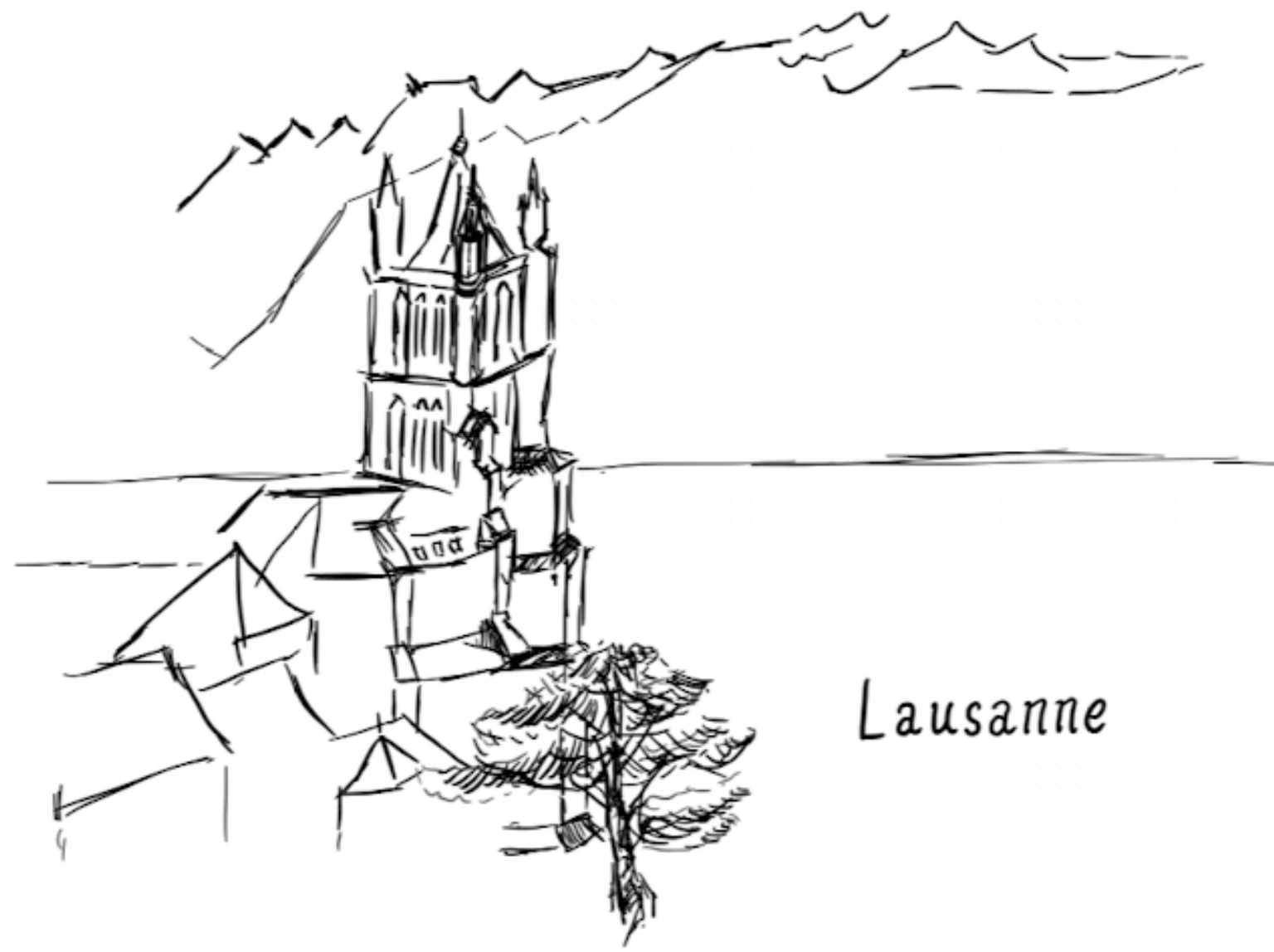


TMO: Transparent Memory Offloading
IOCost: Block IO for Containers

In Datacenters



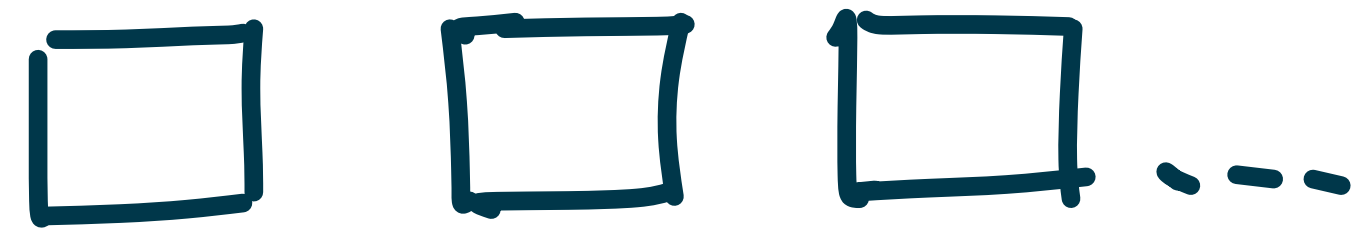
Two papers in ASPLOS 2022

All Drawings by Jing Liu

Memory Offloading

DRAM Apps consume a lot of memory

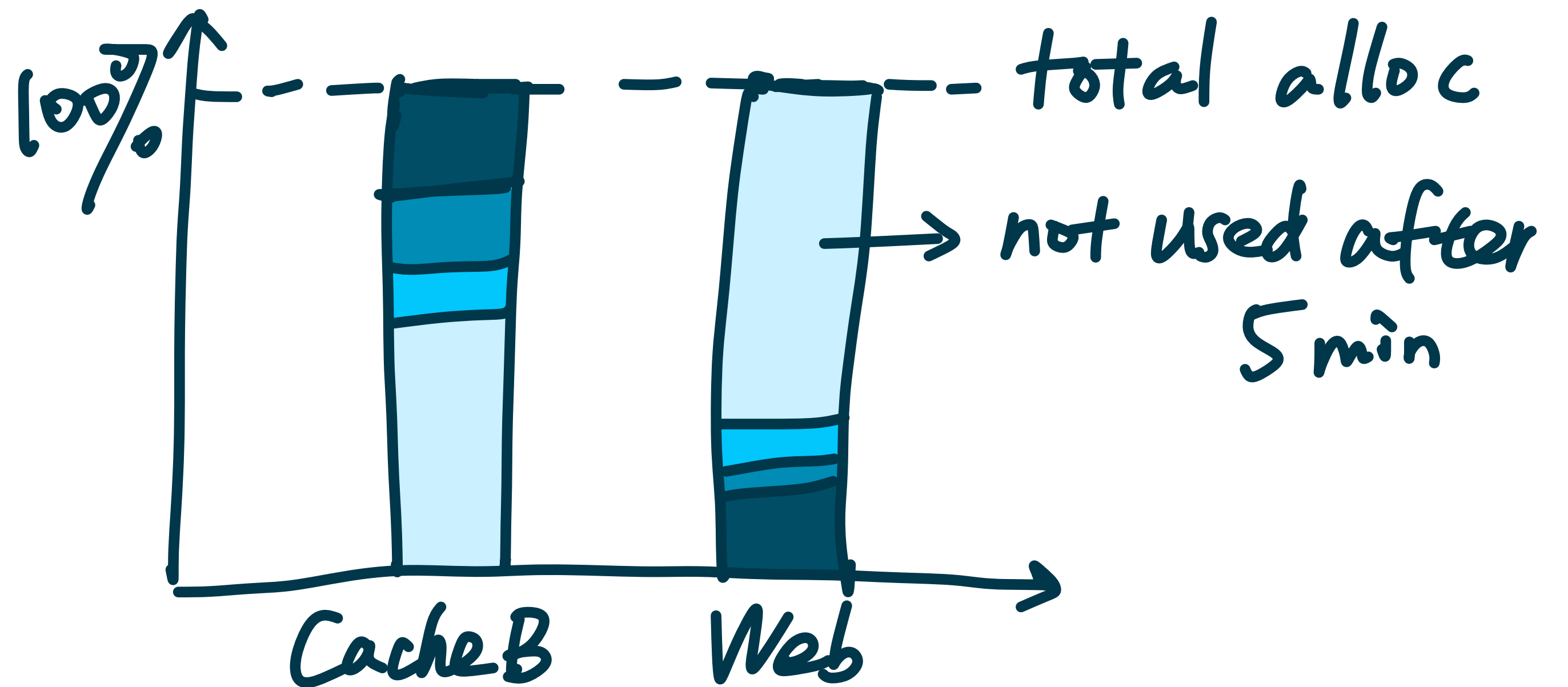
⇓ offload



⊕

Memory tiering

compressed mem
NUMe
SSD



State-of-the-art & Issues

