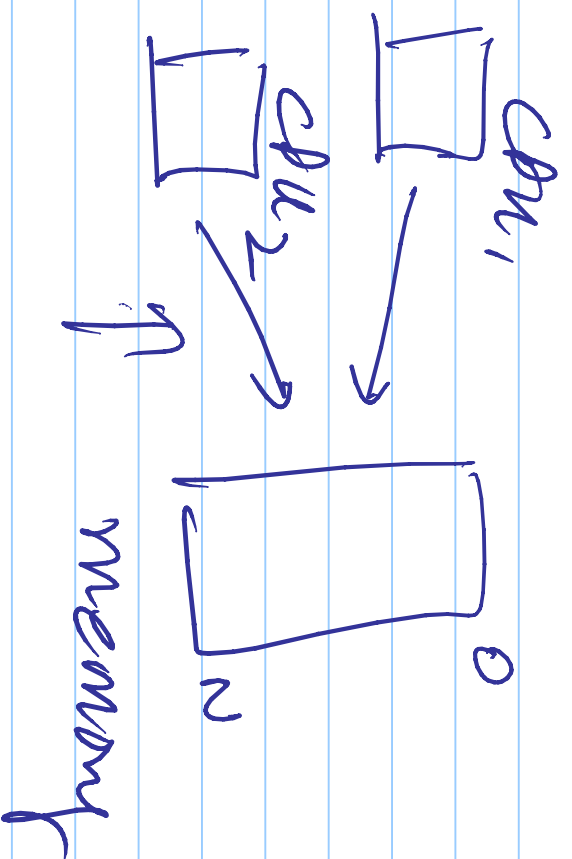


UMA - uniform mem access

↳ SMP - sym multi proc

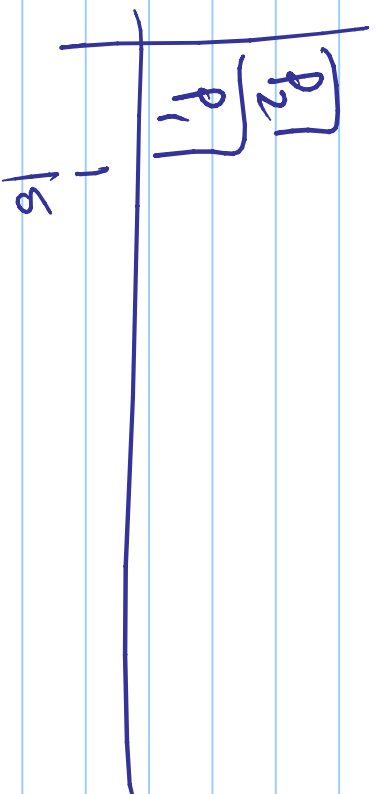
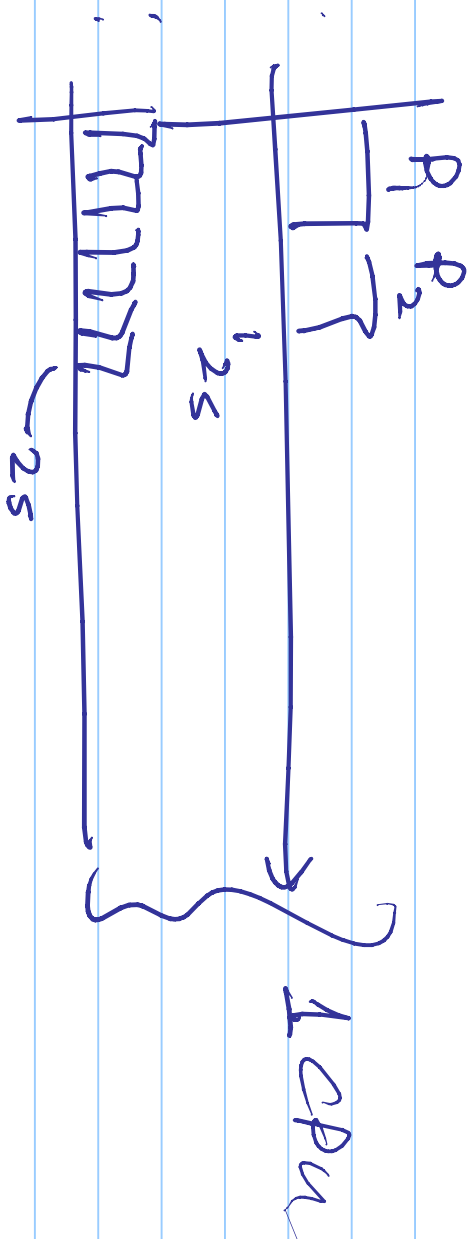
SMP - shared mem multi proc

