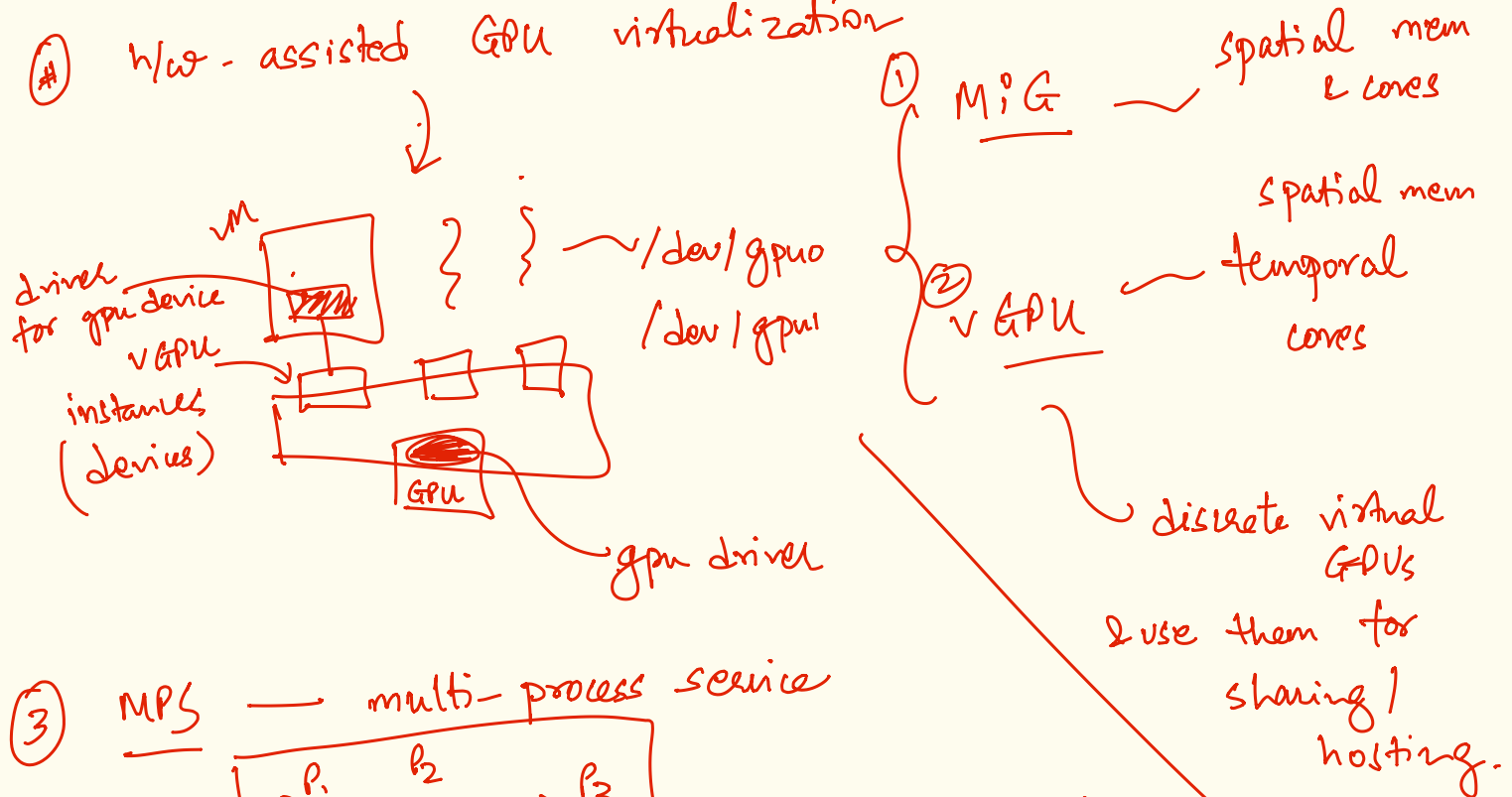


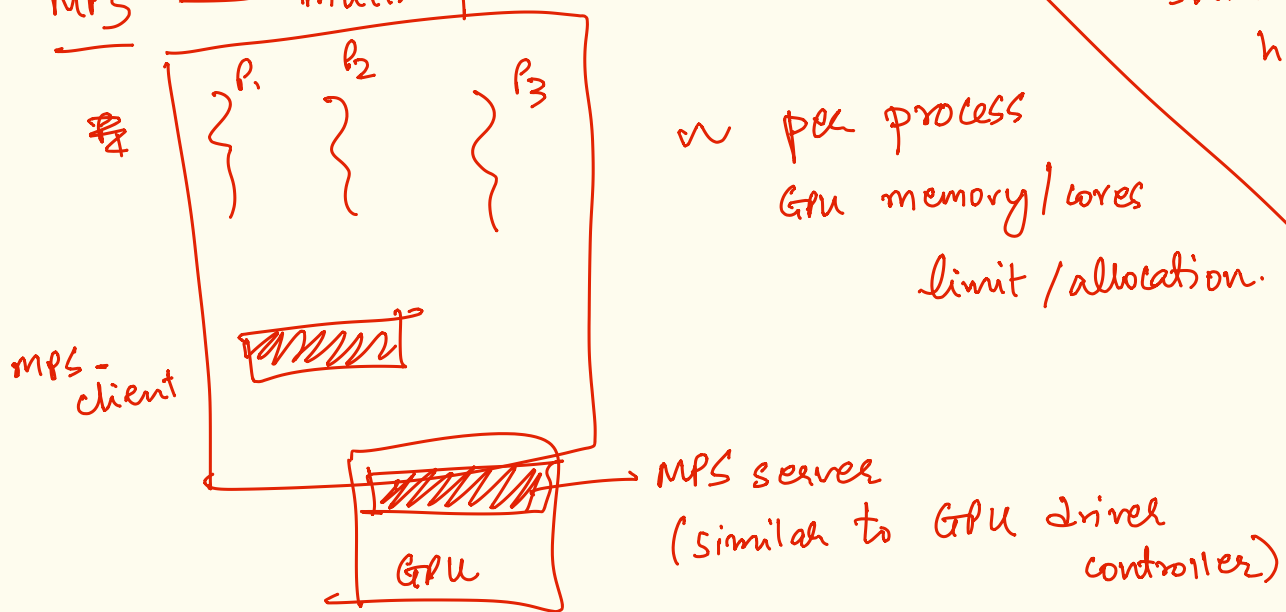
