]
- (Nota.- En la elaboración de este material se han utilizado algunos contenidos del curso CS152 de la UC Berkeley. También se han utilizado figuras de HePa12)

□ SIMD (Single Instruction Multiple Data).

- o Una operación (codificada como una sola instrucción de LM) se ejecuta sobre un conjunto de datos (en contraposición a SISD).

- o Ejemplo:

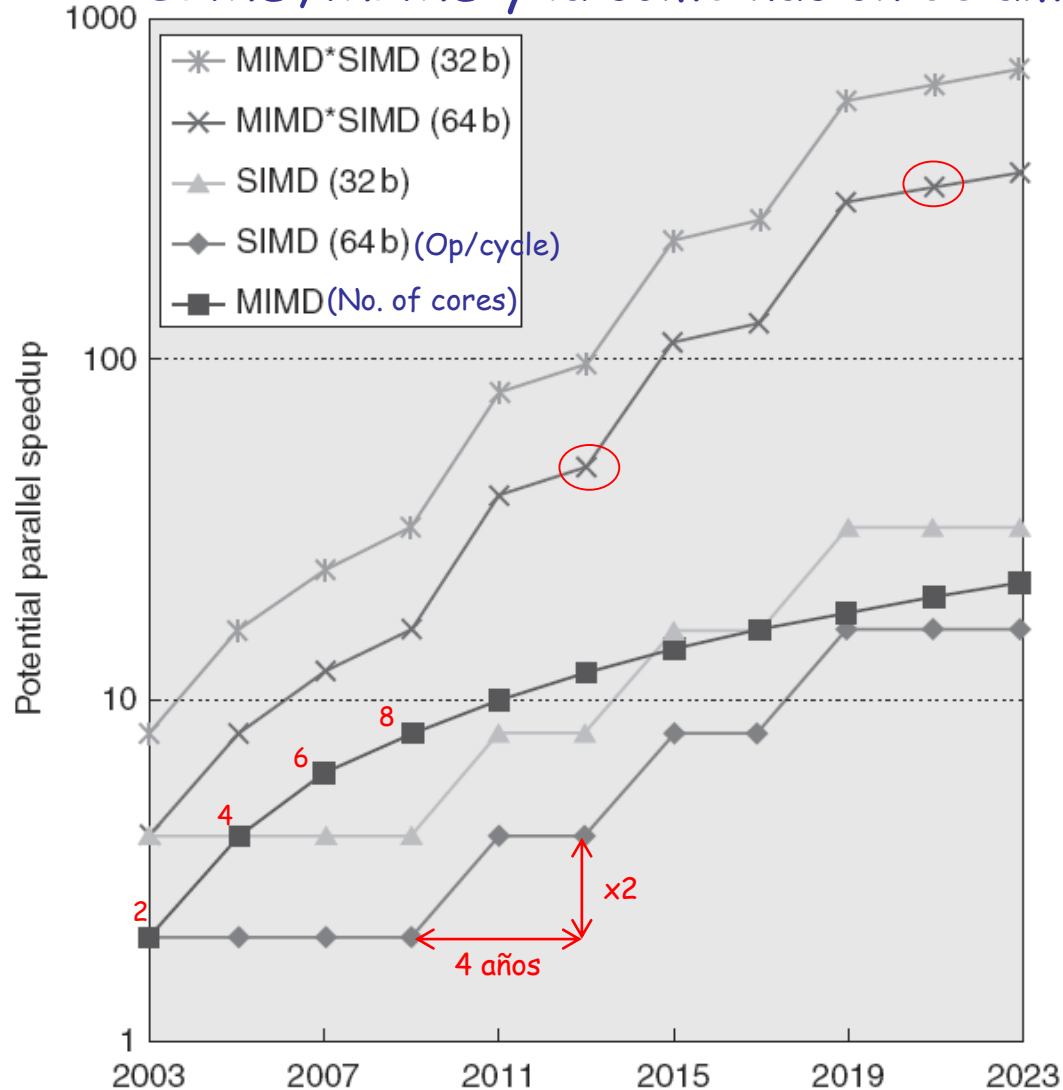
- En lenguaje matemático: $\vec{V3} = \vec{V2} + \vec{V1}$
- En LAN: `for (i = 0; i < N; i++)`
`V3(i) = V2(i) + V1(i);`
- En LM: `ADDV V3,V2,V1`

□ Arquitectura SIMD: puede explotar una cantidad importante de paralelismo de datos en

- o Aplicaciones de ciencia/ingeniería con abundante cálculo matricial (ámbito tradicional)
- o Nuevas aplicaciones: gráficos, visión artificial, comprensión de voz, ...

Introducción

- Predicción de Speedup potencial esperable mediante paralelismos SIMD, MIMD y la combinación de ambos [HePa12]



- Incrementos esperados en arquitecturas 'x86

- o Núcleos : **+2** cada **2 años**
- o Anchura $\rightarrow$ Operaciones de 32/64 bits por ciclo de reloj: **x2** cada **4 años**

- Proyección 2020: En aplicaciones con TLP y DLP $\rightarrow$ Speedup $\sim$ **x10** (Comp 2021 vs. 2013)

- ❑ Eficiencia energética de SIMD: puede ser ventajosa frente a MIMD
 - o Sólo es preciso hacer un “fetch” para operar sobre varios datos.
 - o Ahorro de energía atractivo en dispositivos portátiles
- ❑ En SIMD el programador sigue “pensando” básicamente en un flujo secuencial de instrucciones.
- ❑ Soporte arquitectónico para explotar paralelismo SIMD
 - o Arquitectura vectorial
 - o Extensiones SIMD (extensiones multimedia)
 - o Graphics Processor Units (GPUs)

Arquitectura vectorial

❑ Ideas básicas

- o Leer conjuntos de datos sobre "registros vectoriales"
- o Operar sobre el contenido de estos registros
- o Almacenar los resultados finales en memoria
 - Usar los registros vectoriales para ocultar la latencia de memoria

❑ Características de las operaciones vectoriales

- o Secuencias de cálculos independientes → Ausencia de riesgos
- o Alto contenido semántico
 - Una instrucción → Muchas operaciones
- o Patrón de accesos a memoria conocido
- o Explotación eficiente de memoria entrelazada
- o Disminución de instrucciones de salto (un bucle completo puede transformarse en una instrucción)
 - Reducción conflictos de control

Arquitectura vectorial

❑ En el principio... Seymour Cray - CDC 6600 (1963)



❑ No es una arquitectura vectorial, pero...

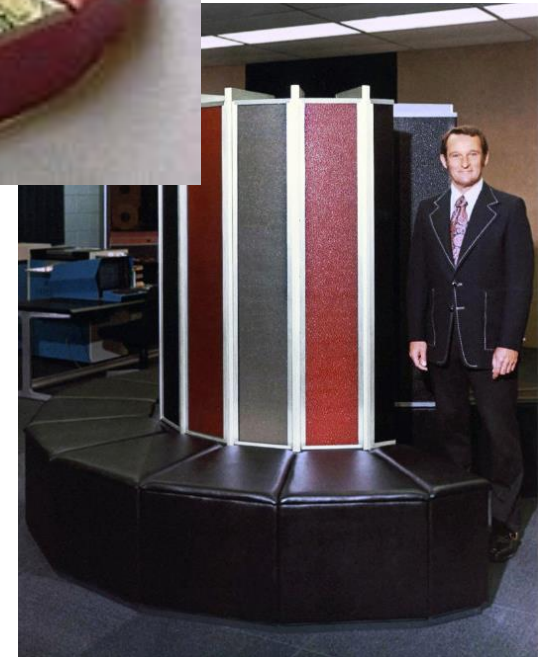
- o Muy segmentada, con palabra de 60 bits
- o Mp de 128 Kword con 32 bancos
- o 10 Fus (paralelas, no segmentadas)
 - PF: sumador, 2 mult, divisor
- o Control cableado
- o Planificación dinámica de instrucciones (scoreboard)
- o 10 procesadores de E/S
- o Clock 10 MHz
 - Muy rápido para la época
 - Suma FP en 4 ciclos
- o >400,000 transistores, 750 sq. ft. (~70 m²), 5 tons, 150 KW, refrigeración freon
- o Máquina más rápida del mundo durante 5 años (hasta el 7600)
- o Ventidas >100 (a \$7-10M c.u.)

- ❑ En el principio... Seymour Cray - CDC 6600 (1963)
- ❑ Thomas Watson Jr., IBM CEO, August 1963:

"Last week, Control Data ... announced the 6600 system. I understand that in the laboratory developing the system there are only 34 people including the janitor. Of these, 14 are engineers and 4 are programmers... Contrasting this modest effort with our vast development activities, I fail to understand why we have lost our industry leadership position by letting someone else offer the world's most powerful computer."
- ❑ To which Cray replied: "It seems like Mr. Watson has answered his own question."

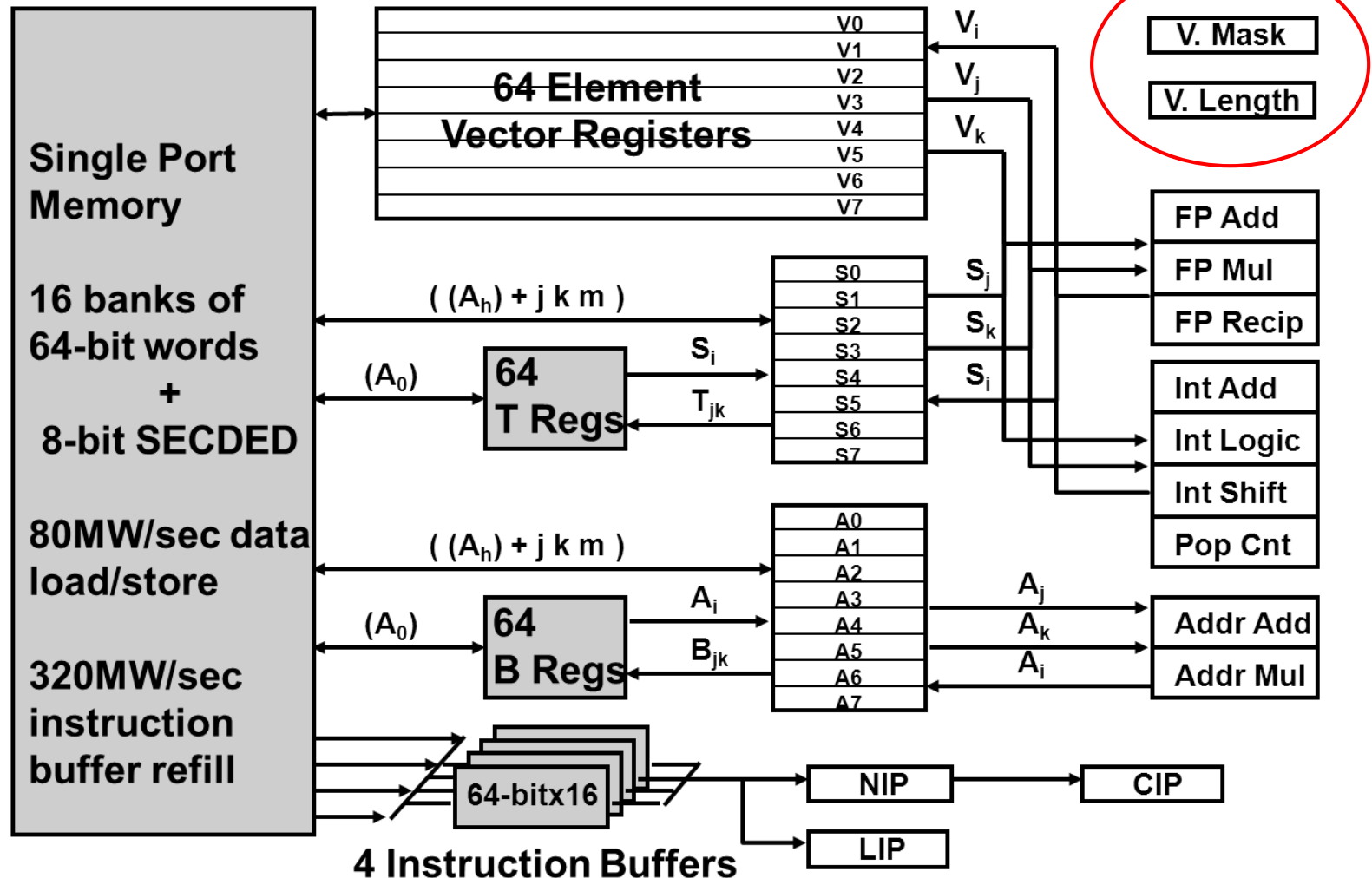
El Cray-1 (1976)

- ❑ Unidad escalar
 - o Arquitectura Load/Store
- ❑ Extensión Vectorial
 - o Registros Vectoriales
 - o Instrucciones Vectoriales
- ❑ Implementación
 - o Control cableado
 - o UF muy segmentadas
 - o Memoria entrelazada
 - o Sin cache de datos
 - o Sin memoria virtual



Arquitectura vectorial

□ Visión global del Cray-1



T ciclo banco memoria: 50 ns, T ciclo procesador: 12,5 ns (80 Mhz)

