

Lezione 12 - Standard IEEE-754 per le operazioni floating-point

Finora abbiamo usato solo estensione QV64M

• Nelle operazioni floating point dobbiamo rappresentare *offers support*

- numeri con decimali, e.g., 3.1416 *con virgola*
- numeri molto piccoli, e.g., .000000001
- numeri molto grandi, e.g., 3.15576×10^9

approssimo numero, convertire numeri

float non sono
numeri reali, appaiono
 $\approx$ *avvicinamento
non come
rappresentato in
memoria*

per F D₁₁
32 bit Floating point
float 64-bit double

• Rappresentazione nel calcolatore:

Scelta rappresentazione

- Partendo dall'idea di memorizzare tali numeri usando cifre binarie, scriveremo il numero floating point x da rappresentare nella forma:

$$x = (-1)^{\text{segno}} \times \text{mantissa} \times 2^{\text{esponente}}$$

*codifica delle quantità
fondamentali*

- dove viene scelto $1 \leq \text{'mantissa'} < 2$ *compresa fra 1,2, non dov'essere
prima cifra da 0 a 1 come 1*
- più cifre binarie (bit) di 'mantissa' danno più accuratezza nelle operazioni
- più cifre binarie (bit) di 'esponente' permettono numeri molto grandi e molto piccoli

vorrei un numero che mi salvi 3 caratteri. fondamentale

$\rightarrow$ più cifre mantissa, più accuratezza $\rightarrow$ più cifre esponente, numeri più grandi o piccoli

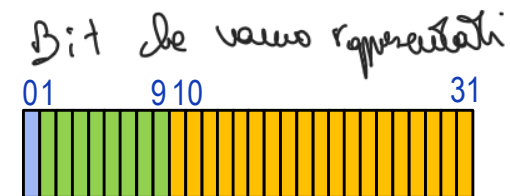
IEEE-754

- IEEE 754 floating point standard:
 - singola precisione ("float" del C) → uso 32 bit per memorizzare:
1 bit per S, 8 bit per E, 23 bit per M**
 - doppia precisione ("double" del C) → uso 64 bit per memorizzare:
1 bit per S, 11 bit per E, 52 bit per M**
 - quadrupla precisione ("__float128" del C) → uso 128 bit:
1 bit per S, 15 bit per E, 112 bit per M**

:= esize

:= msize

Esempio:



S E M

non è il numero, va ricostruito

IEEE-754 Singola Precisione (32 bit totali)

Se $x = 0.6$, la sua rappresentazione floating point IEEE-754 è: $F = 00111111000110011001100110011010$ (=0x3f19999a) *numero non rappresentabile finemente*

separazione
↑
↓
e+e
MANTISSA

**Nota: utilizzando il "trucco" della normalizzazione della mantissa nella forma 1.xxxxx *1.11111111*
le cifre (binarie) effettive per la mantissa sono 24 (singola prec.), 53 (doppia prec.), 113 (quad. prec.)

RELAZIONE FRA NUMERO e SUA RAPPRESENTAZIONE IEEE-754 <S,E,M>

• Sia

$$x = (-1)^{\text{segno}} \times \text{mantissa} \times 2^{\text{esponente}}$$

• Nel formato IEEE-754 la rappresentazione F di x è:

$$(x)_{\text{IEEE-754}} = F = \langle S, E, M \rangle$$

dove

S = segno bit per segno
E = esponente + polarizzazione
M = int[(mantissa-1) × 2^{msize}]

ESPO NENTE IN CODIFICA IN BARRIES. + numero fisso
Esigo l'1, e segue il primo elemento
(MANTISSA NORMALIZZ) numero tra 1 e 2, Esigo 1, 1,53 ~> 0,53 × 2^{msize}

RELAZIONE FRA NUMERO x
E SUA RAPPRESENTAZIONE IEEE-754

(S,E,M sono numeri interi rappresentati su un certo numero di bit ovvero

$$S = 0 \text{ o } 1, M < 2^{\text{msize}}, E < 2^{\text{esize}} - 1, \text{polarizzazione} = 2^{\text{esize}-1} - 1$$