zswap (Google, ASPLOS 19)

① single slow memory tier (compressed mem)

② offline application profiling
metric = page-promotion rate

State-of-the-art & Issues

zswap (Google, ASPLOS 19)

① single slow memory tier (compressed mem)

⇒ add NVMe SSD → cheaper, power

⇒ Compression not applicable (Sometimes)

② offline application profiling

metric = page-promotion rate

⇒ Handle dynamic + diverse apps

⇒ Handle device heterogeneity (diff SSDs)
slow device implies low PPR

App's sensitivity
differs ⇒ e2e
perf

TMO

{ Transparent memory offloading
containerized environment $\Rightarrow$ interesting setting

key questions { How much memory to offload?
What memory to offload?

$\Rightarrow$ mechanism for better information

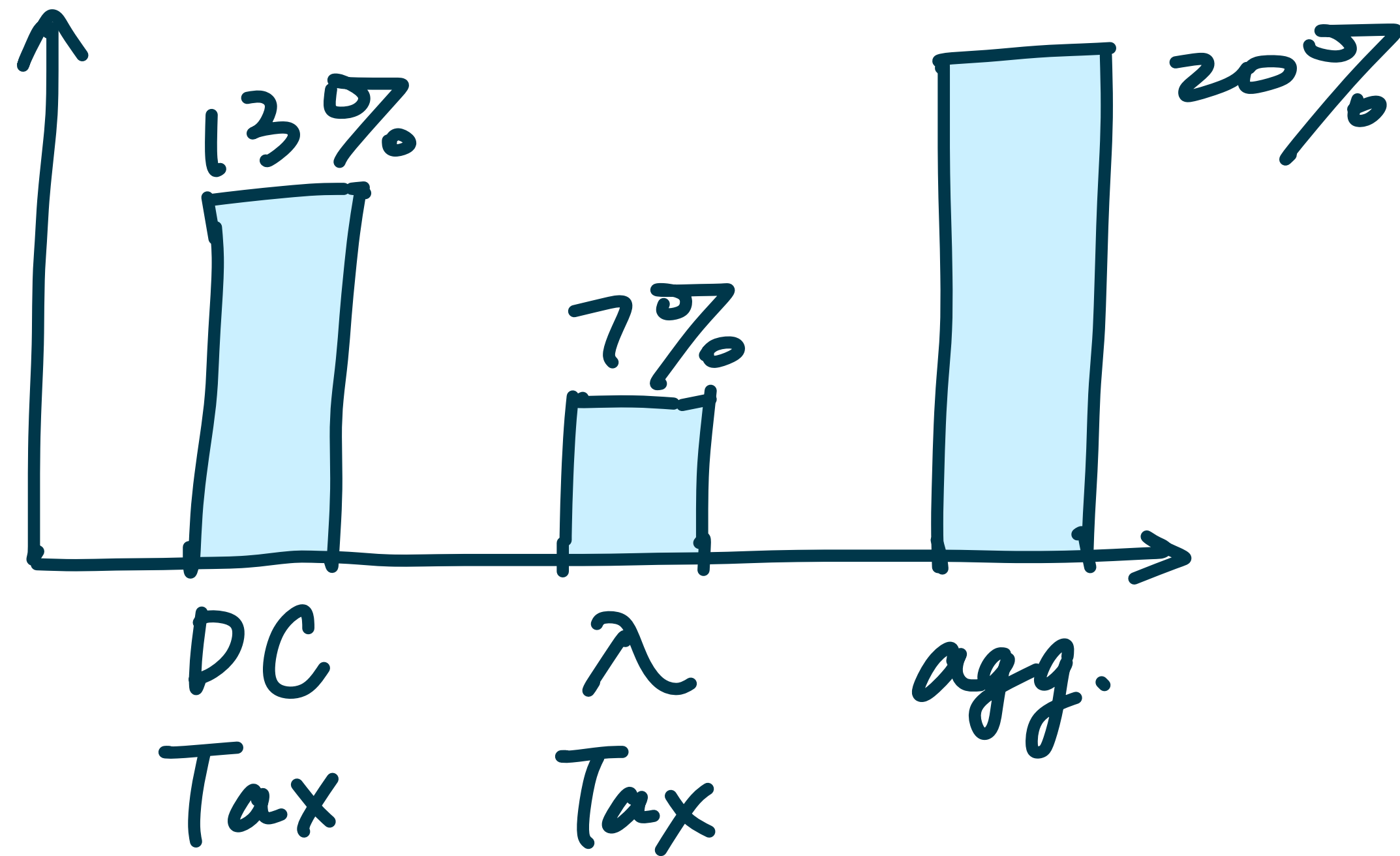
$\Rightarrow$ perf model captures: app-centric slowdown $\oplus$ diff devices