Arquitectura vectorial

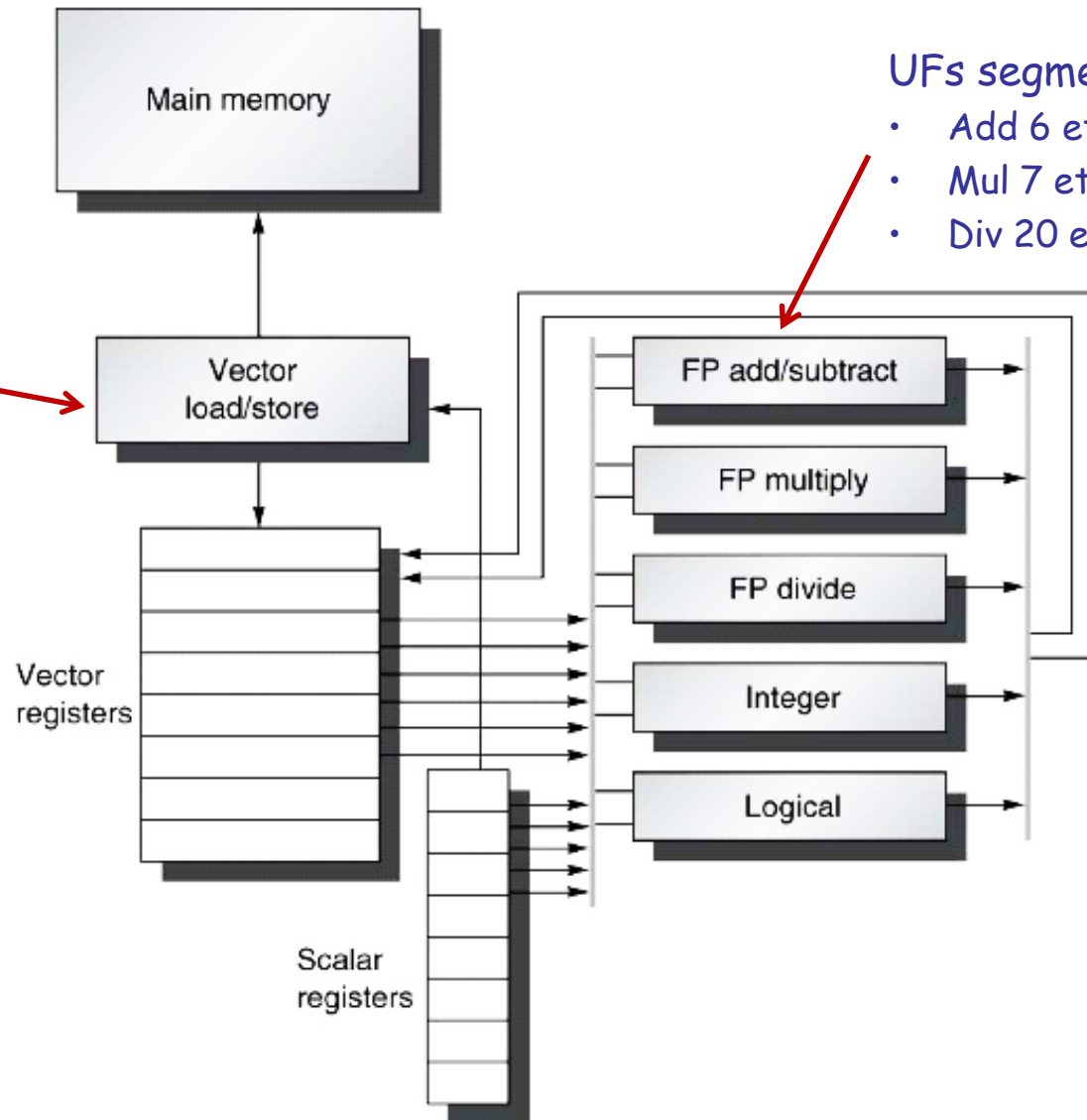
□ Estructura de un procesador vectorial: VMIPS

Funcionalmente
equivalente a un pipe: un
dato por ciclo, después de
latencia inicial (12 ciclos).

UFs segmentadas:

- Add 6 etapas
- Mul 7 etapas
- Div 20 etapas

8 registros vectoriales
1 reg = 64 elementos
1 elemento = 64 bits
16 read, 8 write ports



Arquitectura vectorial: repertorio VMIPS (1)

- ❑ Operaciones vectoriales aritméticas sobre registros
- ❑ Instrucciones especiales de carga/almac de vectores (LV,SV)
- ❑ Modos de direccionamiento especiales para vectores no contiguos. Ejemplos
 - o LVWS: Load Vector With Stride. Carga elementos equiespaciados a una cierta distancia
 - o LVI: Load Vector using Index. El contenido de un registro vectorial indica las posiciones de los elementos a cargar.
- ❑ Registros especiales de longitud vectorial (VLR) y máscara (VM)
 - o Reg VLR: Indica la longitud de los vectores a procesar (≤ 64)
 - o Reg VM: Registro de 64 bits. Para las posiciones con VM a cero, la operación no se ejecuta.
 - Ejecución selectiva de operaciones sobre componentes

Arquitetura vectorial: repertorio VMIPS (2)

Instruction	Operands	Function
ADDVV.D	V1,V2,V3	Add elements of V2 and V3, then put each result in V1.
ADDVS.D	V1,V2,F0	Add F0 to each element of V2, then put each result in V1.
SUBVV.D	V1,V2,V3	Subtract elements of V3 from V2, then put each result in V1.
SUBVS.D	V1,V2,F0	Subtract F0 from elements of V2, then put each result in V1.
SUBSV.D	V1,F0,V2	Subtract elements of V2 from F0, then put each result in V1.
MULVV.D	V1,V2,V3	Multiply elements V2 and V3, then put each result in V1.
MULVS.D	V1,V2,F0	Multiply each element of V2 by F0, then put each result in V1.
DIVVV.D	V1,V2,V3	Divide elements of V2 by V3, then put each result in V1.
DIVVS.D	V1,V2,F0	Divide elements of V2 by F0, then put each result in V1.
DIVSV.D	V1,F0,V2	Divide F0 by elements of V2, then put each result in V1.
LV	V1,R1	Load vector register V1 from memory starting at address R1.
SV	R1, V1	Store vector register V1 into memory starting at address R1.
LVWS	V1,(R1,R2)	Load V1 from address at R1 with stride in R2 (i.e .. $R1 + i \times R2$).
SVWS	(R1,R2),V1	Store V1 to address at R1 with stride in R2 (i.e .. $R1 + i \times R2$).
LVI	V1,(R1+V2)	Load V1 with vector whose elements are at $R1 + V2(i)$ (i.e., V2 is an index).
SVI	(R1+V2) ,V1	Store V1 to vector whose elements are at $R1 + V2(i)$ (i.e., V2 is an index).

Arquitetura vectorial: repertorio VMIPS (3)

Instruction	Operands	Function
CVI	V1,R1	Create an index vector by storing the values 0, 1 x R1, 2 x R1, ... ,63 x R1 into V1.
S--VV.D	V1, V2	Compare the elements (EQ, NE, GT, LT, GE, LE) in V1 and V2. If condition is true, put a 1 in the corresponding bit vector: otherwise put 0. Put resulting bit vector in vector mask register (VM). The instruction S--VS.D performs the same compare but using a scalar value as one operand
S--VS.D	V1, F0	
POP	R1,VM	Count the 1s in vector-mask register VM and store count in R1.
CVM		Set the vector-mask register to all 1s.
MTC1	VLR,R1	Move contents of R1 to vector-length register VL.
MFC1	R1,VLR	Move the contents of vector-length register VL to R1.
MVTM	VM,F0	Move contents of F0 to vector-mask register VM.
MVFM	F0,VM	Move contents of vector-mask register VM to F0.

Código escalar vs. Código vectorial

□ $Y = a * X + Y$ (vectores de 64 elementos)

o Versión escalar

	L.D	F0, a	; load scalar a
	DADDIU	R4,Rx,#512	; last address to load
Loop:	L.D	F2, 0(Rx)	; Load X[i]
	MUL.D	F2, F2, F0	; A x X[i]
	L.D	F4, 0(Ry)	; Load Y[i]
	ADD.D	F4, F4, F2	; A x X[i] + Y[i]
	S.D	F4, 0(Ry)	; Store Y[i]
	DADDIU	Rx, Rx, #8	; increment index X
	DADDIU	Ry, Ry, #8	; increment index Y
	DSUBU	R20, R4, Rx	; loop exhausted?
	BNEZ	R20, Loop	

o Versión vectorial

	L.D	F0, a	; load scalar a
	LV	V1, Rx	; Load vector X
	MULVS.D	V2, V1, F0	; Vector-scalar multiply
	LV	V3, Ry	; Load vector Y
	ADDVV.D	V4, V2, V3	; Vector addition
	SV	V4, Ry	; Store result Y

o Instrucciones ejecutadas: ~ 580 vs. 6 !!

Tiempo de ejecución

- ❑ Dependiente básicamente de tres factores
 - o Longitud de los vectores operandos
 - o Riesgos estructurales: las UF necesarias están ocupadas, no hay puertos del BR disponibles
 - o Dependencias de datos
- ❑ Velocidad de procesamiento
 - o Las FUs del VMIPS consumen (y producen) un elemento por ciclo de reloj
 - o El tiempo de ejecución de una operación vectorial es aproximadamente igual a la longitud del vector
- ❑ Convoy
 - o Se denomina así a un conjunto de (una o varias) instrucciones vectoriales que potencialmente pueden ejecutarse juntas (ausencia de riesgos estructurales). Pueden tener riesgos LDE.
- ❑ Paso (chime)
 - o Unidad de tiempo para ejecutar un convoy
 - o m convoyes se ejecutan en m pasos
 - o Para vectores de longitud n , ejecutar m convoyes requiere (aprox.) $m \times n$ ciclos de reloj (Notación: $T_{\text{chime}}=m$)

□ Convoyes: ejemplo

1: LV	V1, Rx	; load vector X
2: MULVS.D	V2, V1, F0	; vector-scalar multiply
3: LV	V3, Ry	; load vector Y
4: ADDVV.D	V4, V2, V3	; Vector addition
5: SV	V4, Ry	; Store result Y

o Conflictos:

- 1 y 2 no tienen conflictos estructurales
- 3 tiene conflicto estructural con 1 (una sola unidad de Load/Store)
- 4 no tiene conflictos estructurales con 3
- 5 tiene conflicto estructural con 3

o Convoyes resultantes

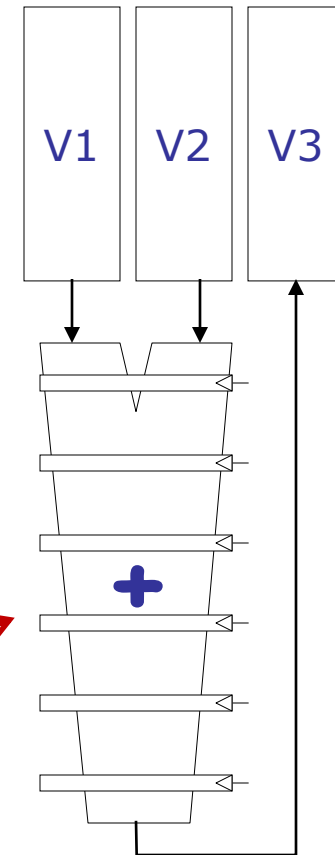
- 1. Formado por LV y MULVS.D
- 2. Formado por LV y ADDVV.D
- 3. Formado por SV

o Tiempo de cálculo aprox para 64 componentes: $3 \times 64 = 192$ ciclos ($T_{chime} = 3$)

Ejecución de operaciones aritméticas

- ❑ Uso de un pipe profundo para la ejecución de las operaciones (reduce el ciclo de reloj)
- ❑ Alta latencia
 - o No demasiado relevante debido a la falta de dependencia entre los cálculos sobre un vector

UF de suma segmentada
en 6 etapas



$$V3 \leftarrow V1 + V2$$

Ejecución de operaciones aritméticas

□ Operaciones independientes

o Ejemplo

MULVV.D V1, V2, V3

ADDVV.D V4, V5, V6

(1 convoy: la UF de * y la de + pueden actuar a la vez)

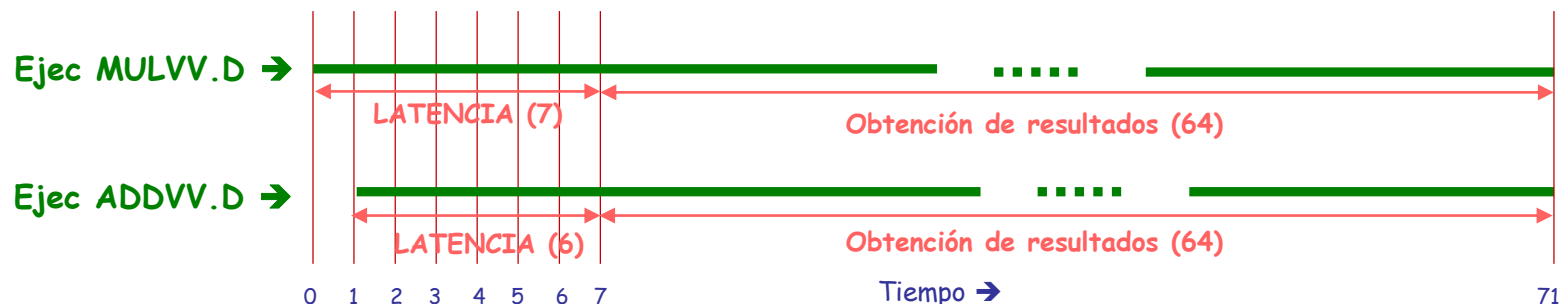
o Comportamiento temporal (1 paso)

- Recordar: MUL 7 ciclos, ADD 6 ciclos

Operación	Inicio	Fin
MULVV.D	0	$7+64 = 71$
ADDVV.D	1	$1+6+64 = 71$

o En ausencia de conflictos lanza una instrucción por ciclo

o Representación



Ejecución de accesos a memoria

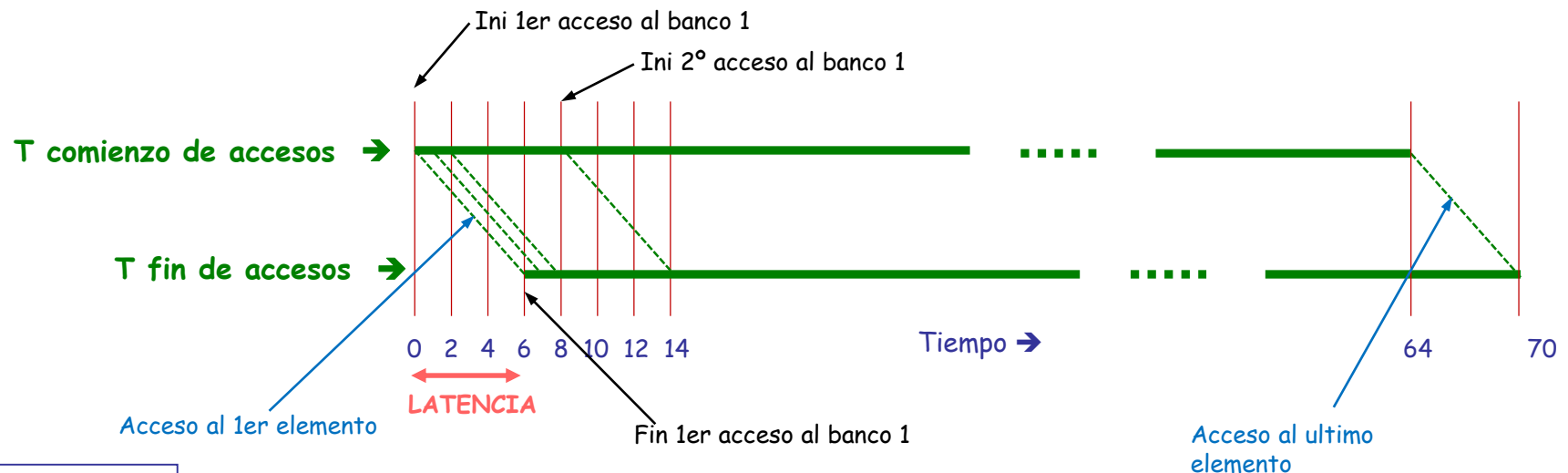
❑ Memoria entrelazada (equivalencia funcional con un pipe)

o Ejemplo: Sup 8 bancos, T acceso a memoria: 6 ciclos.