CS 695

GPU virtualization

① w/o - assisted GPU virtualization



③ MPS — multi-process service



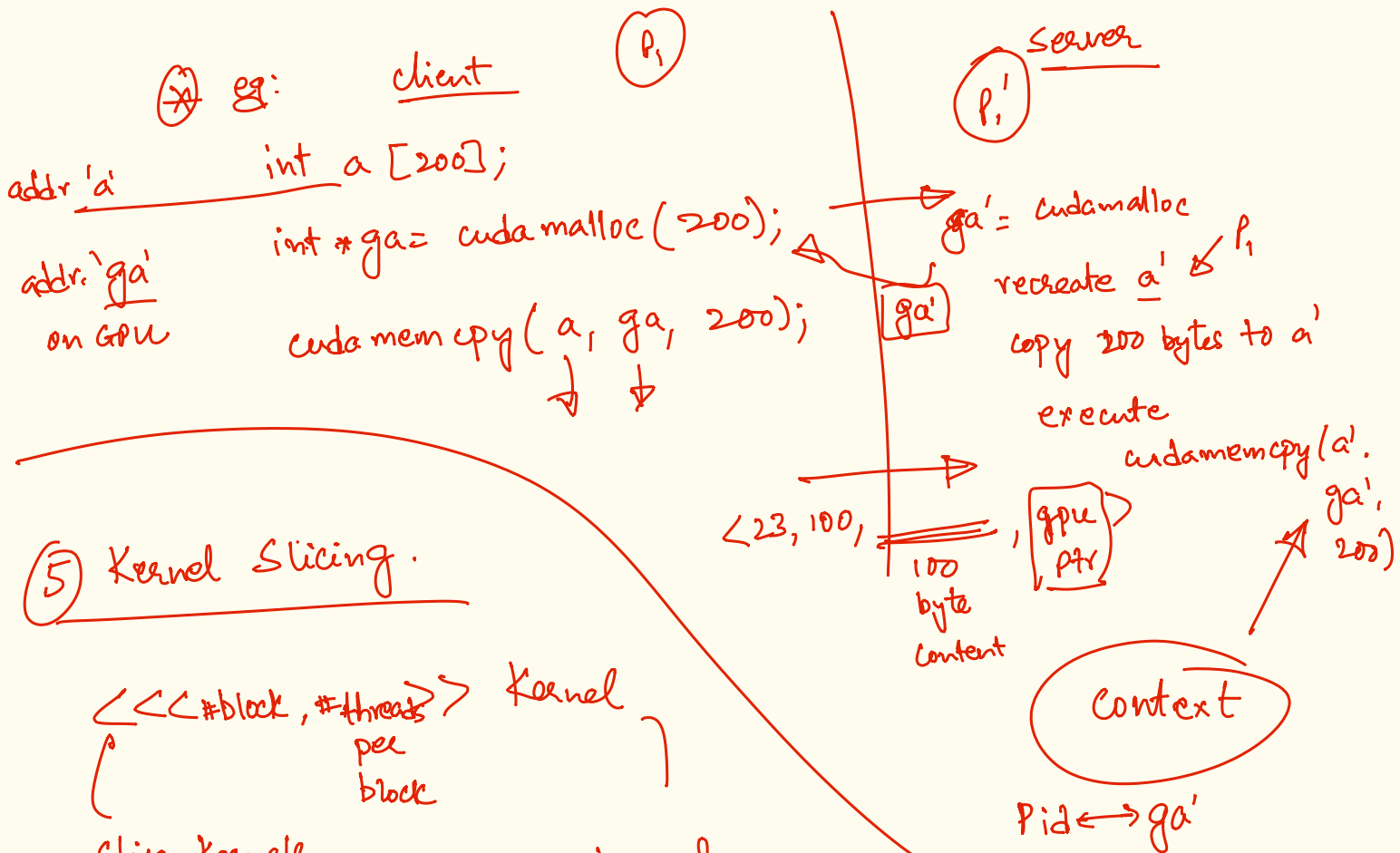