Intero, < 2²³

Exp. < 255

polim. 227

SE INT E' PERIODICO TRONCO

Per es. in «IEEE-754 singola precisione» si ha: polarizzazione=127, msize=23, esize=8 quindi:

$$x = 0.375 = (-1)^0 \times 1.5 \times 2^{-2}$$

moltiplicato per 4 per arrivare a 1,5
NORMALIZZAZIONE

$$\rightarrow \text{segno}=0, \text{esponente}=-2, \text{mantissa}=1.5 \rightarrow S=0, E=-2+127=125, M=(1.5-1) \times 2^{23}=2^{22}$$

$$\rightarrow (x)_{\text{IEEE-754}} = F = \langle 0, 125, 2^{22} \rangle_{\text{base10}} = \langle 0, 0111 \ 1101, 1000 \ 0000 \ 0000 \ 0000 \ 0000 \ 000 \rangle_{\text{base2}}$$

INTERPRETA 0,125 × 2²²
segno esponente MANTISSA come codificatore numero

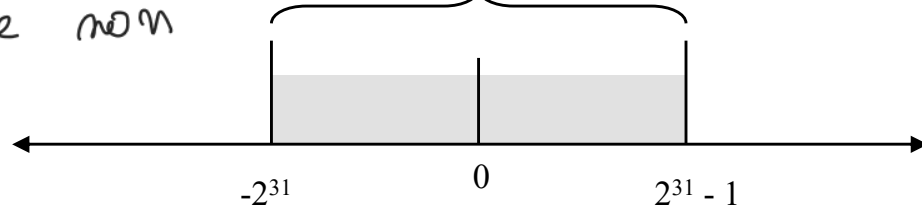
Nota: l'operazione int[] tronca il numero di bit → quindi la rappresentazione IEEE-754 spesso approssima il numero x

Come utilizzare i bit

• Nel caso di interi a 32 bit:

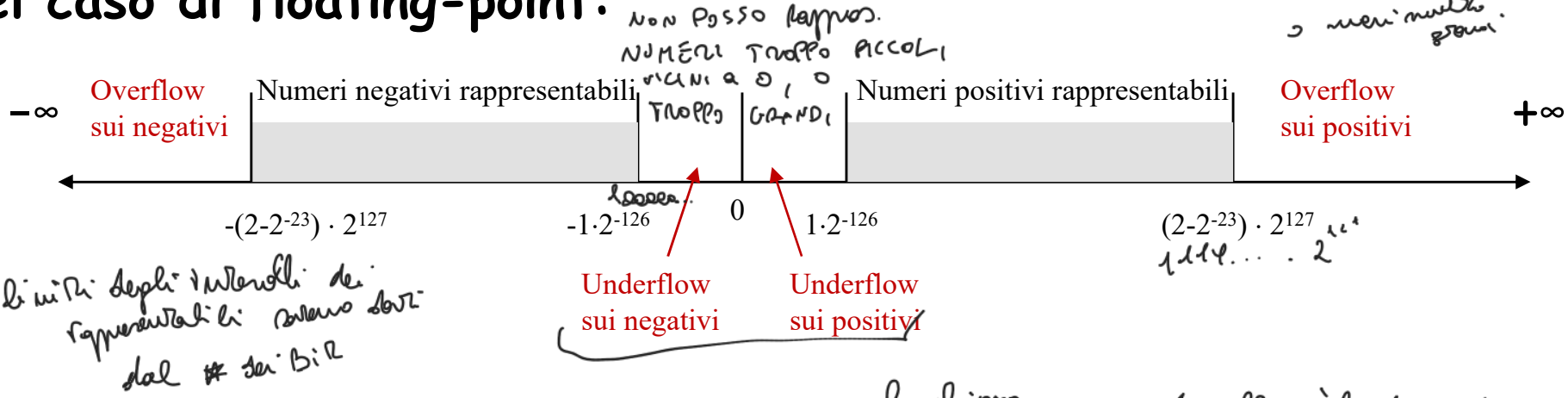
nella parte positiva e negativa
ho numeri rappresentabili e non

equi intero aveva una rappresent. esatta e unica
interi rappresentabili (in complemento a 2)