- Carga de un vector de 64 componentes que comienza en la dirección 136
- Formato de **dirección**: ...xxxx yyy 000 (siendo yyy = n° de banco)

Dir	136	144	152	160	168	176	184	192	200	208	...
Banco	1	2	3	4	5	6	7	0	1	2	...
Tini	0	1	2	3	4	5	6	7	8	9	...
Tfin	6	7	8	9	10	11	12	13	14	15	...

- Diagrama temporal



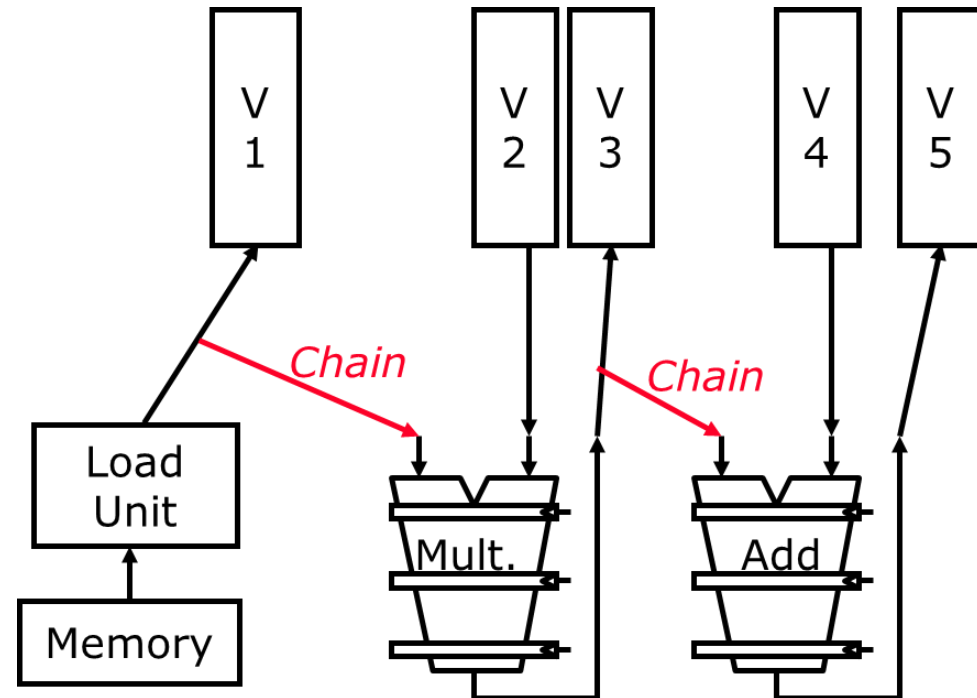
❑ Operaciones dependientes: Tratamiento de dependencias RAW

❑ Problema

LV	V1
MULVV.D	V3, V1, V2
ADDVV.D	V5, V3, V4

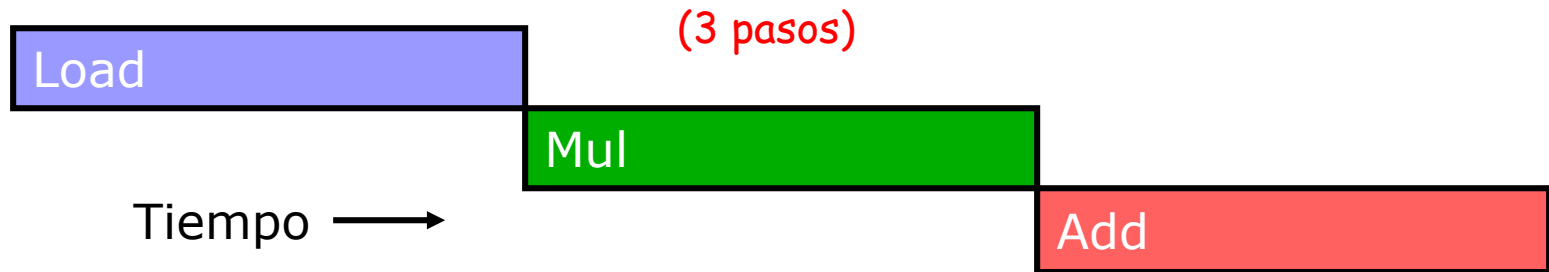
❑ Solución:

- o Cray-1. Extensión del concepto de anticipación de operandos $\Rightarrow$ Encadenamiento de operaciones (chaining)
- o 3 operaciones, pero 1 paso
- o En proc modernos: "encadenamiento flexible"

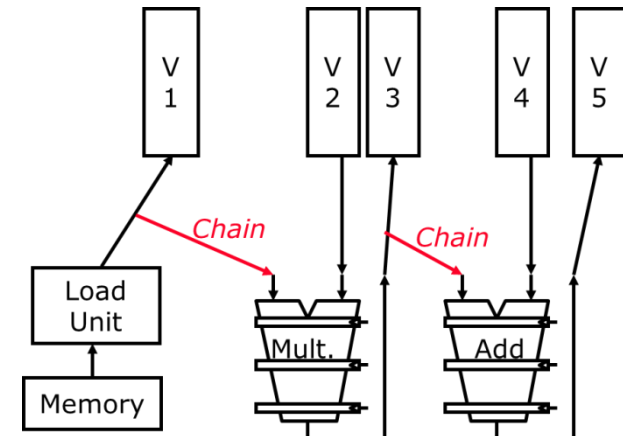
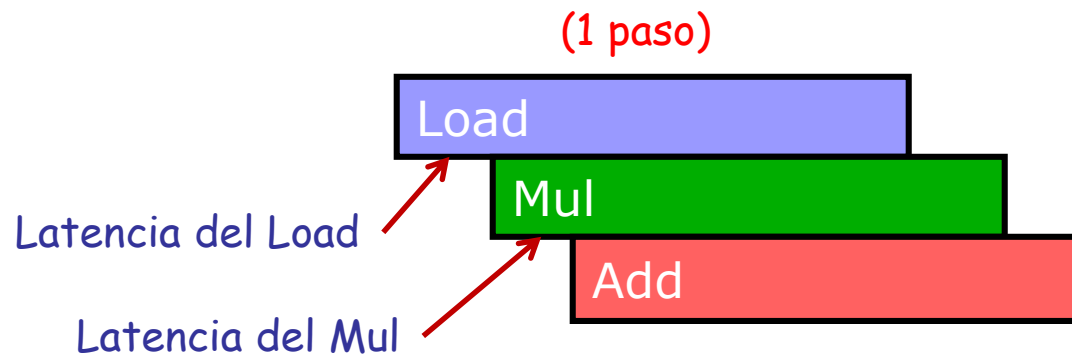


Encadenamiento

- ❑ Sin encadenamiento: esperar hasta que se haya calculado el último elemento de la operación anterior



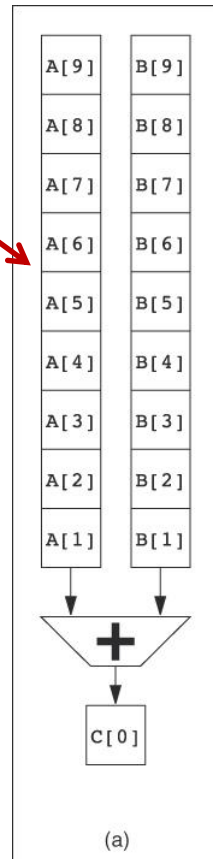
- ❑ Con encadenamiento: Una instrucción puede comenzar cuando está disponible el primer elemento de la operación de la que depende



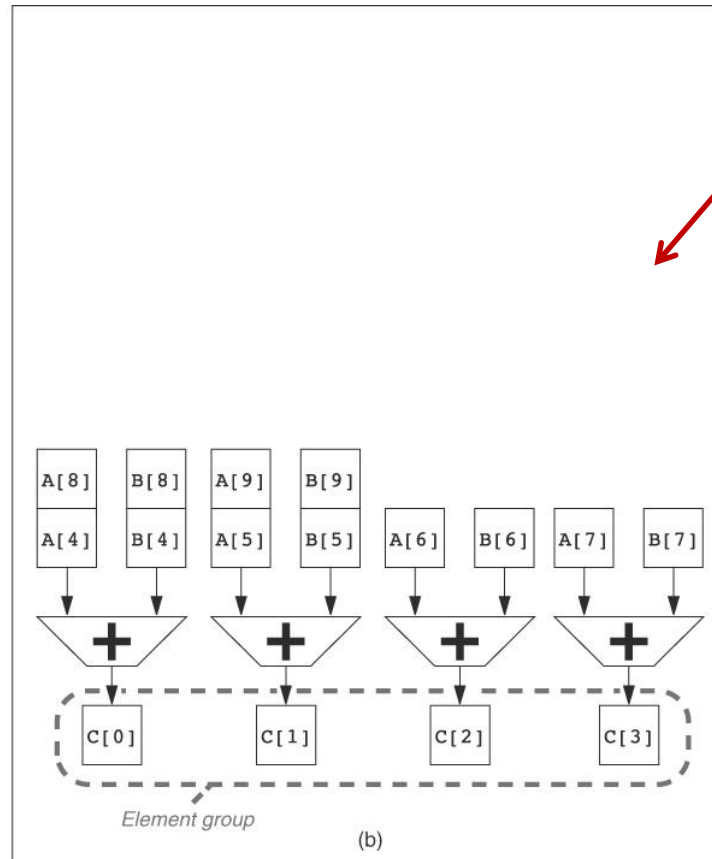
Procesamiento con vías múltiples

- ❑ Aceleración de los cálculos poniendo varias UF segmentadas del mismo tipo
 - o Solamente una parte de las componentes es procesada por cada UF.
 - o Esquema

Ejecución
vectorial
convencional



Ejecución
vectorial con
varias vías



Vectores de longitud arbitraria

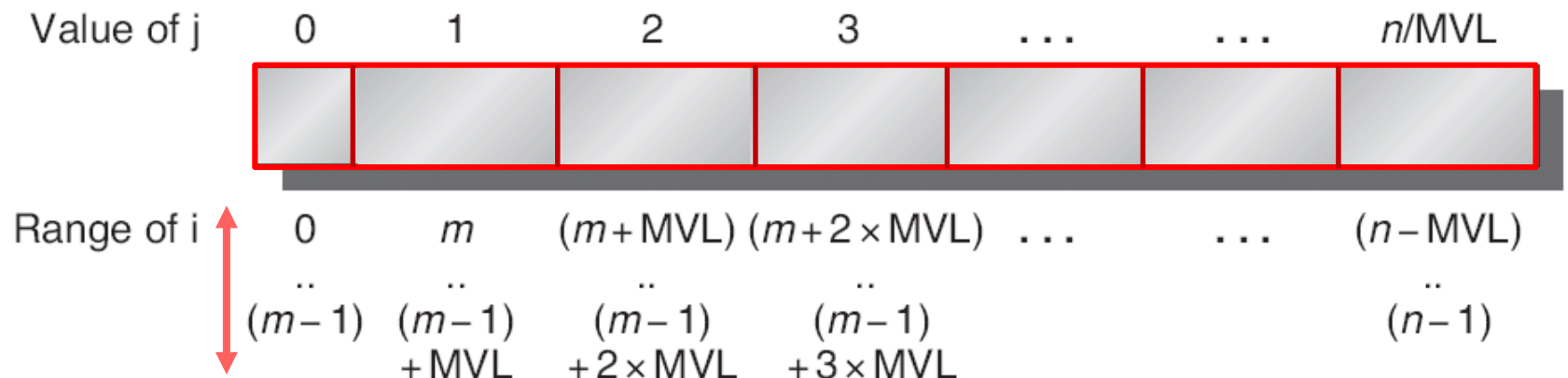
- ❑ Registro de longitud vectorial (VLR)
 - o El valor cargado en VLR determina el n° de componentes sobre los que actúa la op. vectorial lanzada.
- ❑ Vectores cortos: $n \leq \text{long. registros vectoriales (MVL)}$
 - o Cargar registro VLR
 - o Ejecutar operación vectorial
- ❑ Vectores largos: $n > \text{MVL} \rightarrow$ Proceso por bloques (strip mining)
 - o Descomponer operación en varias suboperaciones
 - o n/MVL operaciones vectoriales de longitud MVL
 - o 1 operación vectorial de longitud $(n \bmod \text{MVL})$

Vectores de longitud arbitraria (strip mining)

❑ Ejemplo: Bucles anidados para ejecución de DAXPY

```
low = 0;
VL = (n % MVL);                                /*find odd-size piece using modulo */
for (j = 0; j <= (n/MVL); j=j+1) {             /*outer loop*/
    for (i = low; i < (low+VL); i=i+1)          /*runs for length VL*/
        Y[i] = a * X[i] + Y[i];                /*main operation*/
    low = low + VL;                             /*start of next vector*/
    VL = MVL;                                   /*reset the length to MVL*/
}
```

❑ Diagrama de ejecución de DAXPY ($m = n \bmod MVL$)



Vectores de longitud arbitraria (strip mining)

□ Modelo de rendimiento para operaciones por bloques

- o Inicialización de operaciones (T_{base}): Cálculo de dir iniciales, op escalares de preparación del bucle (1 sola vez)
 - Simplificación: despreciable
- o Penalización por inicialización y control del bucle (1 vez por cada itearación)
 - T_{start} : n° ciclos para el llenado de pipes. Depende de las operaciones vectoriales incluidas en el bucle. Si las instr son independientes, el llenado de pipes se solapa.
 - T_{loop} : actualización de punteros, detección de fin. Simplifiación: 15 ciclos
- o N° de convoyes en el bucle (T_{chime})
- o Tiempo total de cálculo para vectores de n elementos

$$T_n = \left\lfloor \frac{n}{MVL} \right\rfloor \times (T_{loop} + T_{start}) + n \times T_{chime}$$