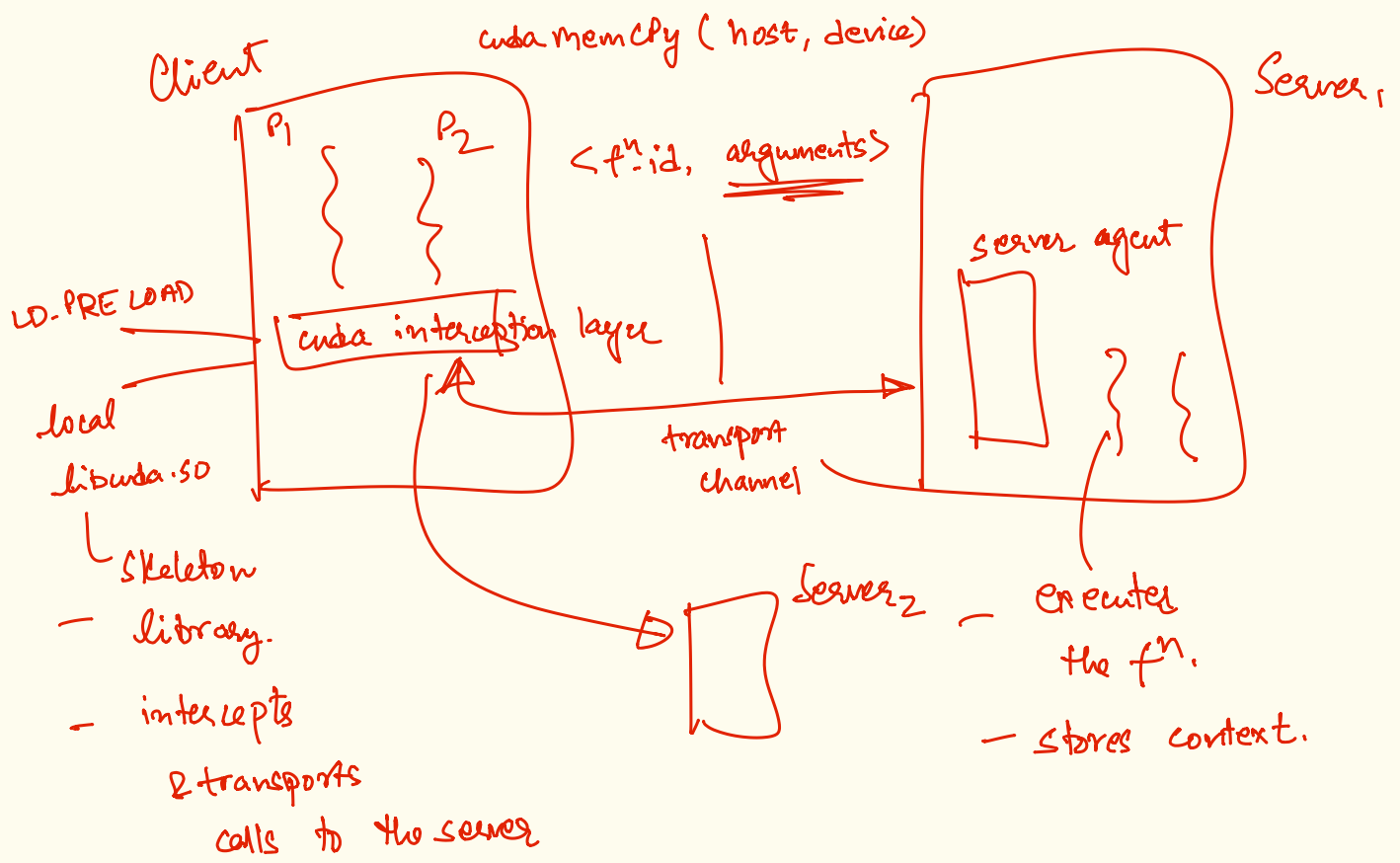
(gpu driver - assisted (c/w) virtualization)