2 processors



2 processes

2 processes $\rightarrow$ each takes 1s of CPU



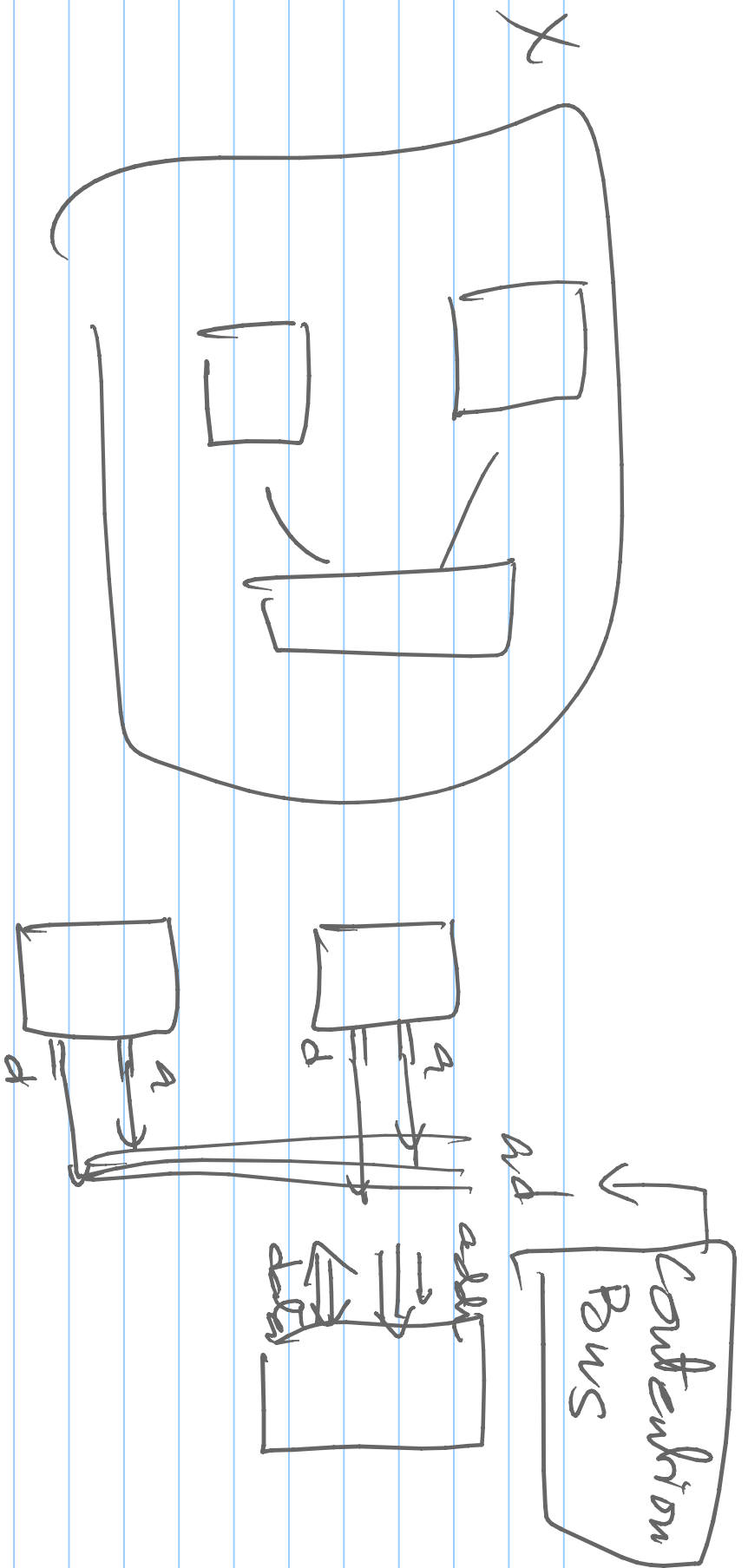
Single proc - 2GHz

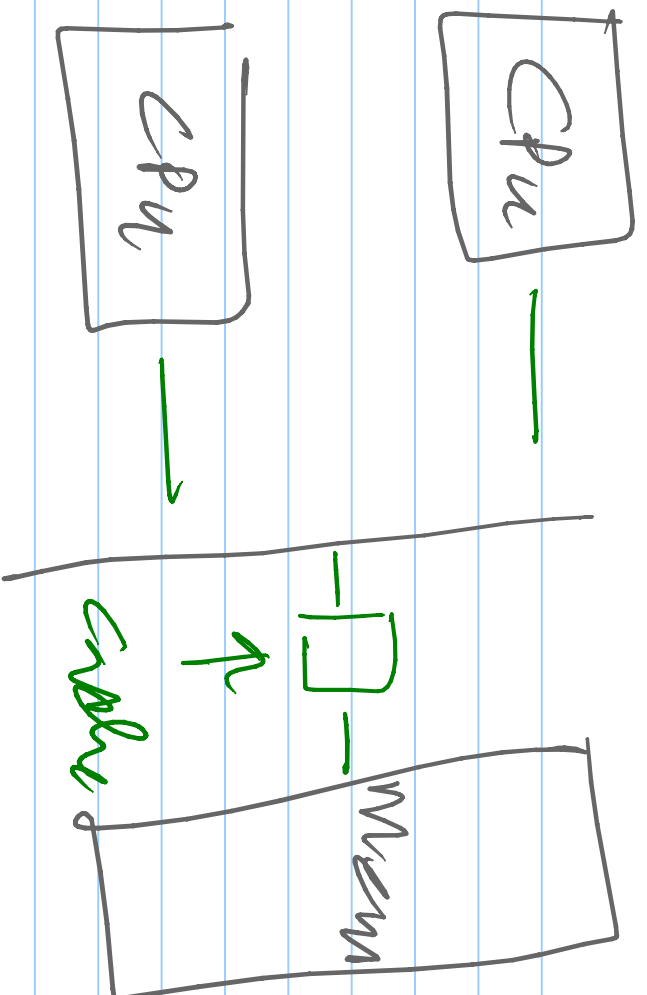
Multi~~proc~~ proc - $2 \times 1\text{GHz}$ }