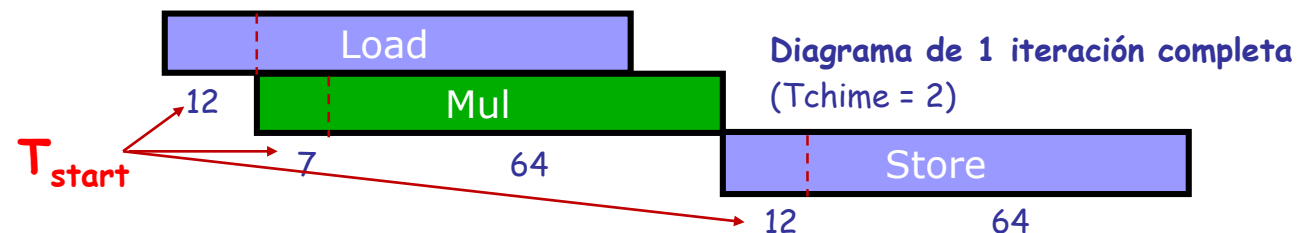
N° de iteraciones

Vectores de longitud arbitraria (strip mining)

□ Ejemplo: $A = B \times s$, para vectores de 200 componentes

	DADDUI	R2,R0,#1600	;total # bytes in vector
	DADDU	R2,R2,Ra	;address of the end of A vector
	DADDUI	R1,R0,#8	;loads length of 1st segment
	MTC1	VLR,R1	;load vector length in VLR
	DADDUI	R1,R0,#64	;length in bytes of 1st segment (8 elements)
	DADDUI	R3,R0,#64	;vector length of other segments (64 elements)
Loop:	LV	V1,Rb	;load B
	MULVS.D	V2,V1,Fs	;vector * scalar
	SV	Ra,V2	;store A (structural hazard)
	DADDU	Ra,Ra,R1	;address of next segment of A
	DADDU	Rb,Rb,R1	;address of next segment of B
	DADDUI	R1,R0,#512	;load byte offset next segment
	MTC1	VLR,R3	;set length to 64 elements
	DSUBU	R4,R2,Ra	;at the end of A?
	BNEZ	R4,Loop	;if not, go back

$$\begin{aligned}
 \square T_n &= \left\lceil \frac{n}{MVL} \right\rceil \times (T_{loop} + T_{start}) + n \times T_{chime} = \\
 &= \left\lceil \frac{200}{64} \right\rceil \times (15 + (12 + 7 + 12)) + 200 \times 2 = 4 \times 46 + 400 = 584 \text{ ciclos}
 \end{aligned}$$



Medidas de rendimiento

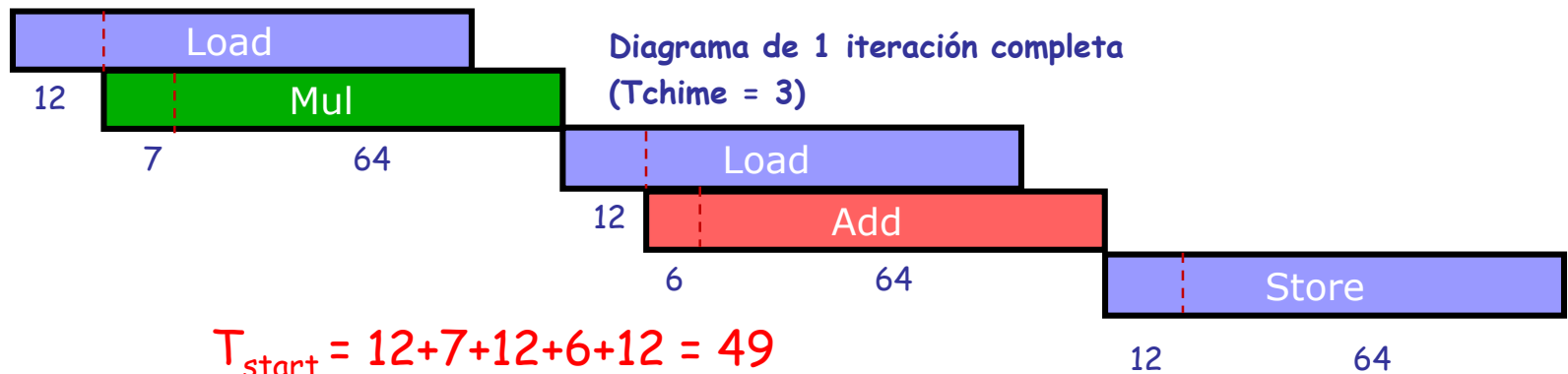
□ Rendimiento asintótico (R_∞): MFLOPS obtenidos para para supuestos vectores de longitud infinita

o Consideremos la op DAXPY sin limitaciones debidas a la longitud de registros vectoriales (tr. 15, 3 convoyes)

▪ $2n$ FLOP en $3n$ ciclos $\Rightarrow R_\infty = 2/3$ FLOP/ciclo.

▪ Si sup $f=500$ MhZ $\Rightarrow R_\infty = 2/3 \times 500 \times 10^6 = 333$ MFLOPS

o Efecto de strip mining: Suponemos $MVL = 64$



$$T_n = \left\lfloor \frac{n}{MVL} \right\rfloor \times (T_{loop} + T_{start}) + n \times T_{chime} = \left\lfloor \frac{n}{64} \right\rfloor \times (15 + 49) + n \times 3$$

$$R_\infty = \lim_{n \rightarrow \infty} \frac{2n}{T_n} = \lim_{n \rightarrow \infty} \frac{2n}{4n} = \frac{1 \text{ FLOP}}{2 \text{ ciclo}} = 250 \text{ MFLOPS !!}$$

Medidas de rendimiento

- Longitud del rendimiento mitad del asintótico ($N_{1/2}$)
- Ejemplo: Sup que se obtiene con $n < MVL \rightarrow 1$ iteración

$$T_n = \left\lfloor \frac{n}{MVL} \right\rfloor \times (T_{loop} + T_{start}) + n \times T_{chime} = 1 \times (15 + 49) + n \times 3 = 64 + 3n$$

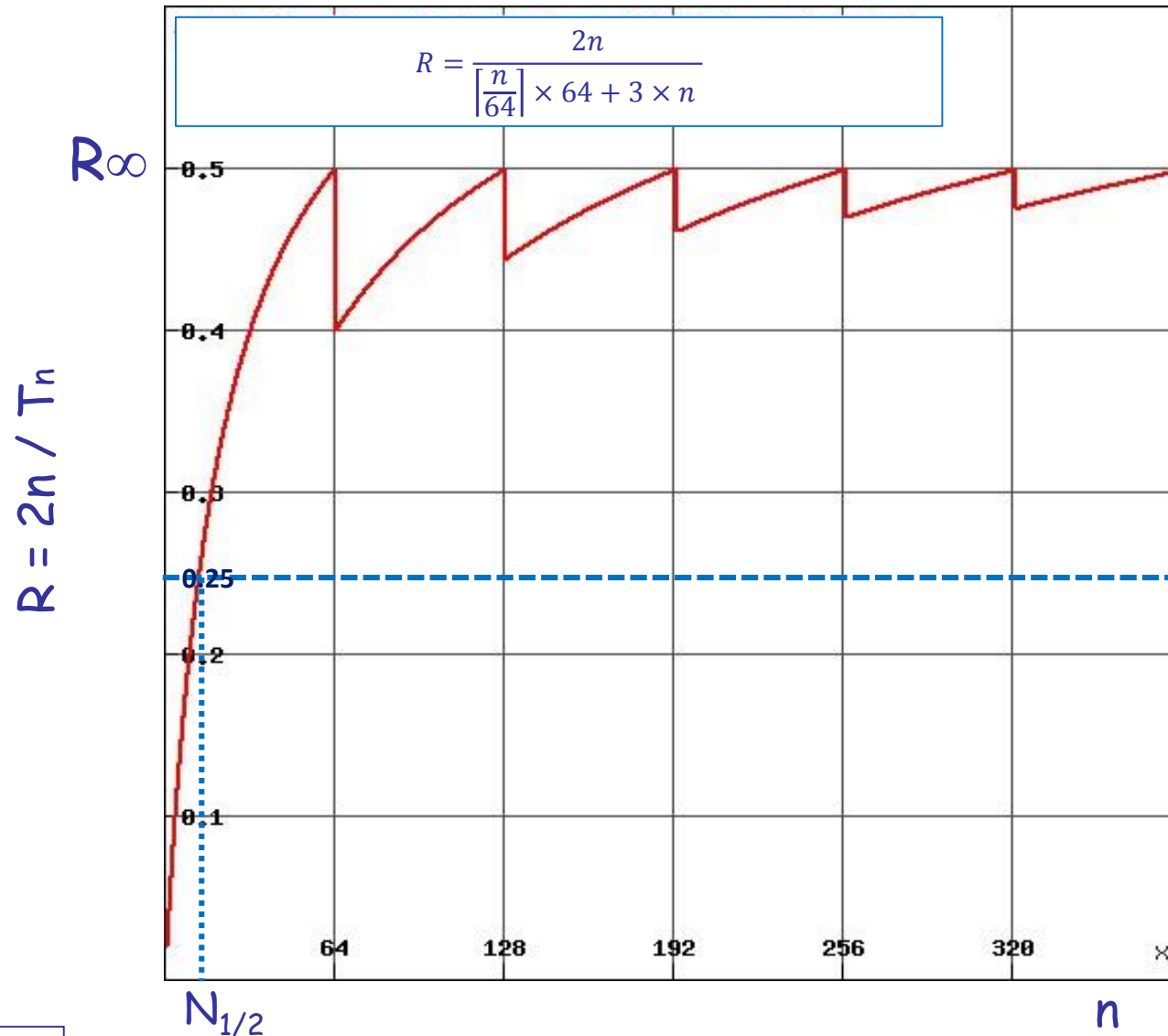
$$\text{En el ejemplo } \frac{R_{\infty}}{2} = \frac{1 \text{ FLOP}}{4 \text{ Ciclo}}$$

$$\text{Por def } R = \frac{2n}{T_n}; \text{ Sustituyendo: } \frac{1}{4} = \frac{2n}{64 + 3n}; n = 12.8; N_{1/2} = 13$$

Es decir, con vectores de sólo 13 componentes ya se obtiene un rendimiento que es la mitad del rendimiento asintótico

Medidas de rendimiento

Ejemplo: Rendimiento obtenido (FLOP/ciclo) en función de n



Operaciones condicionales: registro de máscara

❑ Ejecutar

```
for (i = 0; i < 64; i=i+1)
    if (X[i] != 0)
        X[i] = X[i] - Y[i];
```

❑ El registro de máscara vectorial permite omitir las operaciones sobre los elementos que no cumplen la condición

LV	V1,Rx	;load vector X into V1
LV	V2,Ry	;load vector Y
L.D	F0,#0	;load FP zero into F0
SNEVS.D	V1,F0	; <u>sets VM(i) to 1 if V1(i)!=F0</u>
SUBVV.D	V1,V1,V2	;subtract <u>under vector mask</u>
SV	Rx,V1	;store the result in X

❑ Obviamente, el rendimiento se reduce

- o Se consume el mismo tiempo
- o Se ejecutan menos operaciones útiles

- ❑ El sistema de memoria debe soportar un elevado ancho de banda en la carga y almacenamiento de vectores
- ❑ Idea fundamental: Dispersar los accesos entre múltiples bancos
 - o Control independiente de las direcciones en cada banco
 - o Load y store de palabras no secuenciales
 - o Soporte para que múltiples procesadores vectoriales puedan compartir la misma memoria
- ❑ Ejemplo:
 - o 32 procesadores, cada uno generando 4 loads y 2 stores por ciclo
 - o Tiempo de ciclo del procesador: 2.167 ns, tiempo de ciclo de la SRAM: 15 ns
 - o ¿Cuántos bancos de memoria serían necesarios?

❑ Ejemplo: Producto de matrices

```
for (i = 0; i < 100; i=i+1)
  for (j = 0; j < 100; j=j+1) {
    A[i][j] = 0.0;
    for (k = 0; k < 100; k=k+1)
      A[i][j] = A[i][j] + B[i][k] * D[k][j];
  }
```

❑ Las matrices están ordenadas por filas

- o Para vectorizar el producto de filas de B por columnas de D
 - Acceso a elementos consecutivos de B
 - Acceso a elementos de D separados por 100 palabras de distancia

❑ Soporte: Instrucciones de load/store con "stride" (espaciamiento)

- o Problema: Repetición de acceso al mismo banco de memoria antes del fin de la op anterior
- o Acceso libre de conflicto si:
 - $\text{mcm}(\text{espaciamiento}, n^\circ \text{ bancos}) / \text{espaciamiento} \geq T \text{ acceso a banco}$

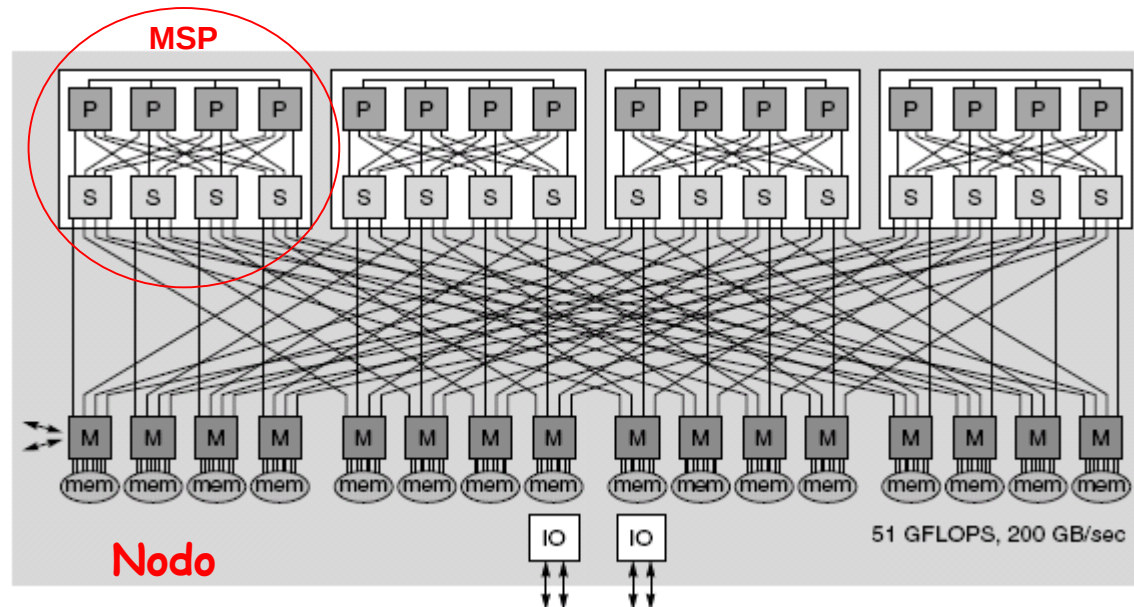
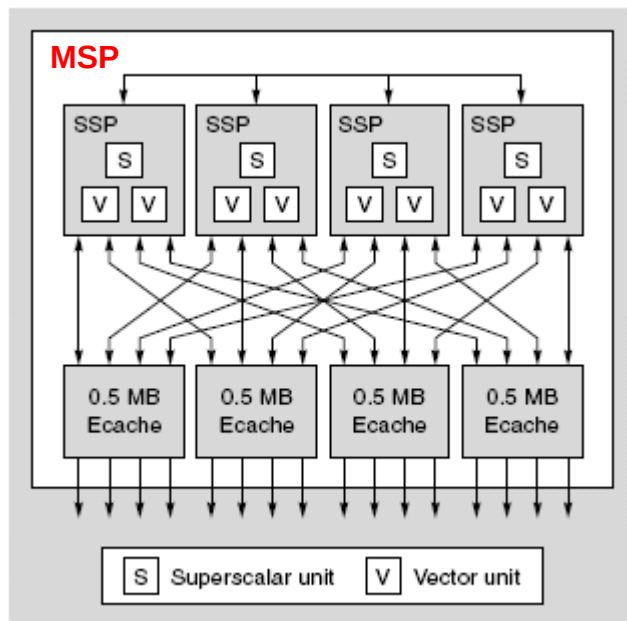
Vectores dispersos