④ API virtualization / remoting

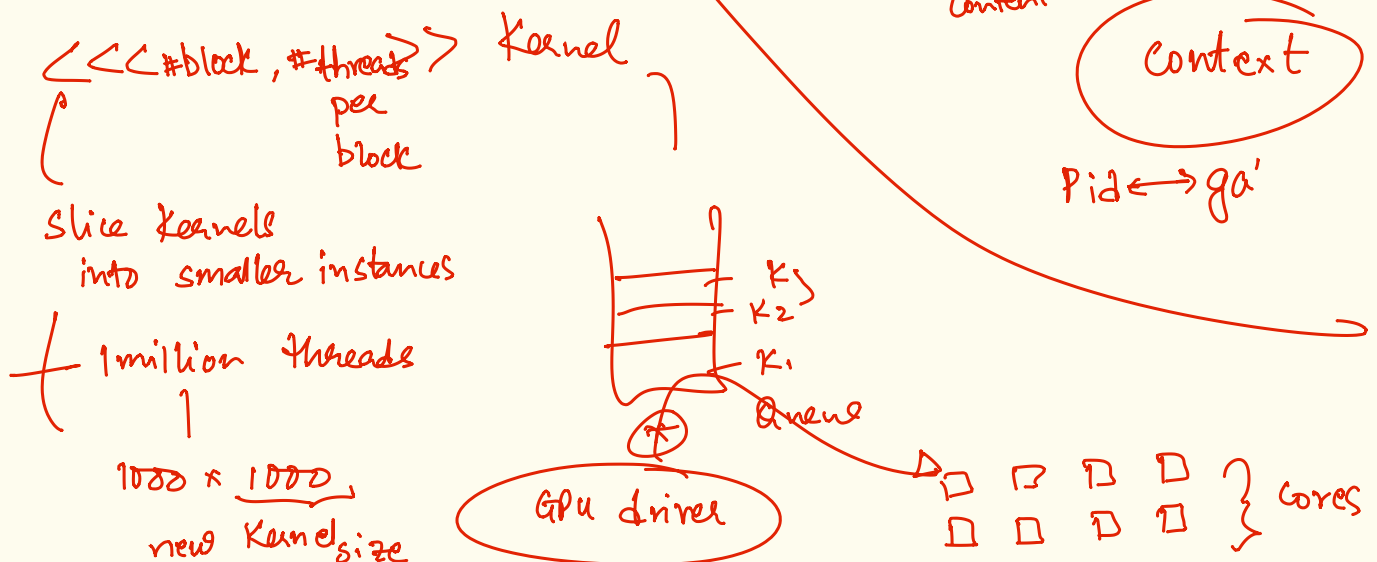
— RPC for GPUs

intercept CUDA calls

rcuda
vrcuda
qcuda
vmcuda



⑤ Kernel Slicing.



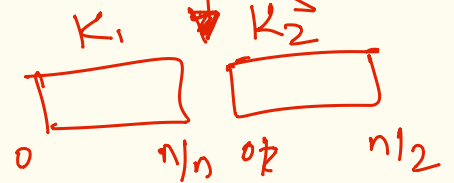
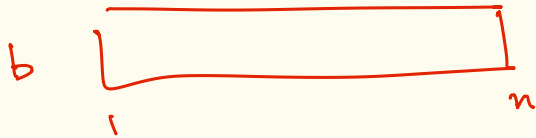
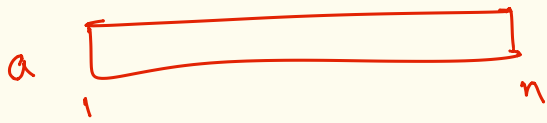
- ⑤ slow down kernel threads per invocation
 & inject ^{artificial} pre-emption at end of each mini-kernel.

for ($i=0$; $i < \# \text{ kernels}$; $i++$) {

$\ll \# \text{ blocks}, \# \text{ tpb} \gg$ Kernel

 slicedid++;
 }

Kernel
add



$\ll \ll 1, 10^6 \gg \text{ add}$

$$= a[tid] + b[tid]$$

$$= \text{tid} = \text{blockid} * \text{tpb} + \text{threadid};$$

$\ll 2, 2 \gg$

blk=0 tid=0
1

blk=1 tid=0
1

⑥ FaaS

└ decouple fn from fn execution

└ GPU functions

└ containers using GPUs
on a server

unikernels ————— VMs ————— isolation.

performance. } Containers
 } processes

Kernel, the libraries
is usually very bloated.

optimized Kernel + stock set of libraries & tools
lean for the target appⁿ.

tradeoff — performance & generalized
runtime execution.

how to build this?

figure out dependencies
& multiple abstractions

enforce them
on the subsystems.

~ My VM is lighter & smaller than your Container.
— SOSP 2017

Idk

objdump

Busybox / debootstrap Debian

} Firecracker VM

Kata Containers

microVMs