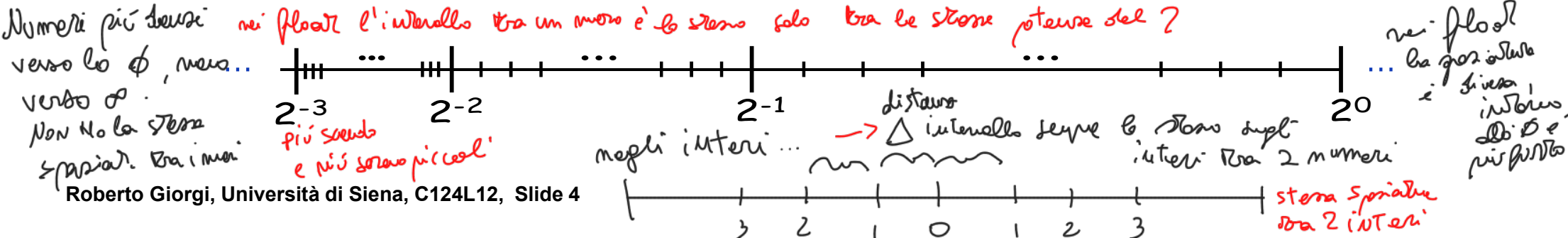


nei float ho
numeri non rappresentabili
più saranno numeri molto piccoli
in un intorno dello 0,
e numeri molto grandi.

• Nel caso di floating-point:



Nota: il numero di valori rappresentati con 32 bit è circa lo stesso nel caso INT e in quello FP... la densità dei punti non è costante, quindi.



$exp = 128, 127$ sono valori speciali

IEEE 754 floating-point mantissa ed esponente

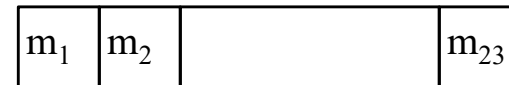
• Mantissa

- Rappresentazione in base due con Normalizzazione (e modifica dell'esponente) in modo tale da riportare un 1 nella cifra più significativa

- (tale unico) 1 prima della virgola non compare nella codifica

- I bit dalla mantissa rappresentano un numero decimale (in base 2) ≥ 1 e < 2

- $M = 1 - \text{mantissa} = m_1 \times 2^{-1} + m_2 \times 2^{-2} + \dots + m_{23} \times 2^{-23}$
numero intero ottenuto moltiplicando bit mantissa per il loro peso



devo approssimare

Nota: come detto, le cifre oltre la m_{23} le butto!

• Esponente

- Rappresentazione "polarizzata" (biased) per rendere più facile il sort

- polarizzazione = 127 per la singola precisione, 1023 per la doppia precisione

- i valori vanno da 1 a $2^e - 2$ ($e = \# \text{bit per l'esponente}$) e codificano gli esponenti da $-2^{e-1} + 2$ a $+2^{e-1} - 1$ (e.g. -126..+127)

- Codifiche riservate

- **esponente con tutti 1 e mantissa $\neq 0$** } $exp = 127 + bias = 255$, mantissa $\neq 0$, NAN
codifica "Not a Number" (NAN) (risultato di 0/0)

- **esponente con tutti 1 e mantissa $= 0$** $exp = 127$ e mantissa $= 0$

- ✓ codifica "Infinity" (INF) (risultato di un numero diverso da zero diviso 0) $\circ - \text{se segno}$

- **esponente con tutti 0 e mantissa $\neq 0$**

- codifica "underflow" (numero denormalizzato) Sono nell'intervallo dello 0 dell'underflow troppo piccolo

- Ricapitolando: per $(-1)^S \times (1 + M) \times 2^{E - \text{polarizzazione}}$ $\rightarrow$ memorizzo $\langle S, M, E \rangle$

- $S = 0$ o 1 , $M = \text{mantissa_normalizzata senza "1." iniziale}$, $E = \text{esponente} + \text{polarizzazione}$