- ❑ Los elementos no equidistan: utilización de vectores de índices
- ❑ **Carga.** Registro vectorial ← elementos dispersos de un array (compresión - *gather*): $LVI \ V1, (R1+V2)$
 - o V2 indica las direcciones de los elementos a cargar en V1 como desplazamientos respecto de una dir inicial indicada por R1
- ❑ **Almacenamiento.** Llevar contenido de un registro vectorial a posiciones dispersas de un array (expansión - *scatter*): $SVI \ (R1+V2), V1$
- ❑ **Ejemplo**
 - o Código fuente: sólo se accede a algunos elementos de A y C
for (i = 0; i < n; i=i+1)
 $A[K[i]] = A[K[i]] + C[M[i]];$
 - o Código máquina: V_k y V_m usados como índices

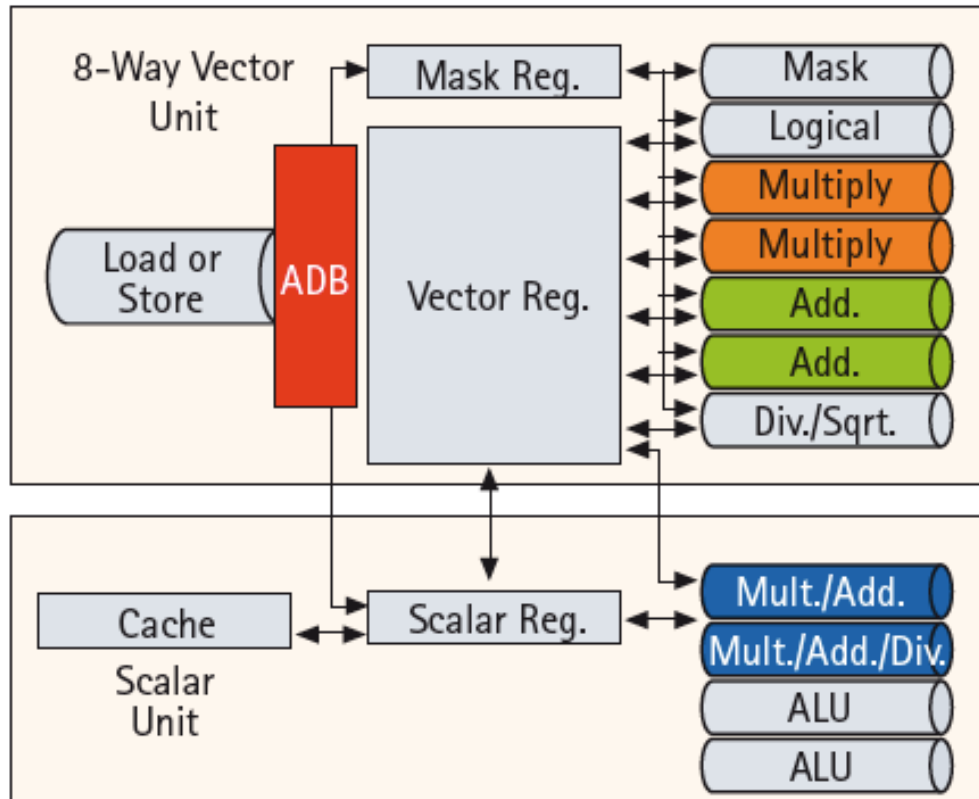
LV	V _k , R _k	;load K
LVI	V _a , (R _a +V _k)	;load A[K[]]
LV	V _m , R _m	;load M
LVI	V _c , (R _c +V _m)	;load C[M[]]
ADDVV.D	V _a , V _a , V _c	;add them
SVI	(R _a +V _k), V _a	;store A[K[]]

Ejemplo de Supercomputador Vectorial: Cray X1

- ❑ Introducido en 2002
- ❑ Elemento básico
 - o Single-Streaming Processor (SSP): μP vectorial con 2 vías + O-o-O procesador superescalar. Cada vía: 1 suma + 1 prod FP por ciclo.
- ❑ MSP: formado por 4 SSP. Rend pico: 12.8 GFLOPS
- ❑ Ecache: AB 1 palabra por vía por ciclo a 800 MHz (sobre 50 GB/s)
- ❑ Nodo: formado por 4 MSP, 16 controladores de memoria, DRAM (AB 200 GB/s)
- ❑ Sistema: hasta 1024 nodos. Red global de alta velocidad. Un único espacio de direcciones.



Ejemplo de Supercomputador Vectorial: NEC SX-9

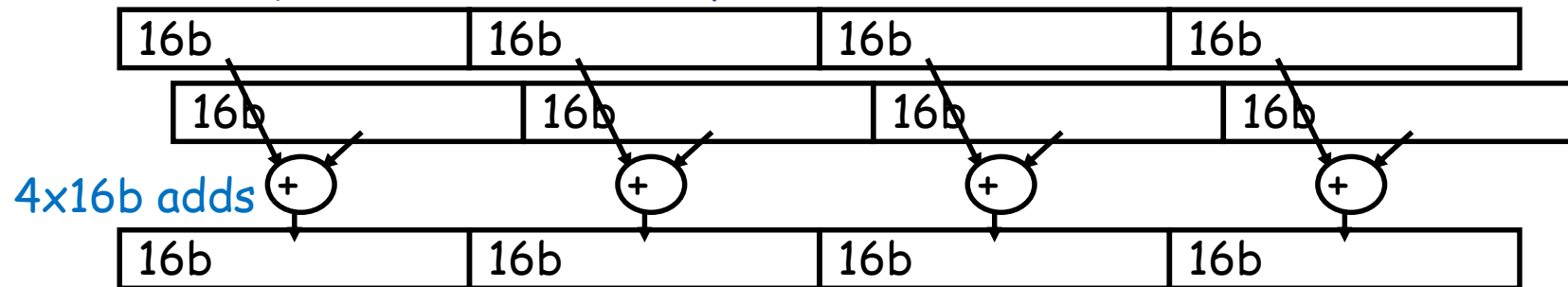


- ❑ Introducido en 2008
- ❑ Tecnología 65nm CMOS
- ❑ Unidad Vectorial (3.2 GHz)
 - o 8 foreground VRegs + 64 background VRegs (256x64-bit elementos/VReg)
 - o UFs de 64-bits: 2 multiply, 2 add, 1 divide/sqrt, 1 logical, 1 mask unit
 - o 8 vías (32+ FLOPS/cycle, 100+ GFLOPS pico por CPU)
 - o 1 load or store unit (8 x 8-byte accesos/ciclo)
- ❑ Unidad escalar (1.6 GHz)
 - o Superescalar 4 vías O-o-O
 - o 64KB I-cache, 64KB data cache

- ❑ AB Memoria: 256GB/s por CPU
 - o Hasta 16 CPUs y hasta 1TB de DRAM forman un **nodo** con mem compartida
 - o AB total: 4TB/s a la DRAM compartida
- ❑ Hasta 512 nodos conectados mediante enlaces de 128 GB/s (Paso de mensajes entre nodos)

Extensiones SIMD

- ❑ También conocidas como "extensiones multimedia"
- ❑ Observación: las aplicaciones multimedia suelen operar sobre datos de menor anchura que las UFs y los registros disponibles
- ❑ Idea: Realizar varias operaciones a la vez
 - o Por ejemplo: desconectar la cadena de propagación de carries
 - o Una sola instrucción de suma opera sobre varios elementos almacenados en un registro → equivale a una op vectorial, aunque con un vector de pocos elementos



- ❑ Limitaciones, comparado con op vectoriales:
 - o La longitud de los vectores se codifica en el Cod_op
 - o No existe direccionamiento sofisticado (stride, gather,...)
 - o No hay reg de máscara

❑ Implementaciones:

- o Intel MMX (1996)
 - Ocho ops enteras de 8-bit o cuatro ops enteras de 16-bit
- o Streaming SIMD Extensions (SSE) (1999-2007)
 - Ocho ops enteras de 16-bit
 - Cuatro ops enteras/FP de 32-bit o dos ops enteras/FP de 64-bit
- o Advanced Vector Extensions (AVX)(2010)
 - Cuatro ops enteras/FP de 64-bit
- o Los operandos deben ser consecutivos y en posiciones de memoria alineadas

Extensiones SIMD

□ Ejemplo: Instrucciones AVX para la arquitectura x86

4 operandos de 64 bits

AVX Instruction	Description
VADDPD	Add four packed double-precision operands
VSUBPD	Subtract four packed double-precision operands
VMULPD	Multiply four packed double-precision operands
VDIVPD	Divide four packed double-precision operands
VFMADDPD	Multiply and add four packed double-precision operands
VFMSUBPD	Multiply and subtract four packed double-precision operands
VCMPxx	Compare four packed double-precision operands for EQ, NEQ, LT, LE, GT, GE, ..
VMOVAPD	Move aligned four packed double-precision operands
VBROADCASTSD	Broadcast one double-precision operand to four locations in a 256-bit register

Como la longitud de los operandos va indicada en el Cod_op, puede dar la impresión de que el n° de instr de las "extensiones multimedia" es mayor de lo que en realidad es.

Ejemplo: código con extensiones SIMD en MIPS

- ❑ Sup: Añadimos instrucciones multimedia SIMD de 256 bits al MIPS (".4D" → op sobre 4 operandos de 64 bits a la vez)
- ❑ Código para DAXPY:

	L.D F0,a	;load scalar a
	MOV F1, F0	;copy a into F1 for SIMD MUL
	MOV F2, F0	;copy a into F2 for SIMD MUL
	MOV F3, F0	;copy a into F3 for SIMD MUL
	DADDIU R4,Rx,#512	;last address to load
Loop:	L.4D F4,0(Rx)	;load X[i], X[i+1], X[i+2], X[i+3]
	MUL.4D F4,F4,F0	;a×X[i],a×X[i+1],a×X[i+2],a×X[i+3]
	L.4D F8,0(Ry)	;load Y[i], Y[i+1], Y[i+2], Y[i+3]
	ADD.4D F8,F8,F4	;a×X[i]+Y[i], ..., a×X[i+3]+Y[i+3]
	S.4D F8,0(Rx)	;store into Y[i], Y[i+1], Y[i+2], Y[i+3]
	DADDIU Rx,Rx,#32	;increment index to X
	DADDIU Ry,Ry,#32	;increment index to Y
	DSUBU R20,R4,Rx	;compute bound
	BNEZ R20,Loop	;check if done

Unidades para Procesamiento Gráfico (GPUs)

- ❑ Las GPUs son económicas, accesibles y contienen una gran cantidad de elementos de cómputo.
- ❑ Se han concebido con el objeto de realizar los procesamiento característicos de las aplicaciones gráficas
- ❑ ¿Cómo poder utilizar la gran potencia de los procesadores gráficos en un espectro de aplicaciones más amplio?
- ❑ Idea básica
 - o Modelo de ejecución heterogéneo (CPU+GPU)
 - o Desarrollar un lenguaje de programación tipo C que permita programar la GPU
 - o Unificar todo el paralelismo de la GPU bajo la abstracción denominada "CUDA Thread"
 - o Modelo de Programación: "Single Instruction Multiple Thread"

□ GPU NVIDIA

- o Multiprocesador compuesto por un conjunto de procesadores SIMD MT

□ NVIDIA vs procesadores vectoriales

o Similaridades

- Funciona bien en problemas con paralelismo de datos
- Transferencias con memoria tipo dispersar/reunir (scatter/gather)
- Registros de máscara
- Existencia de grandes ficheros de registros

o Diferencias

- No hay un procesador escalar
- Utilización de multithreading para ocultar la latencia de memoria
- Existencia de gran cantidad de UFs.
 - Contrasta con la reducida cantidad de UFs muy segmentadas, que es típica de los procesadores vectoriales

CUDA(Compute Unified Device Architecture)

- ❑ CUDA is an elegant solution to the problem of representing parallelism in algorithms, not all algorithms, but enough to matter. It seems to resonate in some way with the way we think and code, allowing an easier, more natural expression of parallelism beyond the task level.

Vincent Natoli

"Kudos for CUDA", *HPC Wire* (July 2010)

http://www.hpcwire.com/hpcwire/2010-07-06/kudos_for_cuda.html

- ❑ CUDA produce código C/C++ para host y dialecto de C y C++ para la GPU
 - o Idea básica: crear un **thread** (hilo) separado para cada elemento de los vectores a procesar
 - Objetivo: generar un gran nº de hilos de cómputo independientes
 - o Los threads se agrupan en **bloques de threads**
 - El número de threads por bloque puede definirlo el programador
 - Cada bloque es ejecutado por un procesador SIMD MT de la GPU
 - Varios bloques pueden ejecutarse en paralelo sobre varios procesadores
 - o El conjunto de bloques que implementan un cálculo vectorial sobre la GPU se denomina **Grid** (malla). La ejecución del cálculo se produce con una llamada similar a una función en C:
 - **nombre_función <<<dimGrid,dimBlock>>>** (... lista de parámetros ...)
 - **dimGrid**: nº de bloques en el Grid
 - **dimBlock**: nº de threads por bloque
 - Los bloques y las mallas pueden tener hasta 3 dimensiones, que se identifican con .x , .y, .z.