$\Rightarrow$ framework + algorithm to adjust at runtime

Data center Setting / Deployment ⊕ Challenges

Memory Tax = { infrastructure-level functions
microservices
routing, proxy...
packaging, logging
profiling...

Memory %



memory tax SLA = relaxed

⇓
primary target of
offloading (first version)

Data center Setting / Deployment ⊕ Challenges

HW heterogeneity of offload backend

Different types of NVMe SSDs { perf
generations { endurance

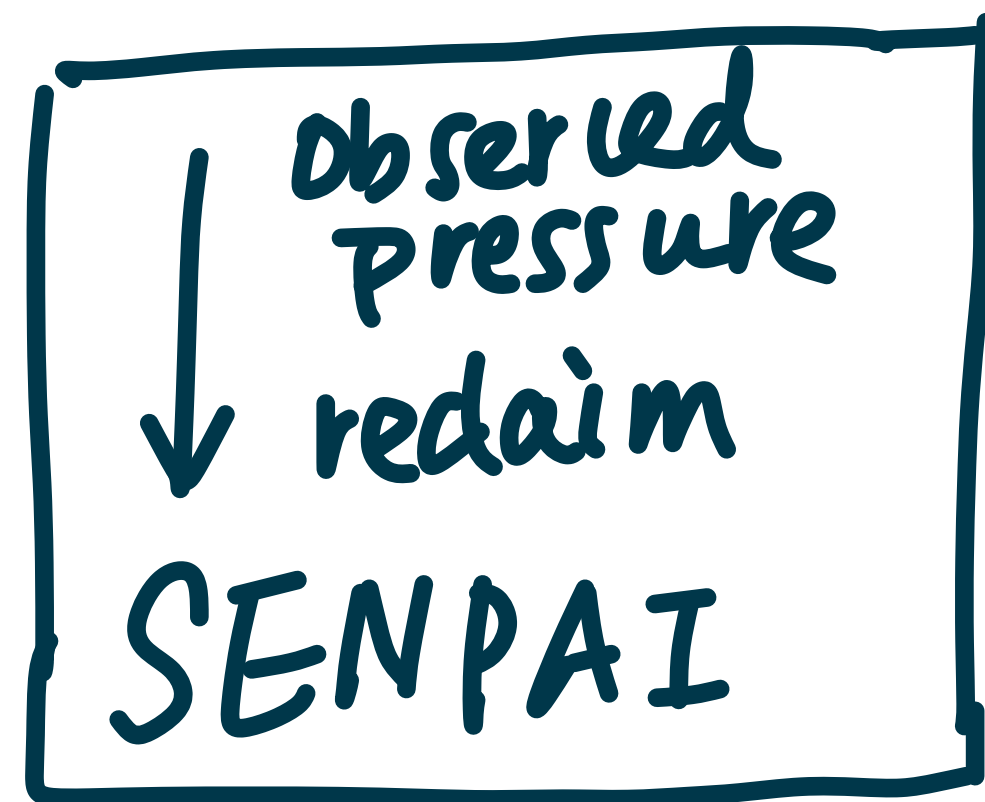
eg. R/W latency: 9.3 ms $\rightarrow$ 470 μ s