Esercizio

- Rappresentare -0.75 in formato IEEE single precision

- rappresentazione decimale: $-0.75 = -3/4 = -3 \times 2^{-2} = -1.5 \times 2^{-1}$
- rappresentazione binaria: $-11 \times 2^{-2} = -1.1 \times 2^{-1}$ *moltiplico per 2*
- rappresentazione binaria normalizzata: -1.1×2^{-1}

- segno=negativo
→ **S=1**

$$M = 0.5 \cdot 2^2 = 2^1$$
$$Exp = -1 + 127 = 126$$

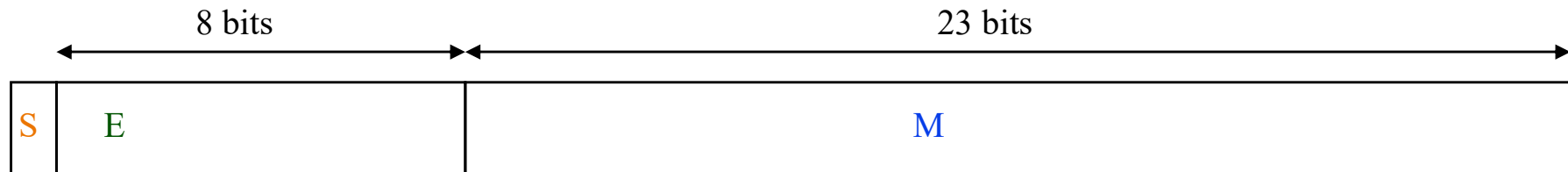
- mantissa_normalizzata=1.1000 0000 0000 0000 0000 000
→ **M = 1000 0000 0000 0000 0000 000**

- esponente=-1
→ **E=esponente+polarizzazione = -1 + 127 = 126 → 01111110**

- IEEE single precision: **10111111010000000000000000000000**

- Nota: la mantissa è stata normalizzata. Per questo tutti i bit più significativi risultano spostati verso SINISTRA.

Floating-Point: single precision (32 bits)



Esempi con mantissa ed esponente positivi/negativi:

float. e' stato implementato perche' risultava piu' veloce gli algoritmi di Sort

$+1.1010001 \cdot 2^{+10100} \rightarrow 0 \overset{+127}{10010011} 101000100000000000000000$
 $-1.1010001 \cdot 2^{+10100} \rightarrow 1 \overset{+127}{10010011} 101000100000000000000000$
 $+1.1010001 \cdot 2^{-10100} \rightarrow 0 \overset{+127}{01101011} 101000100000000000000000$
 $-1.1010001 \cdot 2^{-10100} \rightarrow 1 \overset{+127}{01101011} 101000100000000000000000$

Esempi di numeri "notevoli":

$0 \rightarrow 0 \overset{0+127=127}{00000000} 000000000000000000000000$
 $1 \rightarrow 0 \overset{0+127=127}{01111111} 000000000000000000000000$
 minore rappresentabile (2^{-126}) $\rightarrow 0 \overset{1 \cdot 2^{-126}}{00000001} 000000000000000000000000$
 maggiore rappres. ($(2-2^{-23}) \cdot 2^{127}$) $\rightarrow 0 \overset{1 \cdot 2^{-126}}{11111111} 111111111111111111111111$
 $\infty \rightarrow 0 \overset{1 \cdot 2^{-126}}{11111111} 000000000000000000000000$
 NaN $\rightarrow 0 \overset{1 \cdot 2^{-126}}{11111111} \underbrace{\text{XXXXXXXXXXXXXXXXXXXXXXXXXXXX}}_{\neq 0}$

esponente = float

MANTISSA

1,000...0

numero più piccolo Vaguard.