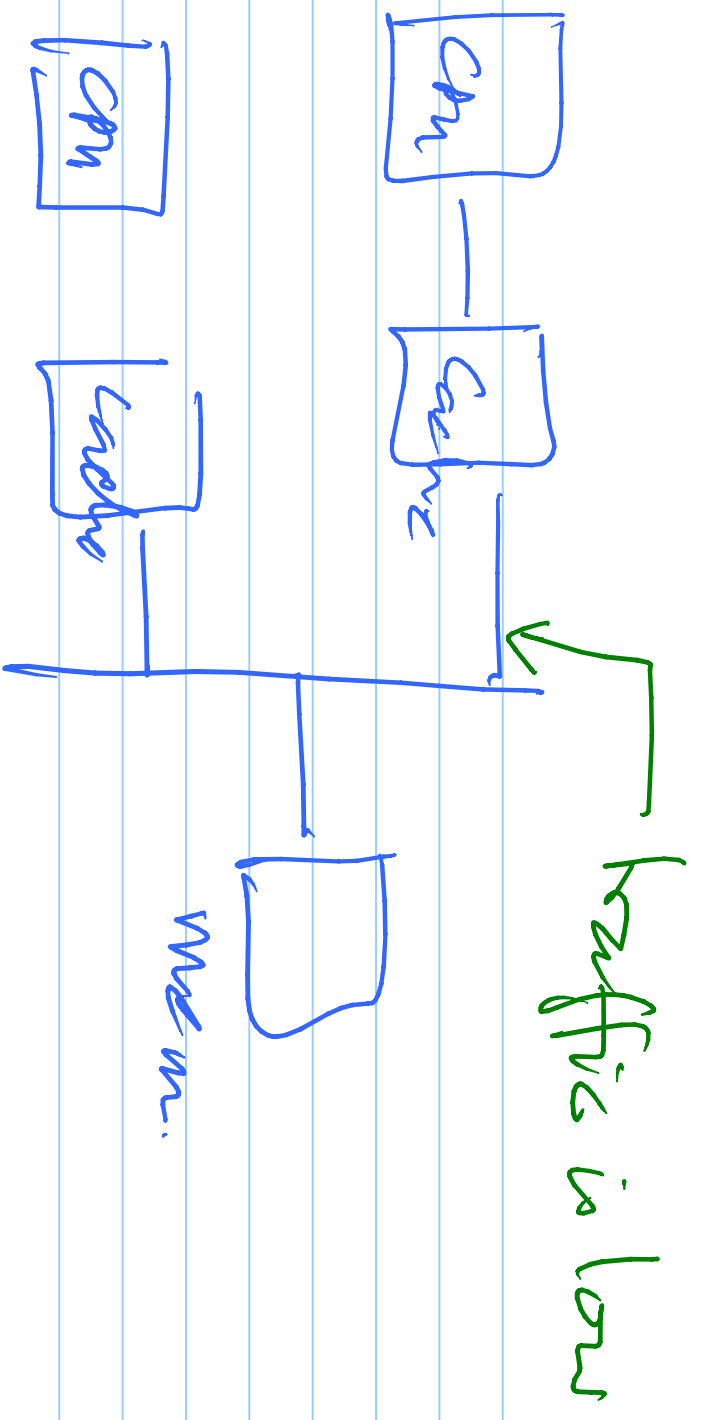
↳ slow on

a single

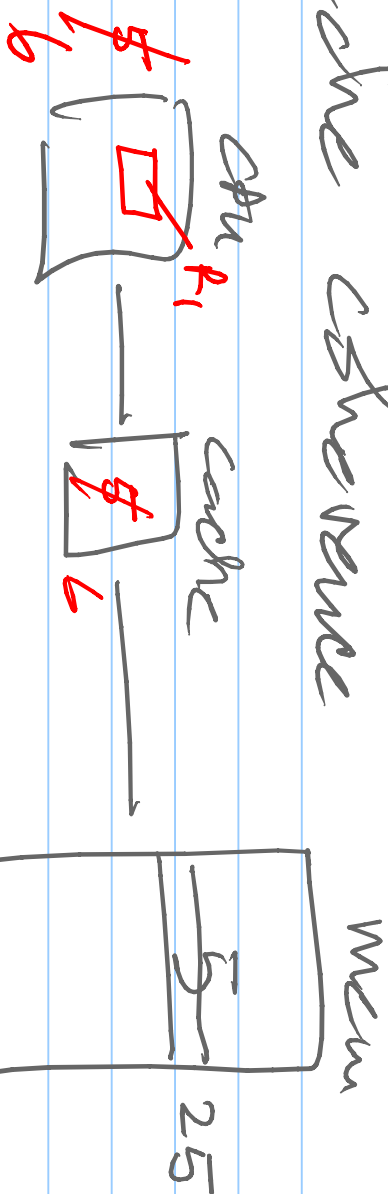
process.