important *

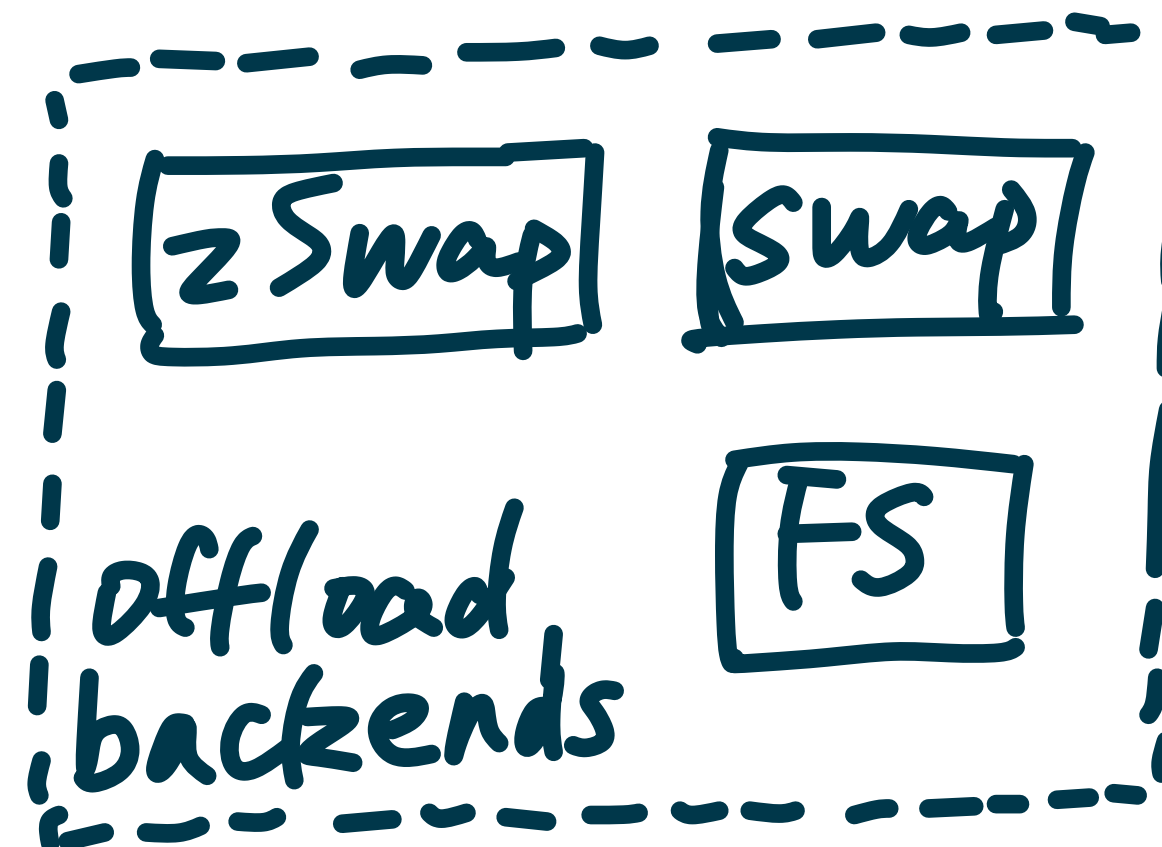