-127 e 128, codificato rinormalizzato

Nota: $2^{-126} \approx 1.18 \cdot 10^{-38}$, $(2-2^{-23}) \approx 3.4 \cdot 10^{38}$

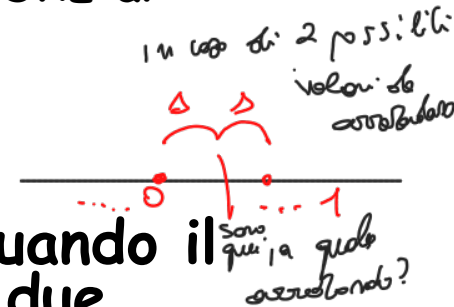
Modalità di arrotondamento IEEE-754

- IEEE-754 prevede 4 modalità di arrotondamento

- **round to nearest**: arrotonda al più vicino rappresentabile *Rappresentabili*
- **round to $+\infty$** : arrotonda al primo numero rappresentabile NON INFERIORE al valore da arrotondare (eventualmente arrotonda a $+\infty$) *a $+\infty$*
- **round to $-\infty$** : arrotonda al primo numero rappresentabile NON SUPERIORE al valore da arrotondare (eventualmente arrotonda a $-\infty$) *a $-\infty$*
- **round to 0**: arrotonda al più vicino rappresentabile con modulo NON SUPERIORE al valore da arrotondare *verso 0, in dipendenza da dove è lo 0*



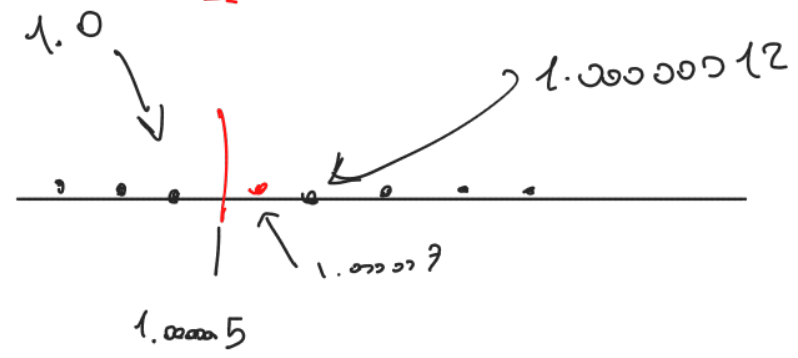
- Nel caso di "round to nearest" ci può essere ambiguità quando il numero da arrotondare sia esattamente equidistante dai due numeri più vicini rappresentabili



- Lo standard adotta la regola "**round ties to even**", ovvero in caso di pareggio va scelto il più vicino rappresentabile PARI
- Il vantaggio rispetto alla modalità "**round ties to away (from 0)**" (supportata nella revisione dello standard IEEE-754-2008) *arrotondo al numero più lontano da 0* è che in media l'errore di arrotondamento si annulla**
- L'implementazione di "round ties to even" è più complessa

Esempi di arrotondamento

numero non rappresentabile	r.to 0	r.to nearest	r.to +inf	r.to -inf
1.00000003	1.0	1.0	1.00000012	1.0
1.00000007	1.0	1.00000012	1.00000012	1.0
-1.00000003	-1.0	-1.0	-1.0	-1.00000012
-1.00000007	-1.0	-1.00000012	-1.0	-1.00000012



Ricapitolazione

- L'aritmetica del calcolatore è limitata da precisione finita
- I pattern di bit non hanno particolare significato ma esistono standard per: *interpretare i bit*
 - complemento a due
 - IEEE 754 floating point
- Le istruzioni determinano il "significato" dei pattern di bit *a seconda delle istruzioni interpretare i bit*
- Le prestazioni e l'accuratezza sono importanti: ci sono diverse complessità nelle macchine reali
- Lo standard è stato rivisto nel Giugno 2008 come IEEE 754-2008:
<http://ieeexplore.ieee.org/stamp/stamp.jsp?arnumber=4610935&isnumber=4610934>
 - Fra le novità il supporto per floating point a 128-bit
 - Aritmetica floating point decimale (oltre che binaria)

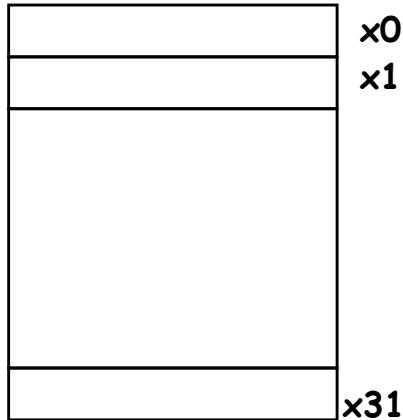
*ho bisogno
delle istruzioni
per i float*