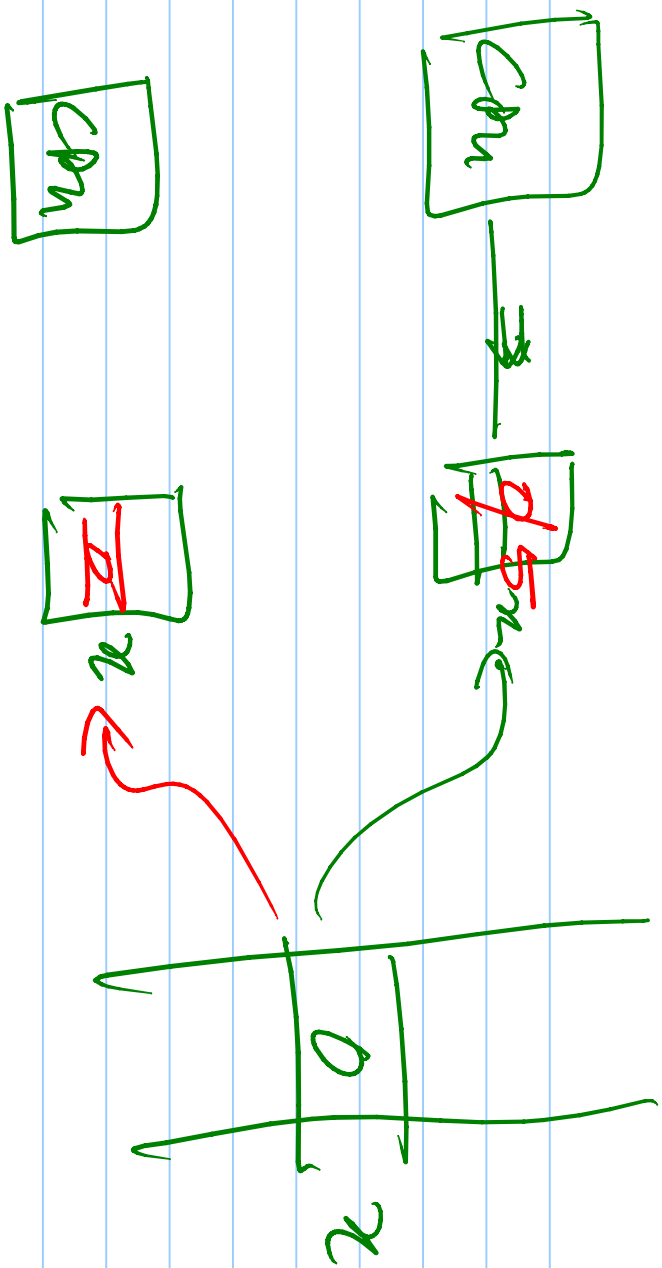
Cache coherence



i = 5
i = 6
print i

load 25 → R1
add 1 to R1
sto R1 → 25
print @25 → 6

a read from
main
returns the
most recent
write



Single copy semantics

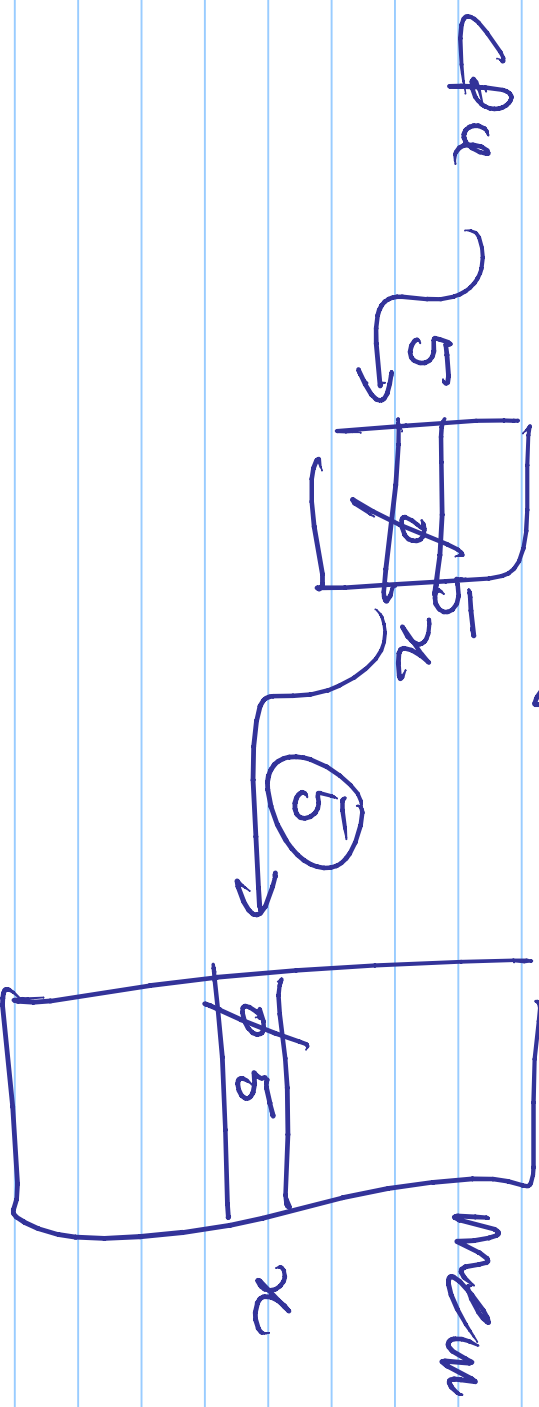
a.k.a coherence

"a real return the
most recent write"