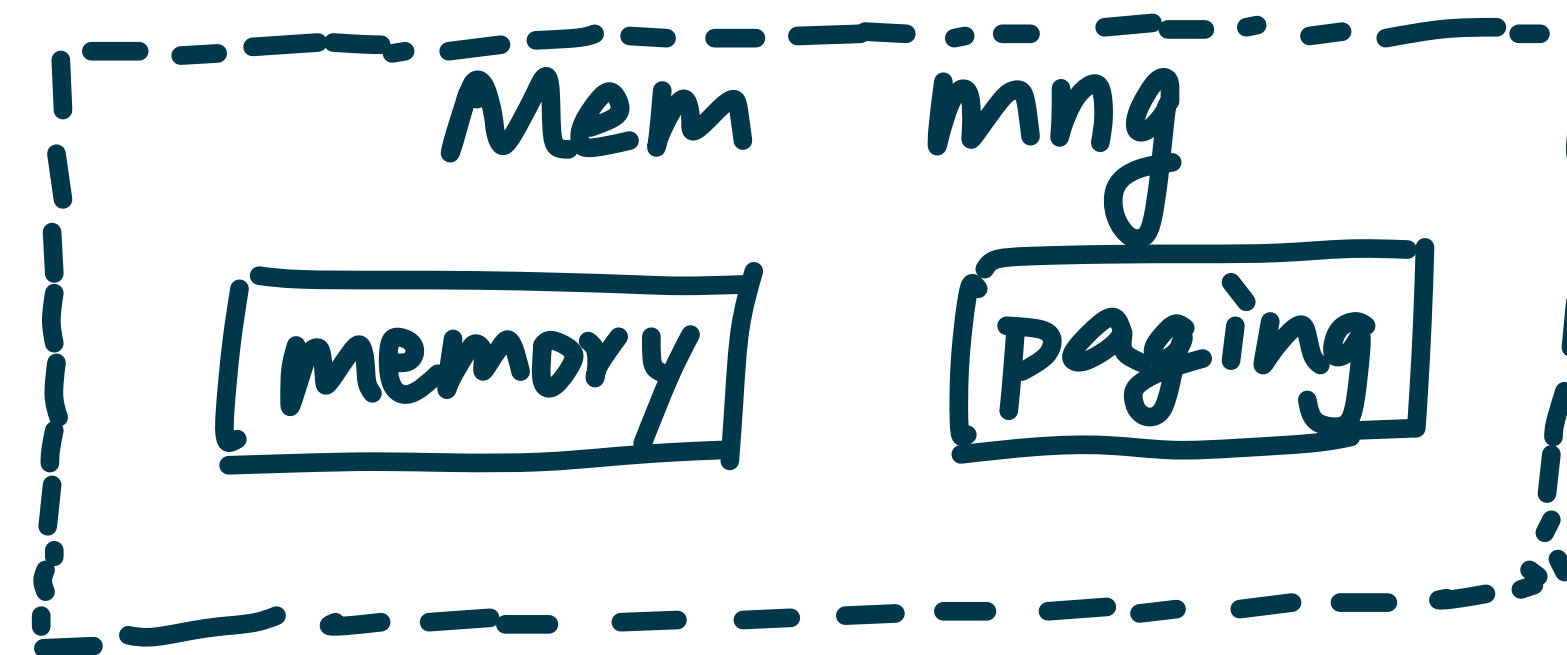
Compressed memory: faster!