□ CUDA vs C. Ejemplo DAXPY

o Versión C

```
// Invocar DAXPY
```

```
daxpy(n, 2.0, x, y);
```

```
// DAXPY en C (bucle escalar: una iteración por elemento)
```

```
void daxpy(int n, double a, double *x, double *y)
```

```
{
```

```
    for (int i = 0; i < n; i++)
```

```
        y[i] = a*x[i] + y[i];
```

```
}
```

❑ CUDA vs C. Ejemplo DAXPY

o Versión CUDA

// Invocar DAXPY con 256 threads por Bloque (dimBlock)

__host__ /* código para la CPU */

int nblocks = (n+ 255) / 256; /* cálculo del nº total de bloques en el Grid (dimGrid) */

daxpy <<<nblocks, 256>>> (n, 2.0, x, y);

// DAXPY en CUDA (representa el cálculo ejecutado para un elemento)

__device__ /* código para los procesadores de la GPU */

void daxpy(int n, double a, double *x, double *y)

{

// ¿Qué thread soy? Calcular i = nº elemento del vector a procesar (= nº de thread), siendo

// nº elemento = (nº de bloque x tamaño de bloque) + (nº de thread dentro del bloque)

int i = blockIdx.x*blockDim.x + threadIdx.x;

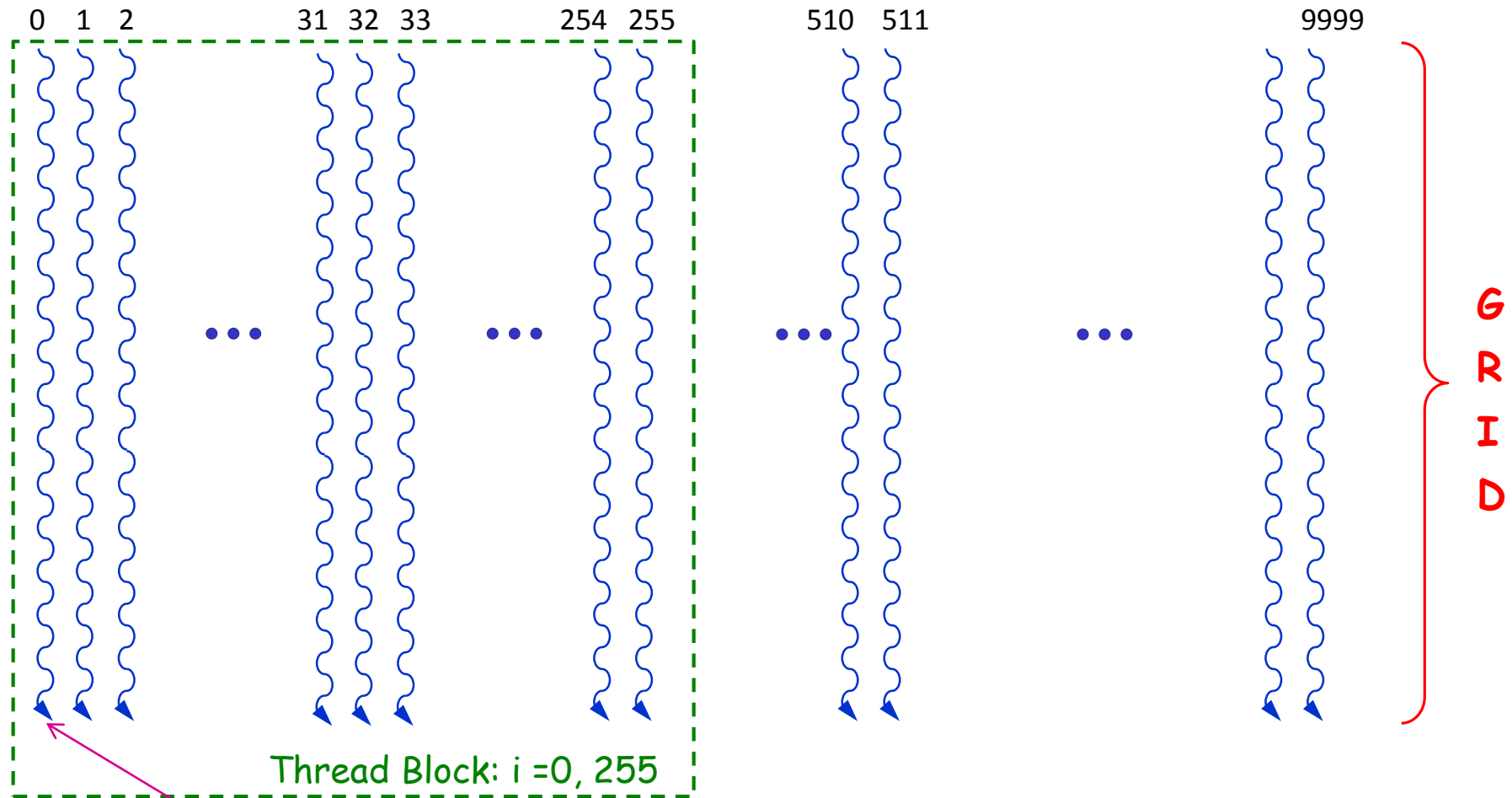
// Si el nº elemento obtenido es mayor que el tamaño del vector, ignorar operación

if (i < n) y[i] = a*x[i] + y[i];

}

CUDA(Compute Unified Device Architecture)

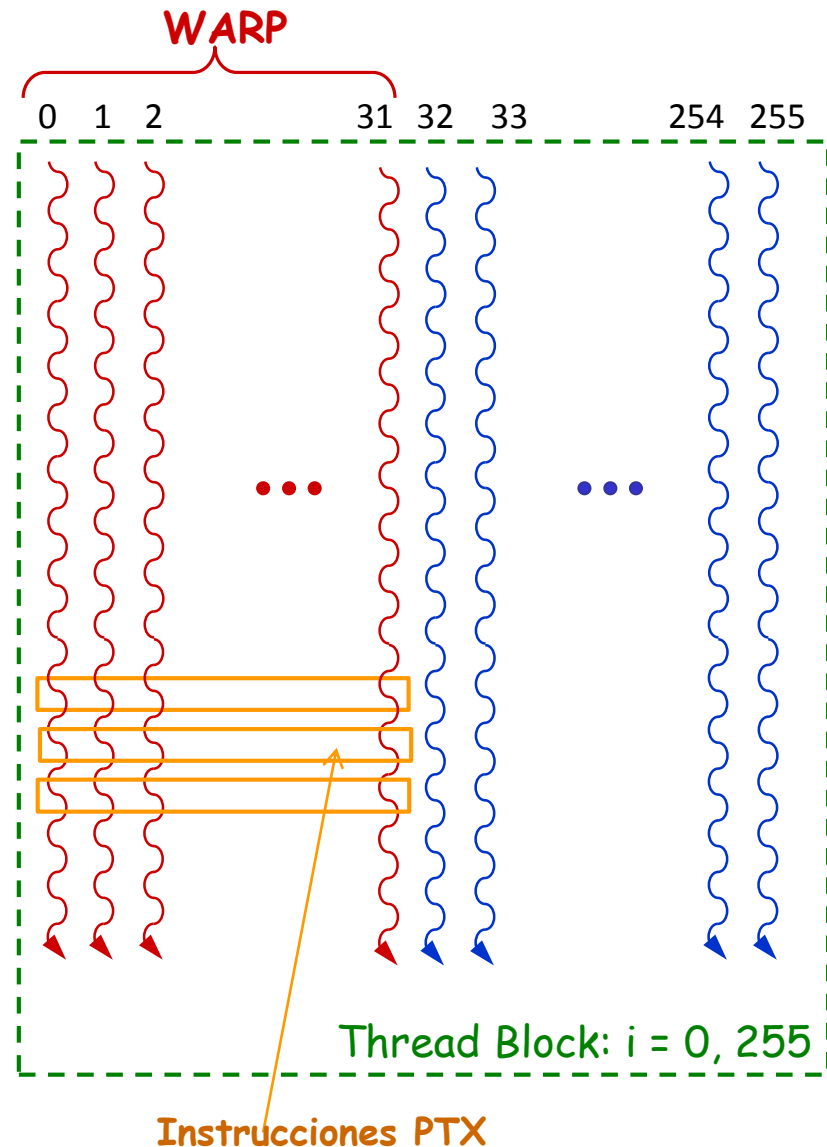
□ Ejemplo: vectores de 10,000 componentes



CUDA thread. Ejemplo: Calcula $Y(0) = a * X(0) + Y(0)$

Thread blocks e Instrucciones PTX

- ❑ Cada Thread Block se ejecuta en un procesador SIMD MT de la GPU.
- ❑ Varios procesadores SIMD MT de la GPU pueden procesar diferentes Thread Blocks en paralelo.
- ❑ Instrucción PTX: Ejecuta un mismo cálculo sobre varios (e.g. 32) datos (Instr SIMD).
 - o Resultados afectados por registro de máscara.
- ❑ WARP: Secuencia (thread) de instr PTX. El procesador ejecuta los WARP de un Thread Block en modo MT.
 - o En el ejemplo hay $256/32 = 8$ WARPs.
 - o Cambios de thread ocultan latencias de acceso a memoria



Código generado por compiladores de NVIDIA

- ❑ Instrucciones PTX (Parallel Thread Execution)
 - o Abstracción del repertorio de instrucciones hw
 - o Formato: `opcode.type dest, src1, src2, src3`
 - o Instrucción que ejecuta una operación elemental sobre múltiples datos (SIMD) utilizando todas las vías del procesador
 - o Ejemplo: un conjunto de instrucciones PTX representativas

Instruction	Example	Meaning	Comments
<code>arithmetic .type = .s32, .u32, .f32, .s64, .u64, .f64</code>			
<code>add.type</code>	<code>add.f32 d, a, b</code>	<code>d = a + b;</code>	
<code>sub.type</code>	<code>sub.f32 d, a, b</code>	<code>d = a - b;</code>	
<code>mul.type</code>	<code>mul.f32 d, a, b</code>	<code>d = a * b;</code>	
<code>mad.type</code>	<code>mad.f32 d, a, b, c</code>	<code>d = a * b + c;</code>	multiply-add
<code>div.type</code>	<code>div.f32 d, a, b</code>	<code>d = a / b;</code>	multiple microinstructions
<code>setp.cmp.type</code>	<code>setp.lt.f32 p, a, b</code>	<code>p = (a < b);</code>	compare and set predicate
<code>numeric .cmp = eq, ne, lt, le, gt, ge; unordered cmp = equ, neu, ltu, leu, gtu, geu, num, nan</code>			
<code>mov.type</code>	<code>mov.b32 d, a</code>	<code>d = a;</code>	move
<code>selp.type</code>	<code>selp.f32 d, a, b, p</code>	<code>d = p? a: b;</code>	select with predicate
<code>memory.space = .global, .shared, .local, .const; .type = .b8, .u8, .s8, .b16, .b32, .b64</code>			
<code>ld.space.type</code>	<code>ld.global.b32 d, [a+off]</code>	<code>d = *(a+off);</code>	load from memory space
<code>st.space.type</code>	<code>st.shared.b32 [d+off], a</code>	<code>*(d+off) = a;</code>	store to memory space

- ❑ Ejemplo: secuencia de instrucciones PTX para una iteración del bucle DAXPY
 - o Usa reg virtuales: Ri (32 bits), RDi (64 bits)
 - o Asigna reg físicos en el momento de la carga del programa

shl.u32	R8, blockIdx, 8	; Thread Block ID * Block size (256 or 2 ⁸)
add.u32	R8, R8, threadIdx	; R8 = i = my CUDA Thread ID
shl.u32	R8, R8, 3	; byte offset
ld.global.f64	RD0, [X+R8]	; RD0 = X[i]
ld.global.f64	RD2, [Y+R8]	; RD2 = Y[i]
mul.f64	RD0, RD0, RD4	; Product in RD0 = RD0 * RD4 (scalar a)
add.f64	RD0, RD0, RD2	; Sum in RD0 = RD0 + RD2 (Y[i])
st.global.f64	[Y+R8], RD0	; Y[i] = sum (X[i]*a + Y[i])

Ojo! Recordar que cada instrucción PTX procesa 32 elementos

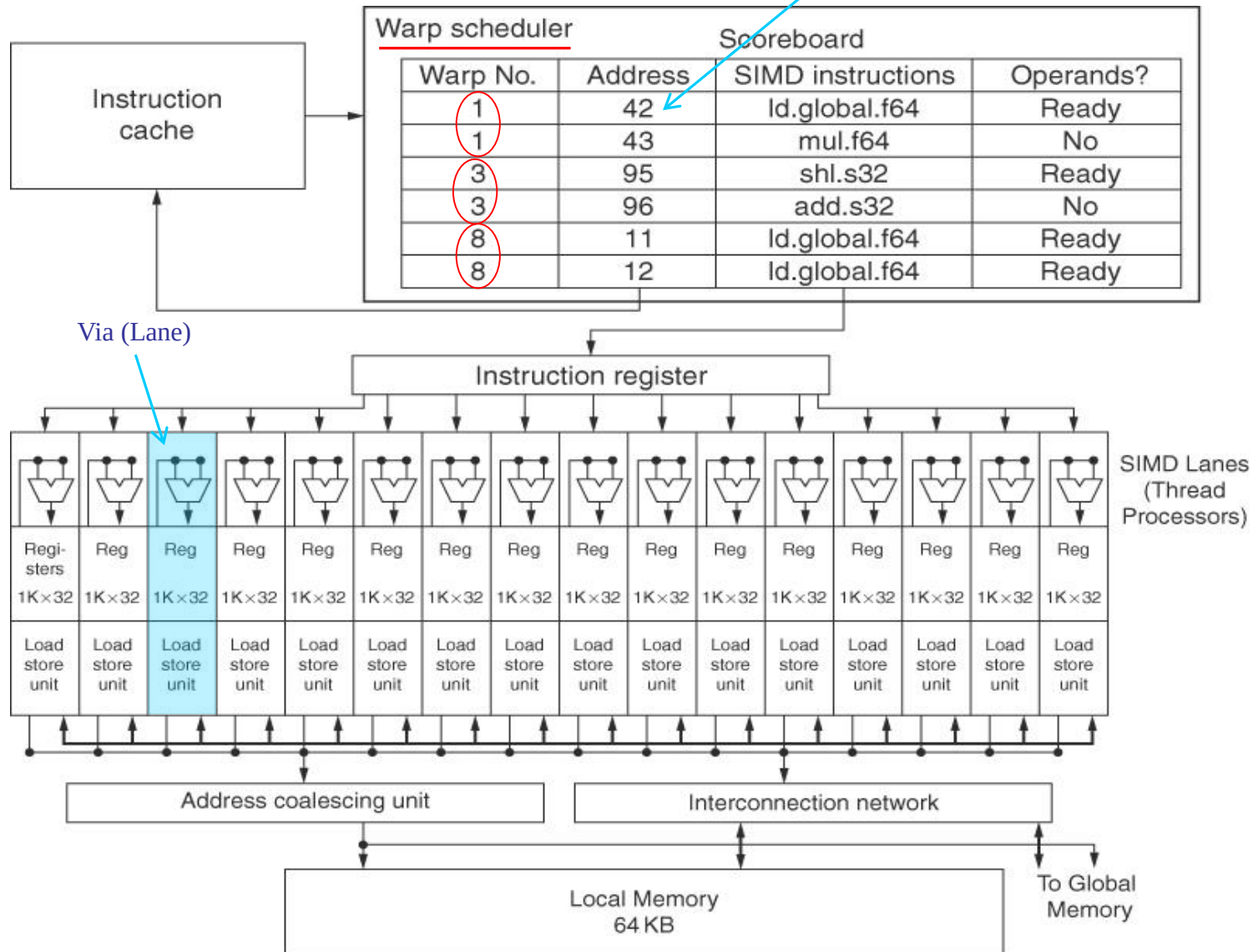
Resumen de terminología: arquitectura vectorial vs. GPUs