⇒ needed

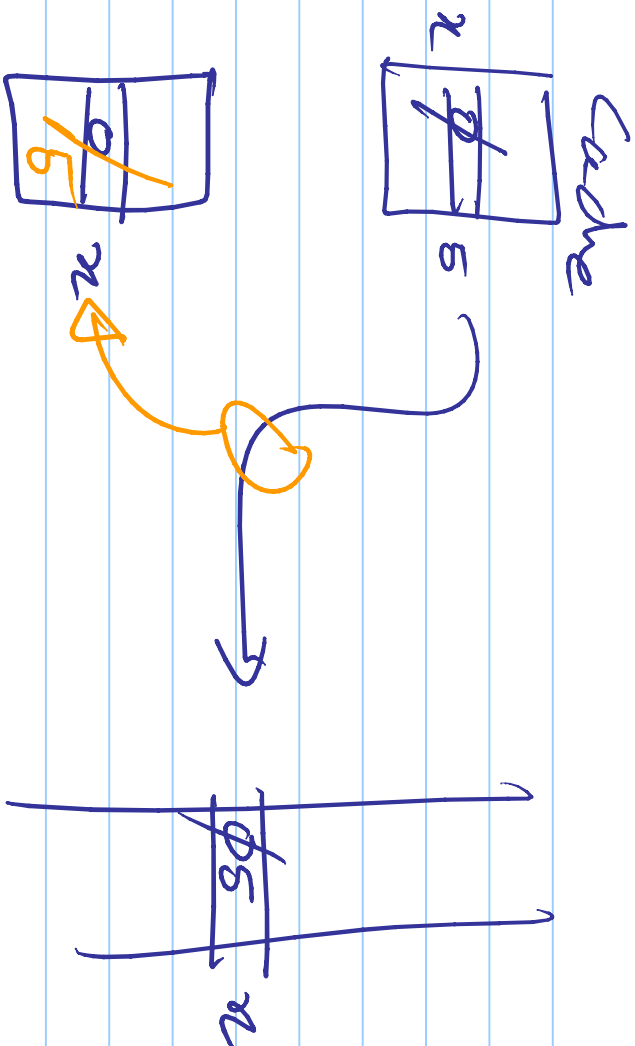
Cache coherence

① use write through caches



Snoopy caches

↓
each
cache
snoops
the bus



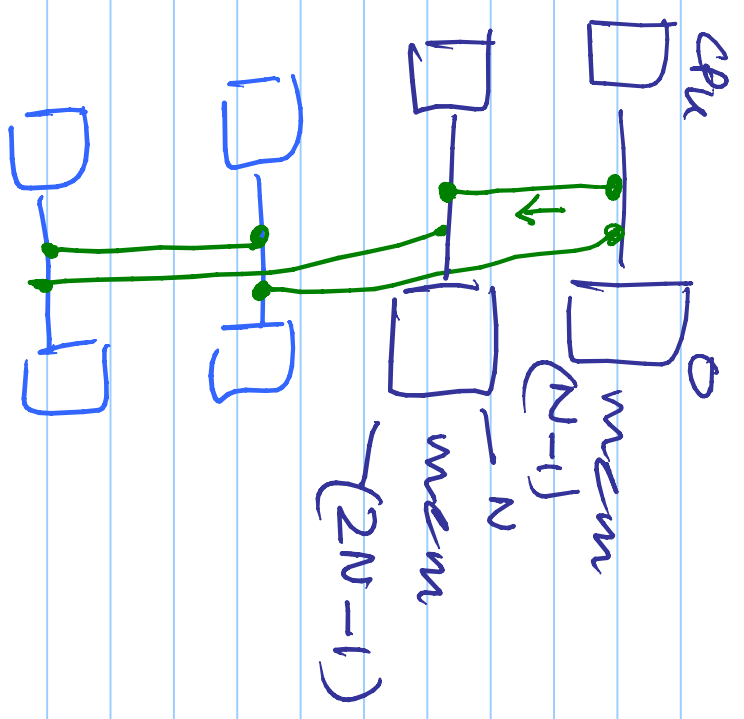
Shared mem multicores (SMT)

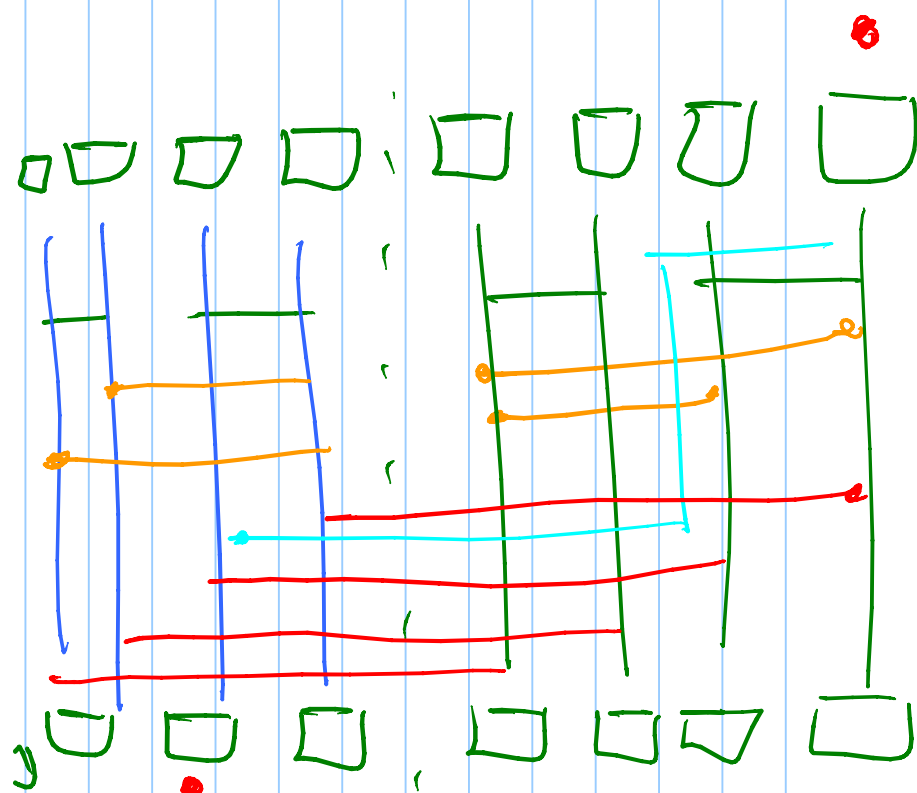
- has bottlenecks $\rightarrow$ adder bus traffic
- " hardware overhead $\rightarrow$ coherence