OPERAZIONI FLOATING-POINT NEL PROCESSORE RISC-V

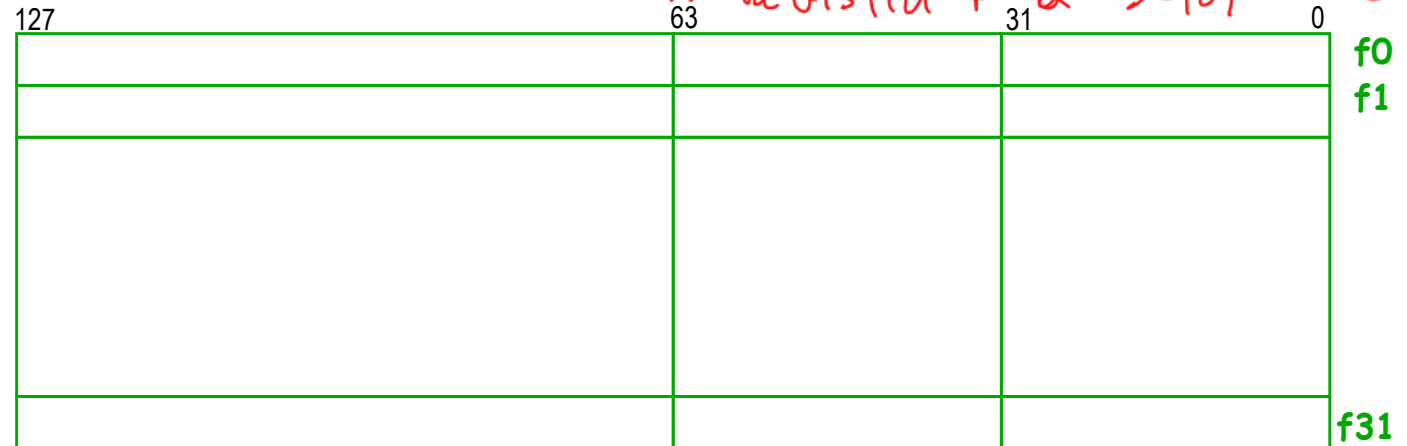
numero 64 bit, rappresentato a 32, preciso,
rimane sempre rappresentabile, o meno
di salvare la info del segno

Supporto architetturale per le op. floating-point nel RISC-V^{64/128}_{li2}

REGISTRI PER OPERAZIONI SU INTERI:



REGISTRI AGGIUNTIVI *aliqui registri F*
PER OPERAZIONI SU FLOATING-POINT (s=32 bit, d=64bit, q=128bit):



64 bit

←→

single destination operand

• fadd.s f0, f1, f2

• fadd.d f0, f2, f4

• fadd.q f0, f2, f4

A diagram illustrating the bit lengths of the three components of the SHA-256 hash. It consists of three horizontal arrows pointing to the right, stacked vertically. The top arrow is labeled "128 bit", the middle arrow is labeled "64 bit", and the bottom arrow is labeled "32 bit".

lavora sui registri fp a 32 bit	(estensione 'F' del RISC-V)
lavora su registri fp a 64 bit	(estensione 'D' del RISC-V) <i>allora ha a disposizione</i>
lavora su registri fp a 128 bit	(estensione 'Q' del RISC-V) <i>128</i>