Type	More descriptive name	Closest old term outside of GPUs	Official CUDA/ NVIDIA GPU term	Book definition
Program abstractions	Vectorizable Loop	Vectorizable Loop	Grid (Malla) Ej. $i = 0..9999$	A vectorizable loop, executed on the GPU, made up of one or more Thread Blocks (bodies of vectorized loop) that can execute in parallel.
	Body of Vectorized Loop	Body of a (Strip-Mined) Vectorized Loop	Thread Block Ej. $i = 0..255$	A vectorized loop <u>executed on a multithreaded SIMD Processor</u> , made up of one or more threads of SIMD instructions. They can communicate via Local Memory.
	Sequence of SIMD Lane Operations	One iteration of a Scalar Loop	CUDA Thread Ej. $i = 12$	A vertical cut of a thread of SIMD instructions <u>corresponding to one element executed by one SIMD Lane</u> . Result is stored depending on mask and predicate register.
Machine object	A Thread of SIMD Instructions	Thread of Vector Instructions	Warp (Trama)	A <u>traditional thread, but it contains just SIMD instructions</u> that are executed on a multithreaded SIMD Processor. Results stored depending on a per-element mask.
	SIMD Instruction	Vector Instruction	PTX Instruction Ej. $i = 0..31$	A single SIMD instruction executed across SIMD Lanes.

Procesadores de una GPU

❑ Procesador SIMD MT

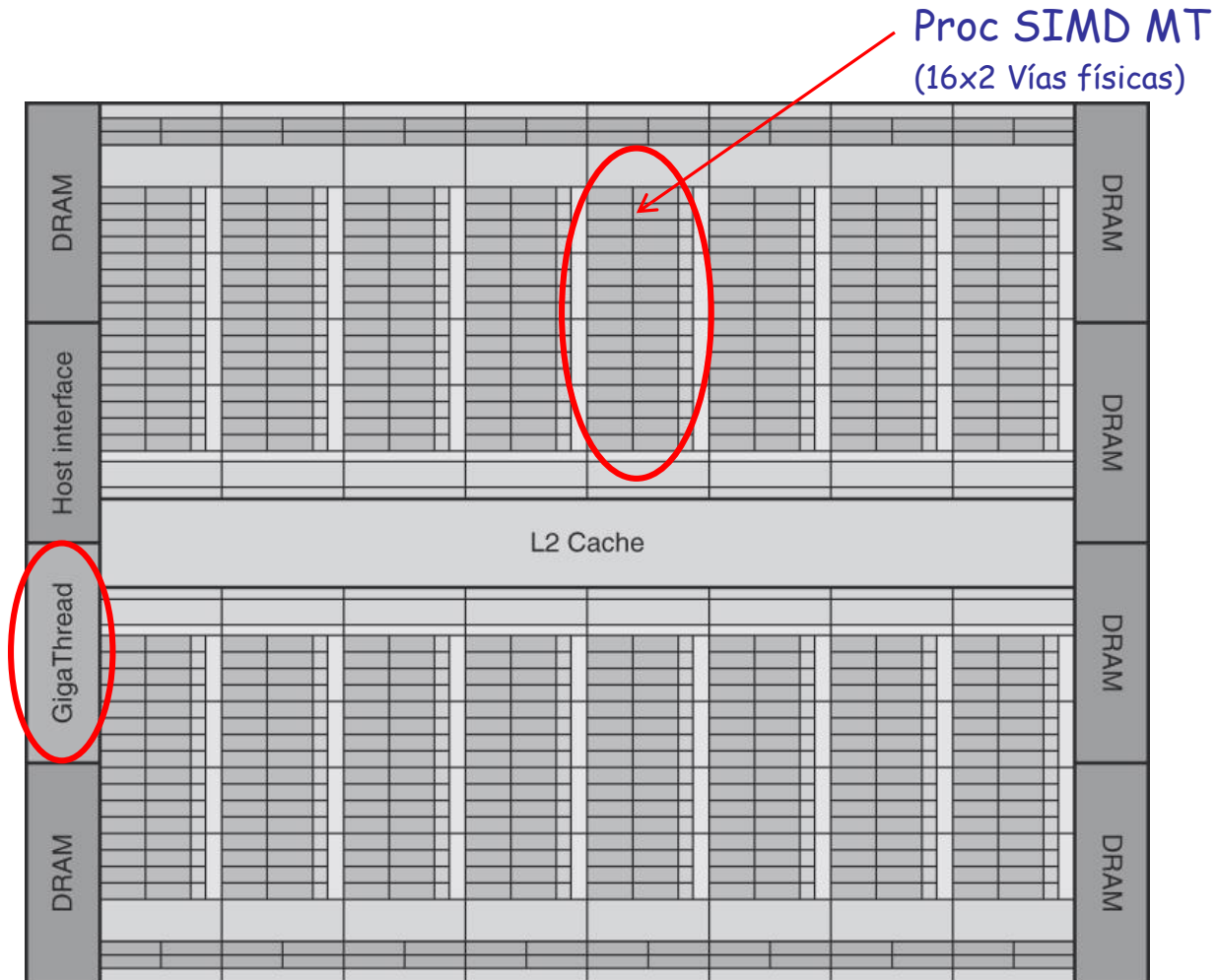
Cada Warp (th. de instrucciones SIMD) tiene su PC



- Todas las vías ejecutan la misma instrucción
- Según la máscara, unas guardan el resultado y otras no

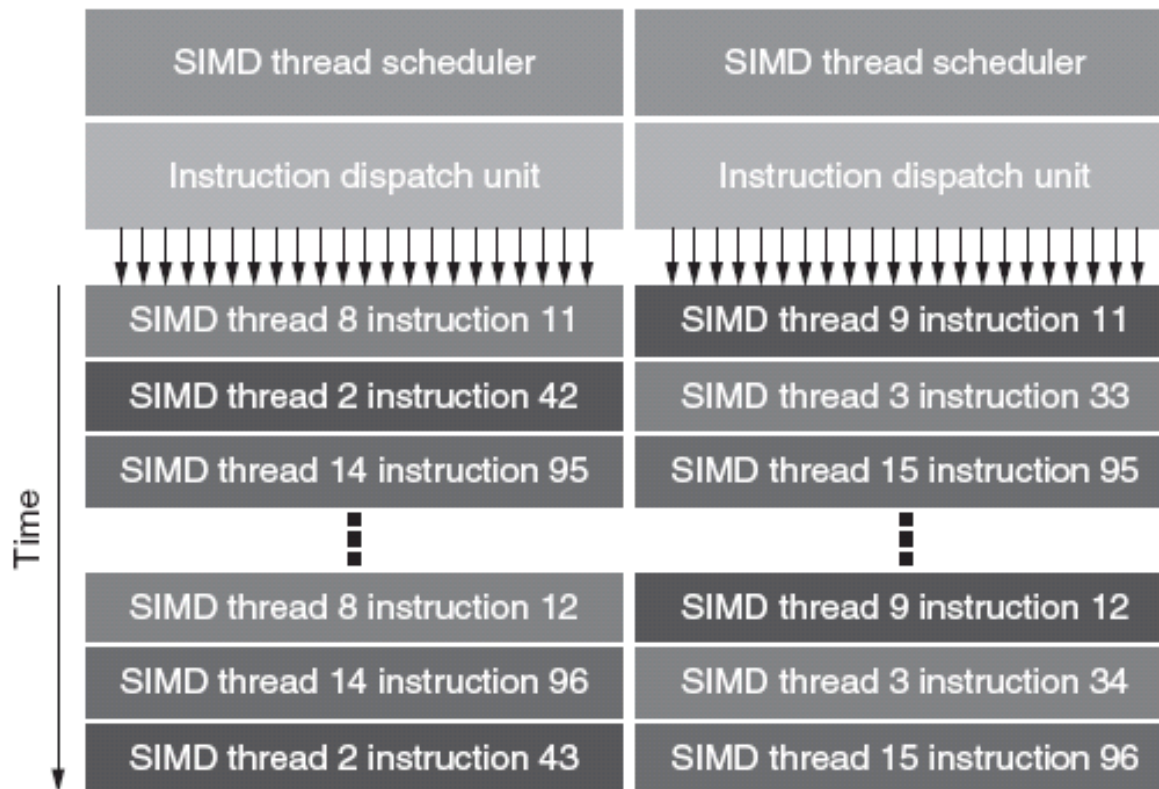
Procesadores de una GPU

- ❑ Ejemplo: Fermi GTX 480 (16 procesadores SIMD MT)
 - o GigaThread: Distribuye bloques a procesadores
 - o Hasta 6 GB de memoria



❑ Fermi: Thread Scheduler

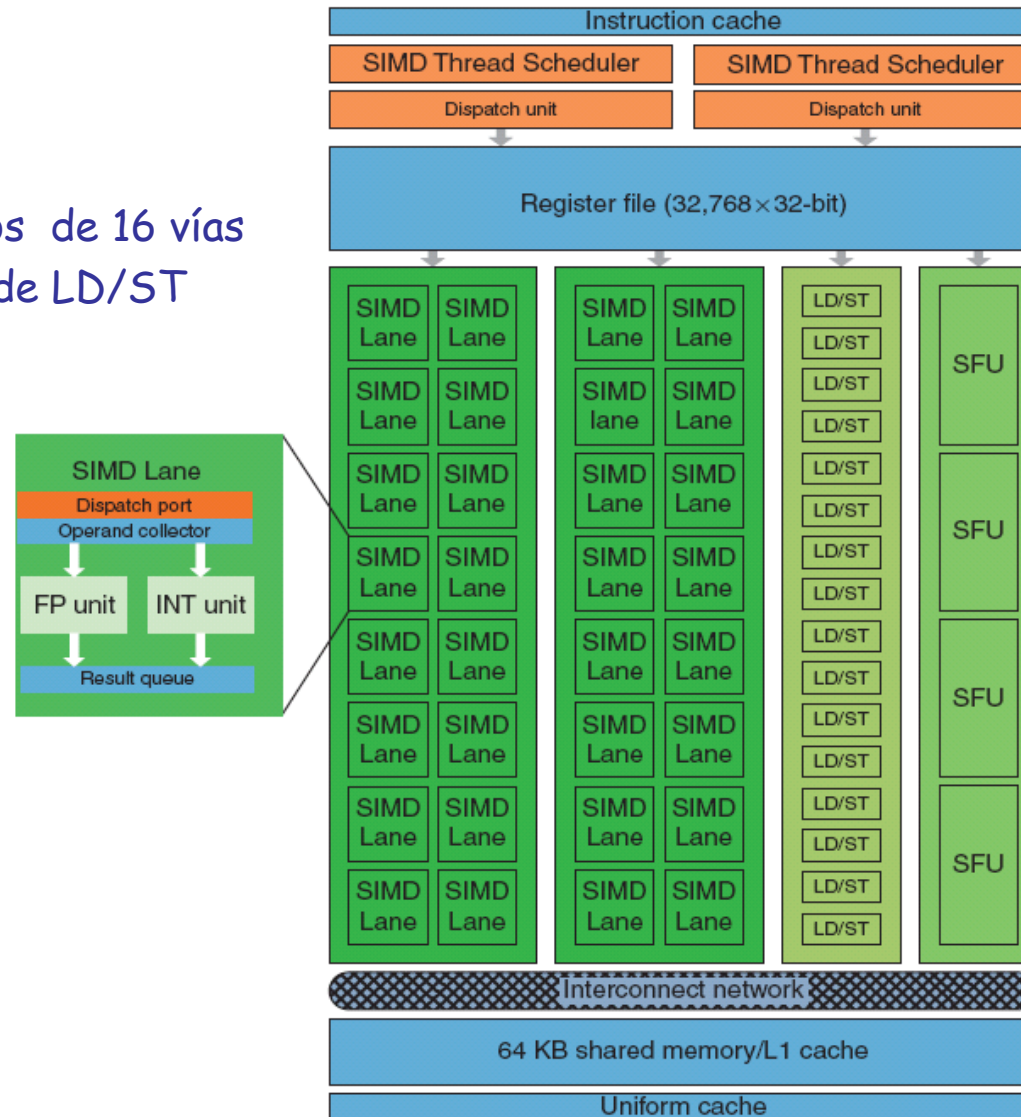
- o Cada procesador: 2 thread schedulers en paralelo
- o Selecciona una instrucción de cada thread y las envía a dos conjuntos de 16 vías físicas
- o Cada instrucción SIMD procesa 32 elementos (y se ejecuta sobre 16 vías físicas) → 2 ciclos por instrucción



Procesadores de una GPU

□ Fermi: Esquema de un procesador SIMD MT

- Dos conjuntos de 16 vías
- 16 unidades de LD/ST



- ❑ Gestión de registros de máscara (predicado) similar a los procesadores vectoriales.
 - o Para componentes enmascaradas, el resultado no se guarda en el registro destino
 - o Permite implementar construcciones IF...THEN mediante una instrucción PTX "compare and set predicate" (setp)
 - o Construcciones IF... THEN...ELSE: mecanismo similar, pero para la parte ELSE el registro de máscara se complementa.
 - o Impacto en rendimiento
- ❑ Instrucciones PTX de salto: Permiten implementar construcciones condicionales anidadas
 - o Formato: **@p branch target**
 - Si todos los bits de máscara (registros p) del procesador están a 1, el procesador pasa a ejecutar la instrucción con etiqueta "target"
 - o Además para preservar la máscara existente antes de entrar en un IF...THEN...ELSE, existe un stack de máscaras.
 - Apilar máscara (push) antes de entrar al IF...THEN...ELSE, desapilar (pop) al salir. Complementar (comp) máscara actual al entrar en parte ELSE
 - Marcadores de sincronización: *push, *pop, *comp
 - o El flujo de programa externo al IF...THEN...ELSE no continua hasta que todas los threads han finalizado