$\Rightarrow$ Does not scale!

How to scale?

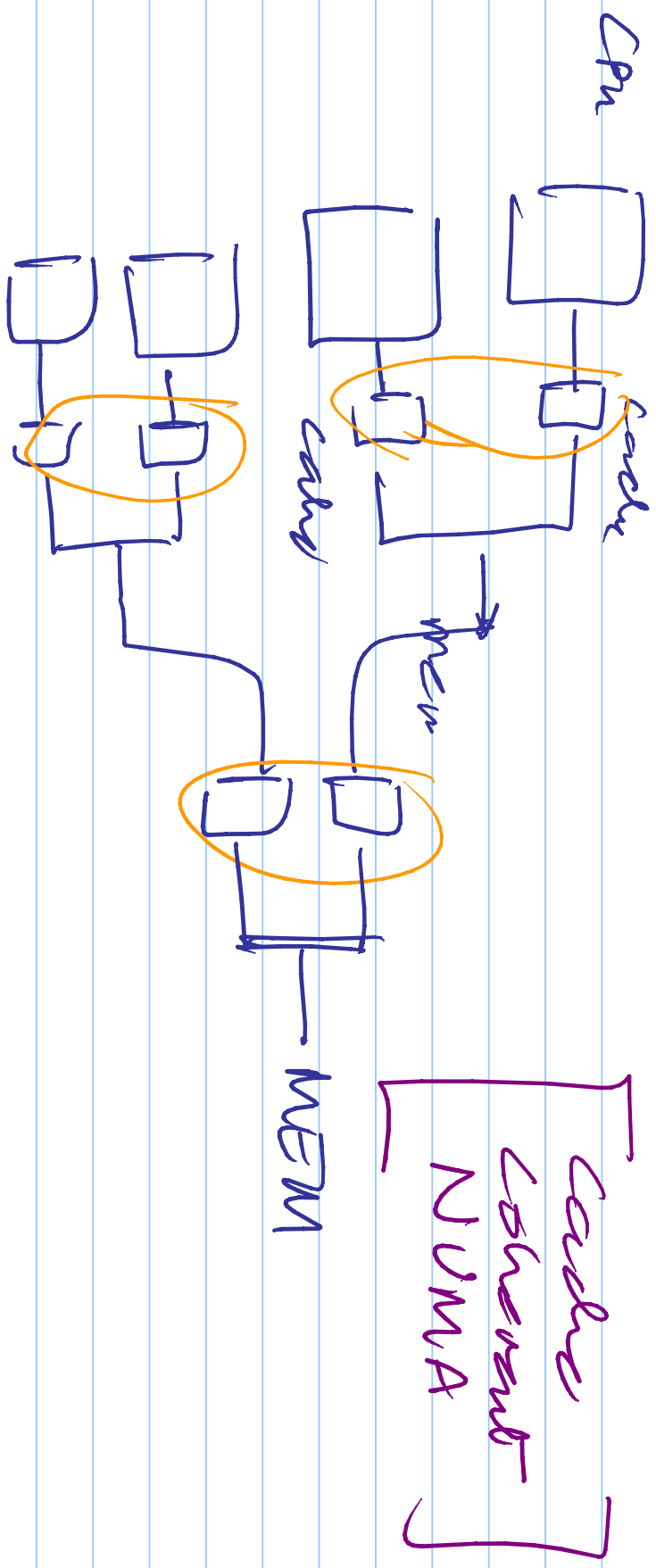
① NUMA



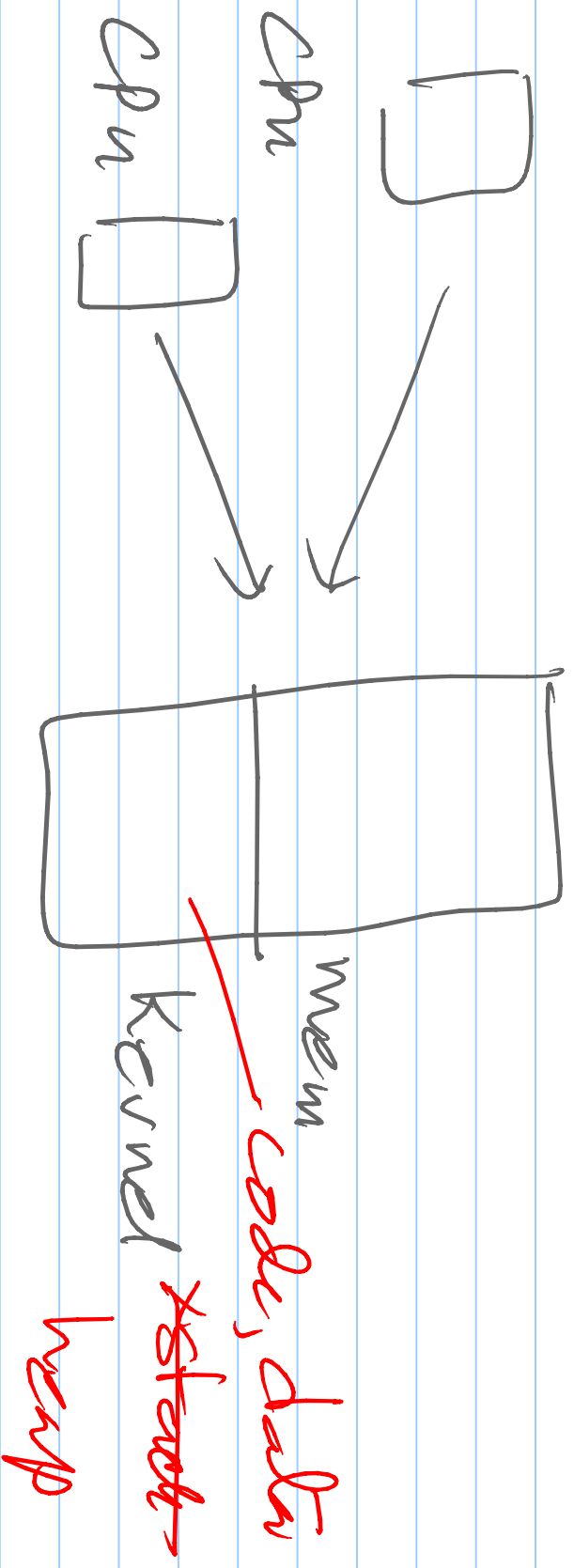


hierarchical caches

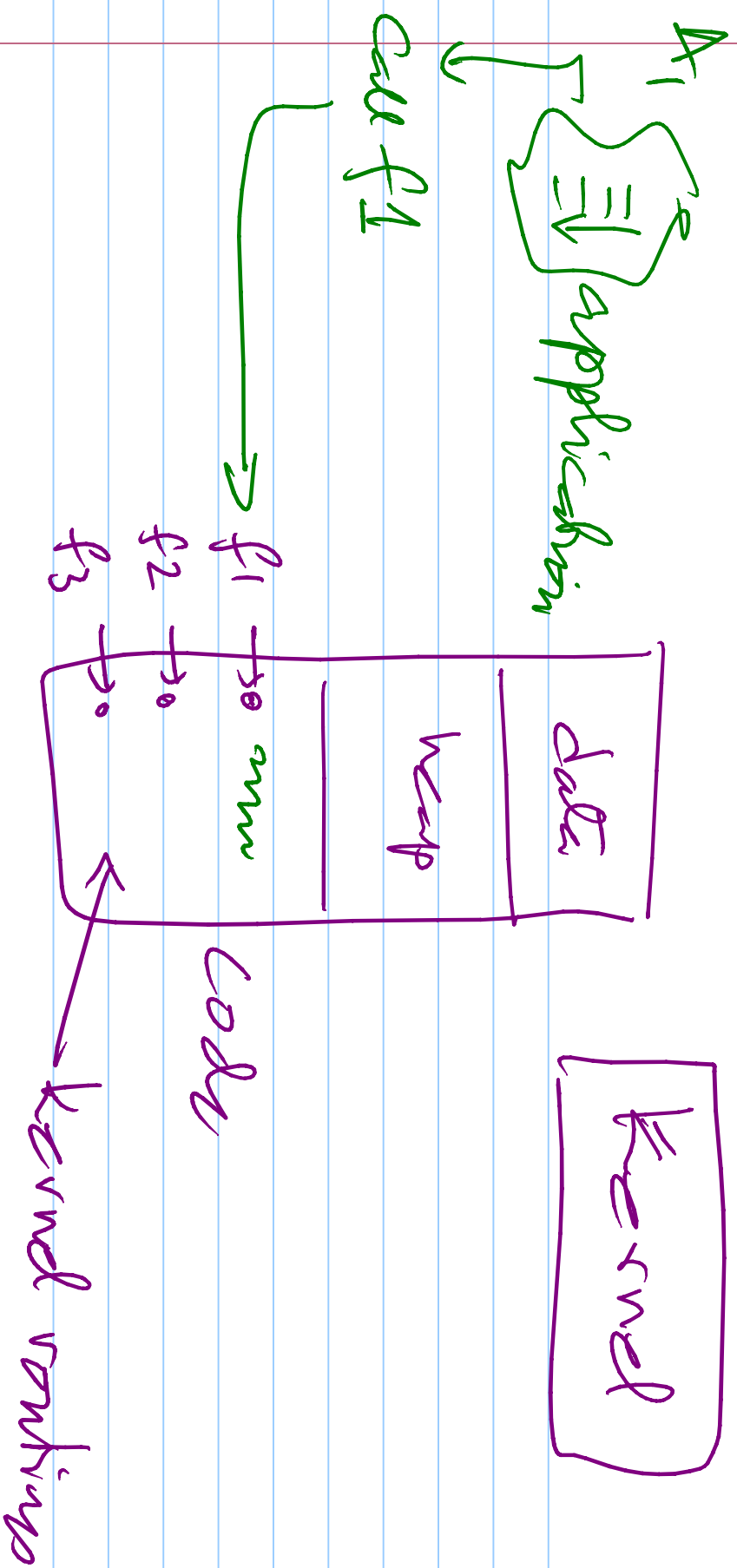
CC-NUMA



The shared memory problem