Istruzioni floating point nel processore RISC-V

- ARITMETICHE (dove 'z' può essere s per single, d per double, q per quad)

fadd.z f1, f2, f3 *z = s, d, q*

somma

fsub.z

sottrazione

fmul.z

moltiplicazione

fdiv.z

divisione

fsqrt.z f1, f2

radice quadrata

fmv.z f1, f2 *qui non è una pseudo inst.*

spostamento fra reg. fp

fabs.z f1, f2

valore assoluto

fneg.z f1, f2

valore contrario (negato) *non è booleano*

In blu-chiaro: pseudoistruzioni basate su FSIGNJ

- MEMORIA

non perdere della memoria come a 32, 64, 128 bit
flw / fld / flq f1, 100(x5)

flw:32bit, fld:64bit, flq:128bit

fsw / fsd / fsq f1, 100(x5) *sono in complemento a 2*

- SPOSTAMENTO *facciamo riferimento a parti del circuito di verso* int $\leftrightarrow$ fp

intero 32 bit
fmv.x.w x5, f1

fmv.w.x f1, x5

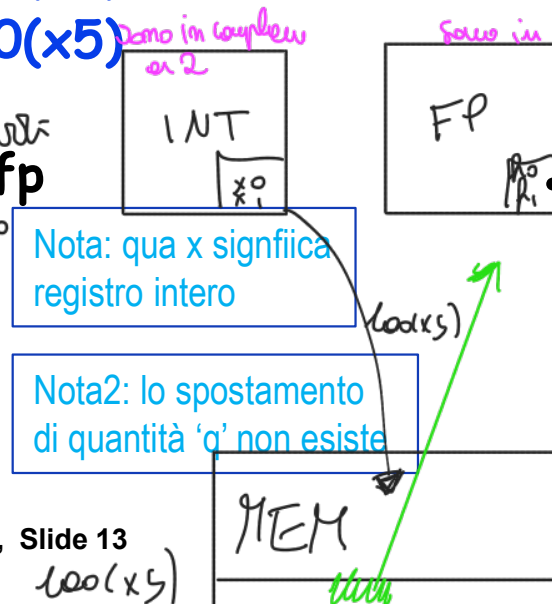
32 bit in 64 bit float
fmv.x.d x5, f1

fmv.d.x f1, x5

facciamo anche interpretazione

Roberto Giorgi, Università di Siena, C124L12, Slide 13

sempre conversione da float a int



Converti da float a int.
SPOST. CON CONVERSIONE

fcvt.s.w / fcvt.s.wu f1, x5

fcvt.w.s / fcvt.wu.s x5, f1

fcvt.d.w / fcvt.d.wu f1, x5

fcvt.w.d / fcvt.wu.d x5, f1

W = intero con segno (word)
WU=intero senza segno (word-unsigned)

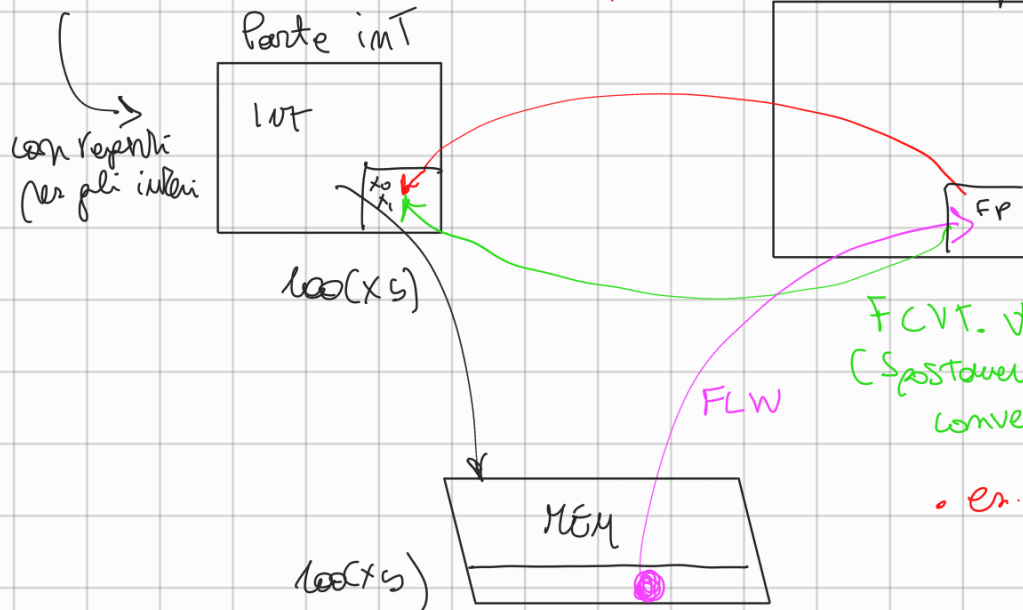
Parte del Processore
che elabora gli interi