Saltos condicionales en GPUs

□ Ejemplo

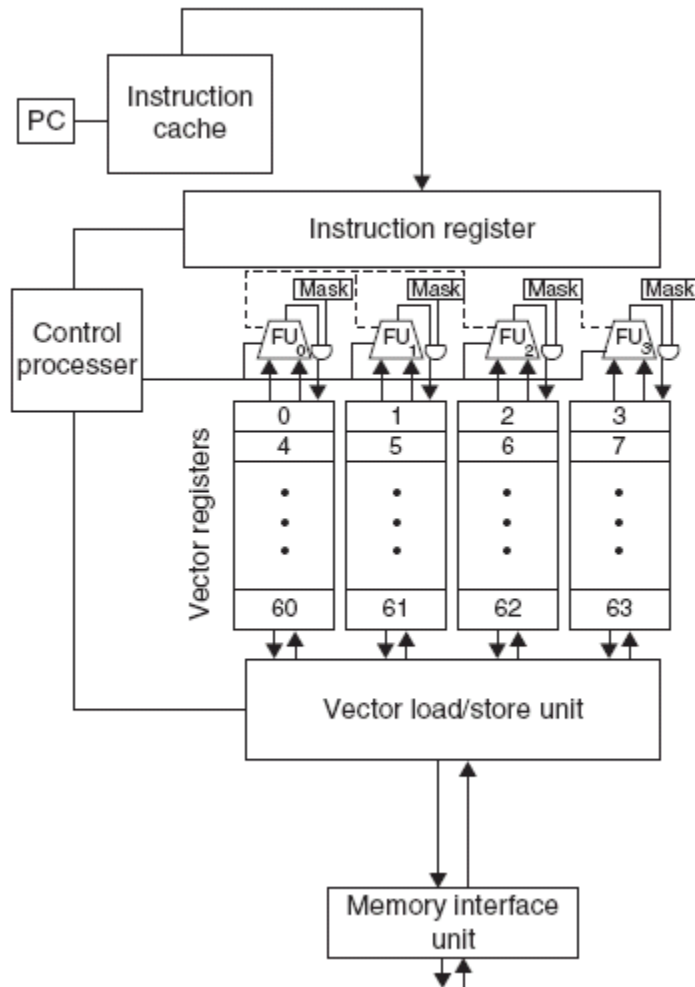
```
if (X[i] != 0)
    X[i] = X[i] - Y[i];
else X[i] = Z[i];
```

o Código PTX

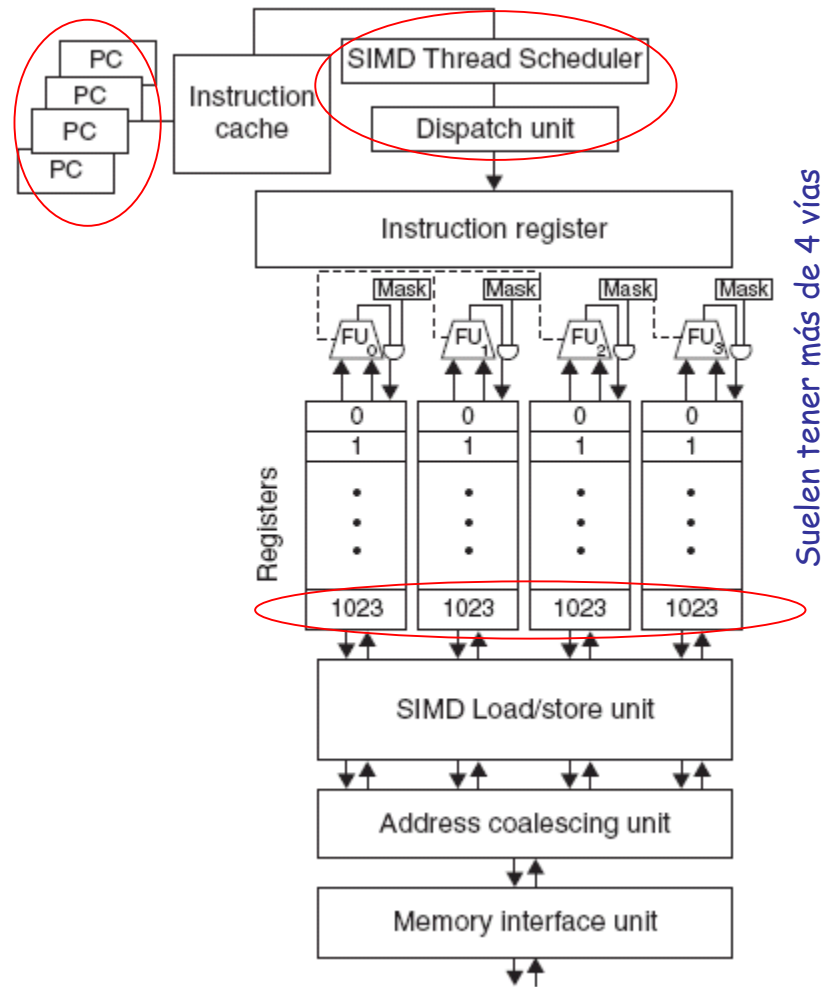
	ld.global.f64	RD0, [X+R8]	; R8 actualizado de acuerdo con Thread Id
	setp.neq.s32	P1, RD0, #0	; RD0 = X[i]
	@!P1, bra	ELSE1, <i>*Push</i>	; P1 is predicate register 1
			; Push old mask, set new mask bits
			; if P1 false, go to ELSE1
	ld.global.f64	RD2, [Y+R8]	; RD2 = Y[i]
	sub.f64	RD0, RD0, RD2	; Difference in RD0
	st.global.f64	[X+R8], RD0	; X[i] = RD0
	@P1, bra	ENDIF1, <i>*Comp</i>	; complement mask bits
			; if P1 true, go to ENDIF1
ELSE1:	ld.global.f64	RD0, [Z+R8]	; RD0 = Z[i]
	st.global.f64	[X+R8], RD0	; X[i] = RD0
ENDIF1:	<next instruction>	, <i>*Pop</i>	; pop to restore old mask

Comparación Procesador Vectorial - GPU

Procesador vectorial
con cuatro vías



Procesador SIMD MT (4
PCs) con cuatro vías



Paralelismo a nivel de bucle: vectorización

- ❑ Las dependencias RAW entre sentencias de una misma iteración no impiden la vectorización eficiente (mediante encadenamiento de pipes)
 - o Dependencia **directa**: obliga a ejecutar en el orden dado
- ❑ ¿Qué ocurre si los datos de una iteración son dependientes de los resultados generados en iteraciones previas?
 - o Dependencia **en el espacio de iteraciones** (loop-carried)
 - o Puede impedir la vectorización
 - o Pueden existir reordenaciones de sentencias válidas

❑ Ejemplo 1

```
for (i=999; i>=0; i=i-1){  
    x[i] = x[i] + s;      /* S1 */  
    z[i] = z[i] + x[i];   /* S2 */  
}
```

Dep directa: no impide la vectorización. Orden: 1º S1, luego S2

```
x[1:999] = x[1:999] + s;  
z[1:999] = z[1:999] + x[1:999]
```

Paralelismo a nivel de bucle: vectorización

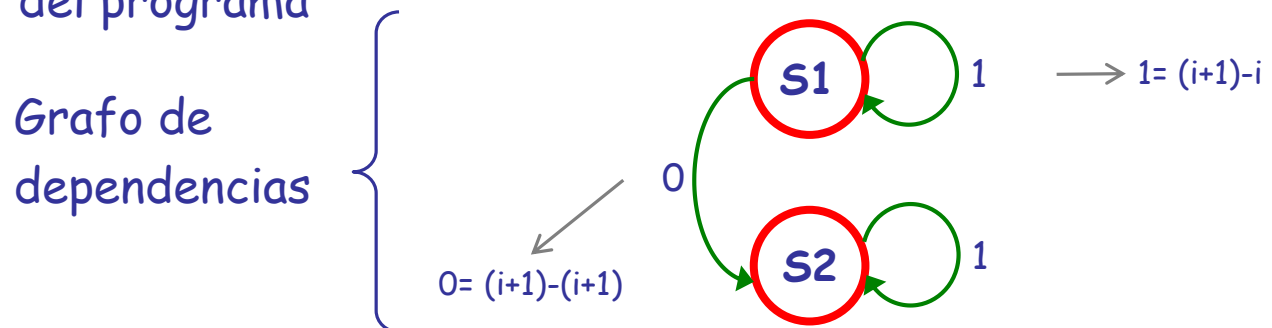
❑ Ejemplo 2:

```
for (i=0; i<100; i=i+1) {  
    A[i+1] = A[i] + C[i];          /* S1 */  
    B[i+1] = B[i] + A[i+1];       /* S2 */  
}
```

❑ S1 y S2 usan valores calculados por ellas mismas en la iteración previa → ejecución en serie

❑ S2 usa resultados de S1 en la misma iteración.

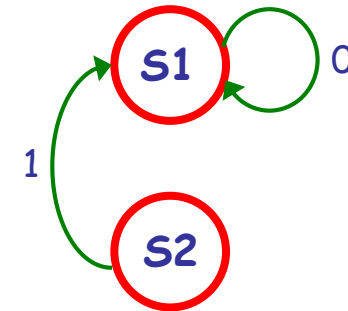
- o Si ésta fuera la única dependencia (no loop-carried), el bucle sería vectorizable
- o En todo caso, el orden de ejecución de S1 y S2 debe mantenerse. Si se cambian de orden se altera la semántica del programa



Paralelismo a nivel de bucle: vectorización

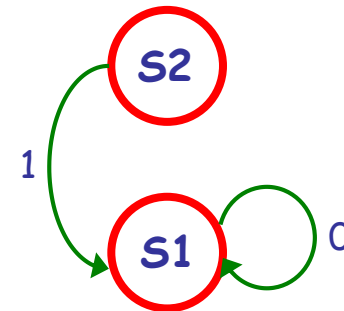
❑ Ejemplo 3

```
for (i=0; i<100; i=i+1) {  
    A[i] = A[i] + B[i];           /* S1 */  
    B[i+1] = C[i] + D[i];        /* S2 */  
}
```



- ❑ S1 usa un valor calculado por S2 en la iteración previa, pero la dependencia no es circular → Vectorizable...Cómo?
- ❑ Única flecha hacia arriba tiene $d > 0$ → Transformable a (orden de sentencias inverso):

```
for (i=0; i<100; i=i+1) {  
    B[i+1] = C[i] + D[i];        /* S2 */  
    A[i] = A[i] + B[i];          /* S1 */  
}
```



Detección de dependencias

- ❑ Sup: los índices de los bucles toman valores de acuerdo con una función afín
 - o $a*i + b$ (siendo i el índice)
- ❑ Sup:
 - o Almacenar en la posición $a*i + b$ de un vector. Después:
 - o Cargar desde la posición $c*i + d$ del mismo vector
 - o i toma valores desde m hasta n
- ❑ Existe una dependencia si:
 - o Dados j, k tales que $m \leq j \leq n, m \leq k \leq n$
 - o Almacenar en $a*j + b$, cargar desde $a * k + d$, y
$$a * j + b = c * k + d$$
- ❑ Test del MCD
 - o Si existe dep entonces $MCD(c,a)$ es un divisor de $(d-b)$
- ❑ Ejemplo 4:

```
for (i=0; i<100; i=i+1) {  
    X[2*i+3] = X[2*i] * 5.0;  
}
```

 - o $a=2, c=2, b=3, d=0$. $MCD(2,2)=2$. $d-b=-3 \rightarrow$ No dependencia

Antidependencias y dependencias de salida

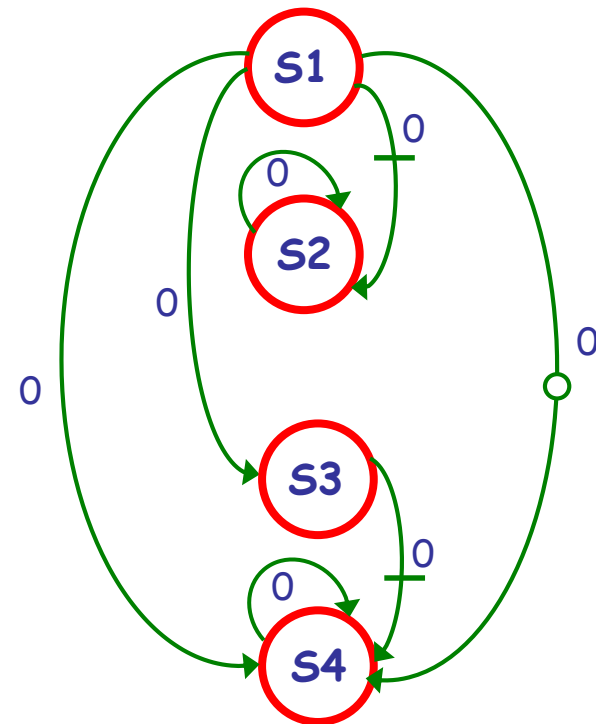
- ❑ Las dependencias de nombre pueden evitarse renombrando variables

- ❑ Ejemplo 5

```
for (i=0; i<100; i=i+1) {  
    Y[i] = X[i] / c;      /* S1 */  
    X[i] = X[i] + c;      /* S2 */  
    Z[i] = Y[i] + c;      /* S3 */  
    Y[i] = c - Y[i];      /* S4 */  
}
```

o Transformar a:

```
for (i=0; i<100; i=i+1) {  
    T[i] = X[i] / c;      /* Y renamed to T to remove output dependence */  
    X1[i] = X[i] + c;     /* X renamed to X1 to remove antidependence */  
    Z[i] = T[i] + c;      /* Y renamed to preserve true dependence */  
    Y[i] = c - T[i];  
}
```



- ❑ Ejemplo de operación de reducción :

```
for (i=9999; i>=0; i=i-1)
    sum = sum + x[i] * y[i]; /* no vectorizable */
```

- ❑ Transformar a...

```
for (i=9999; i>=0; i=i-1)
    sum [i] = x[i] * y[i]; /* vectorizable */
```

```
for (i=9999; i>=0; i=i-1)
    finalsum = finalsum + sum[i]; /* no vectorizable */
```

- ❑ Suma final se puede acelerar. Si tenemos 10 procesadores (p=0..9), procesar 1000 elementos en cada uno:

```
for (i=999; i>=0; i=i-1)
    finalsum[p] = finalsum[p] + sum[i+1000*p];
```

- ❑ Se está asumiendo que la operación es asociativa