kernel



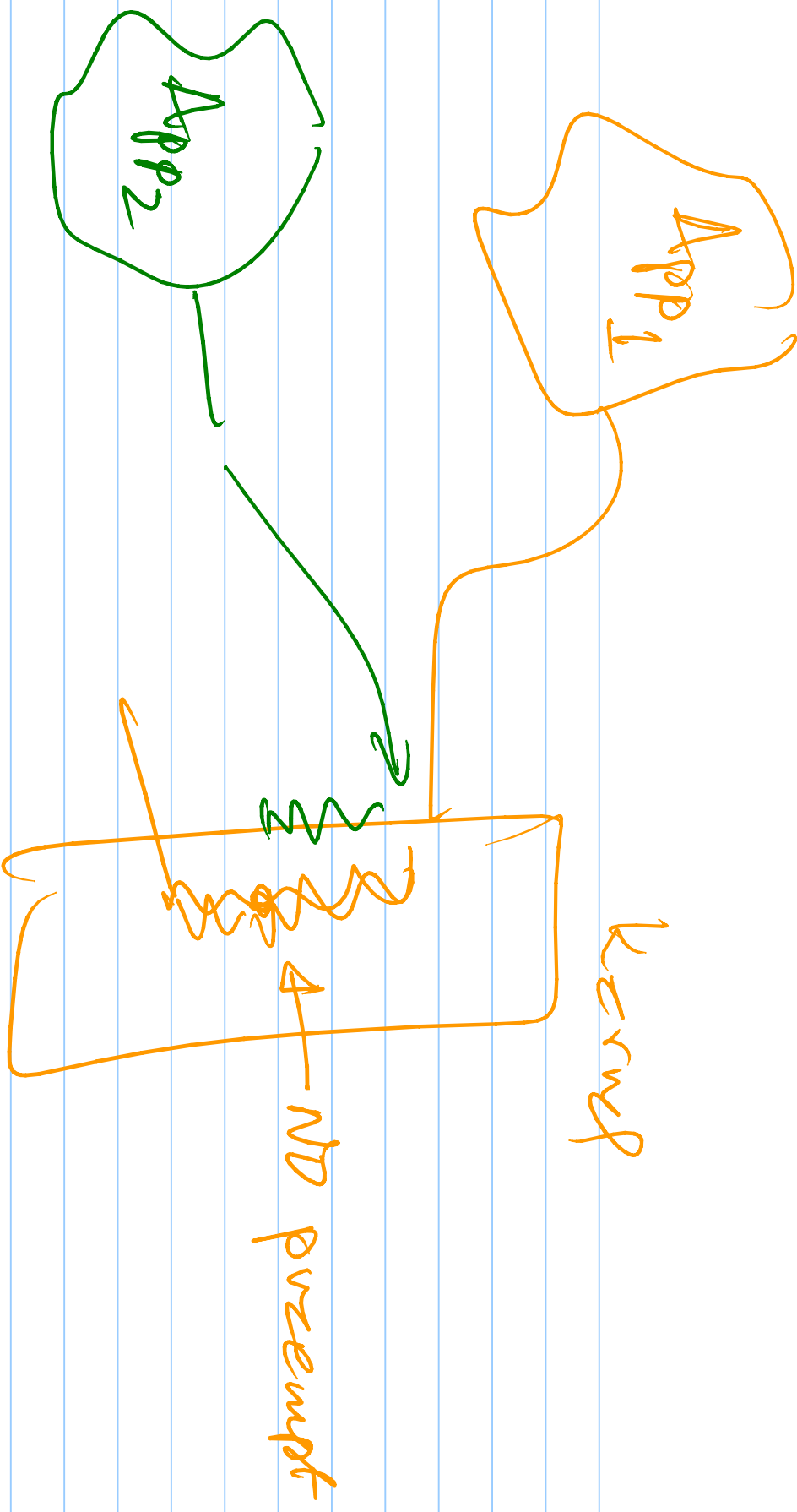
kernel preemption?

context
switch.

kernel
fn

NO $\rightarrow$ non preemptible
kernel

stop
preemption
app
x preempt
— OK to
preempt



kernels

└─ UNIProc

① not preemptible (v1)

② ~~mp~~ → "big lock" (v2)

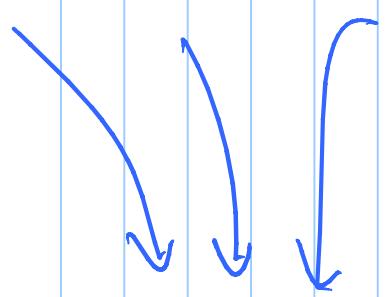
③ No big lock

→ (v3)

Single / multiple cpu

`int x` | `global`

`func`



{

-

`x++`

- application process

- threads

}

Single proc