$F.M.W. \cdot X.W$

Parte float; con registri float

2 cose diverse

Caso di MIX



• es. gli indirizzi, sempre
sono interi

• Nel caso dell'accesso alla memoria uso sia float che interi, registri X per gli indirizzi della memoria ma carico un valore che è un float

Dalla memoria prendere con load e store
elementi a 32, 64, 128 bit

Istruzioni condizionali floating-point

- Sono nella forma

`fcc.s x5, f2, f4`

dove '**cc**' indica una condizione di confronto e può essere:

'lt' → less than

'le' → less than or equal

'eq' → equal

- Perché non ci sono i casi (gt, ne, ge)?

Nel RISC-V si è scelto di non inserire istruzioni per tali casi in quanto basta... scambiare gli operandi* per trattarli!

- Il risultato va su un registro intero

- Esempio:

`flt.s x5, f2, f4`

`beq x5, x0, L1`

il risultato di confronto è vero o falso
risultato booleano, in registro int

* oppure, es., usare bne invece beq nell'uso del valore booleano della condizione

Istruzione FCLASS.x *non la usiamo*

- FCLASS.z permette di sapere se il numero FP è normalizzato o meno, +/-INF, +/-0, NaN
 - dal manuale "RISC-V User-Level ISA V2.2"

The FCLASS.S instruction examines the value in floating-point register *rs1* and writes to integer register *rd* a 10-bit mask that indicates the class of the floating-point number. The format of the mask is described in Table 8.5. The corresponding bit in *rd* will be set if the property is true and clear otherwise. All other bits in *rd* are cleared. Note that exactly one bit in *rd* will be set.

31	27 26	25 24	20 19	15 14	12 11	7 6	0
funct5	fmt	rs2	rs1	rm	rd	opcode	
5	2	5	5	3	5	7	
FCLASS	S	0	src	001	dest	OP-FP	

<i>rd</i> bit	Meaning
0	<i>rs1</i> is $-\infty$.
1	<i>rs1</i> is a negative normal number.
2	<i>rs1</i> is a negative subnormal number.
3	<i>rs1</i> is -0 .
4	<i>rs1</i> is $+0$.
5	<i>rs1</i> is a positive subnormal number.
6	<i>rs1</i> is a positive normal number.
7	<i>rs1</i> is $+\infty$.
8	<i>rs1</i> is a signaling NaN.
9	<i>rs1</i> is a quiet NaN.

RISC-V Floating-Point ABI (Application Binary Interface)

- Definisce la convenzione di uso dei registri nel software
- Valido nel caso single e in quello double precision

- f10,f11 (chiamati **fa0,fa1**) → argomenti delle funzioni o valori restituiti

- f12-f17 (chiamati **fa2,fa7**) → argomenti ulteriori di funzioni

- f8,f9,f18-f27 (chiamati **fs0-fs11**) → registri **"saved"**
per i registri FP non ho registri speciali
non salvati se li uso

- f0-f7,f28-f31 (chiamati **ft0-ft11**) → registri **temporanei**

*fo è un registro
generico
R, ecc. manuale*

*non è detto che FO è $\emptyset$, non ha
registri speciali;
non mi serve*
X0 è $\emptyset$

Alcune syscall che coinvolgono i numeri floating point

- 'ecall' serve per chiedere un servizio al sistema operativo (chiamata anche 'system call')
 - PARAMETRI DI INGRESSO/USCITA
 - **a7** deve contenere il numero del servizio
 - 2: stampa un single-precision float a video (fa0 contiene il float da stampare)
 - 6: leggi un single-precision float dalla tastiera (fa0 in uscita, contiene il float letto in binario)
 - 3: stampa un double-precision float a video (fa0 contiene il double da stampare)*
 - 7: leggi un double-precision float dalla tastiera (fa0 contiene il float letto in binario)*