TMD Design

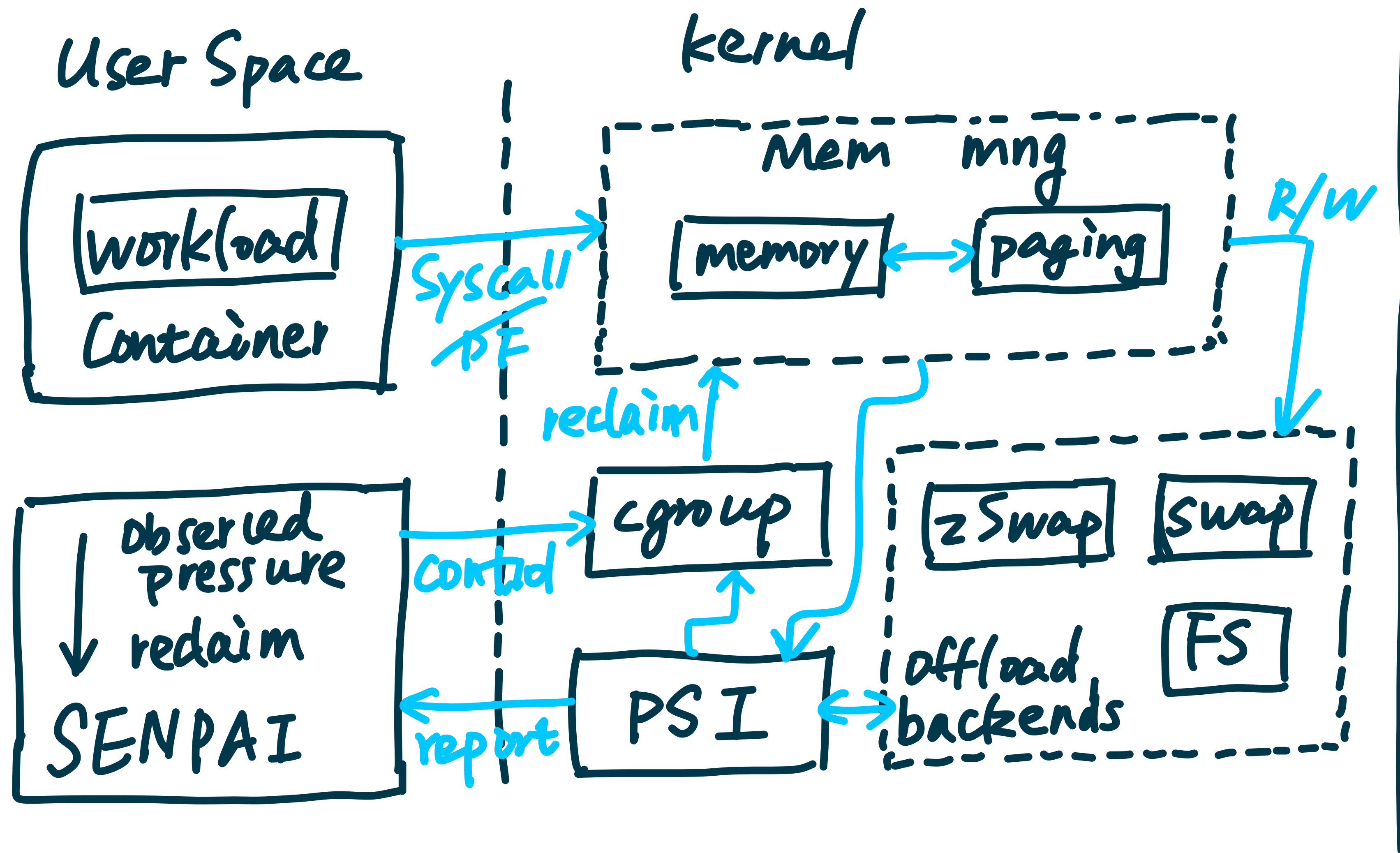
User Space



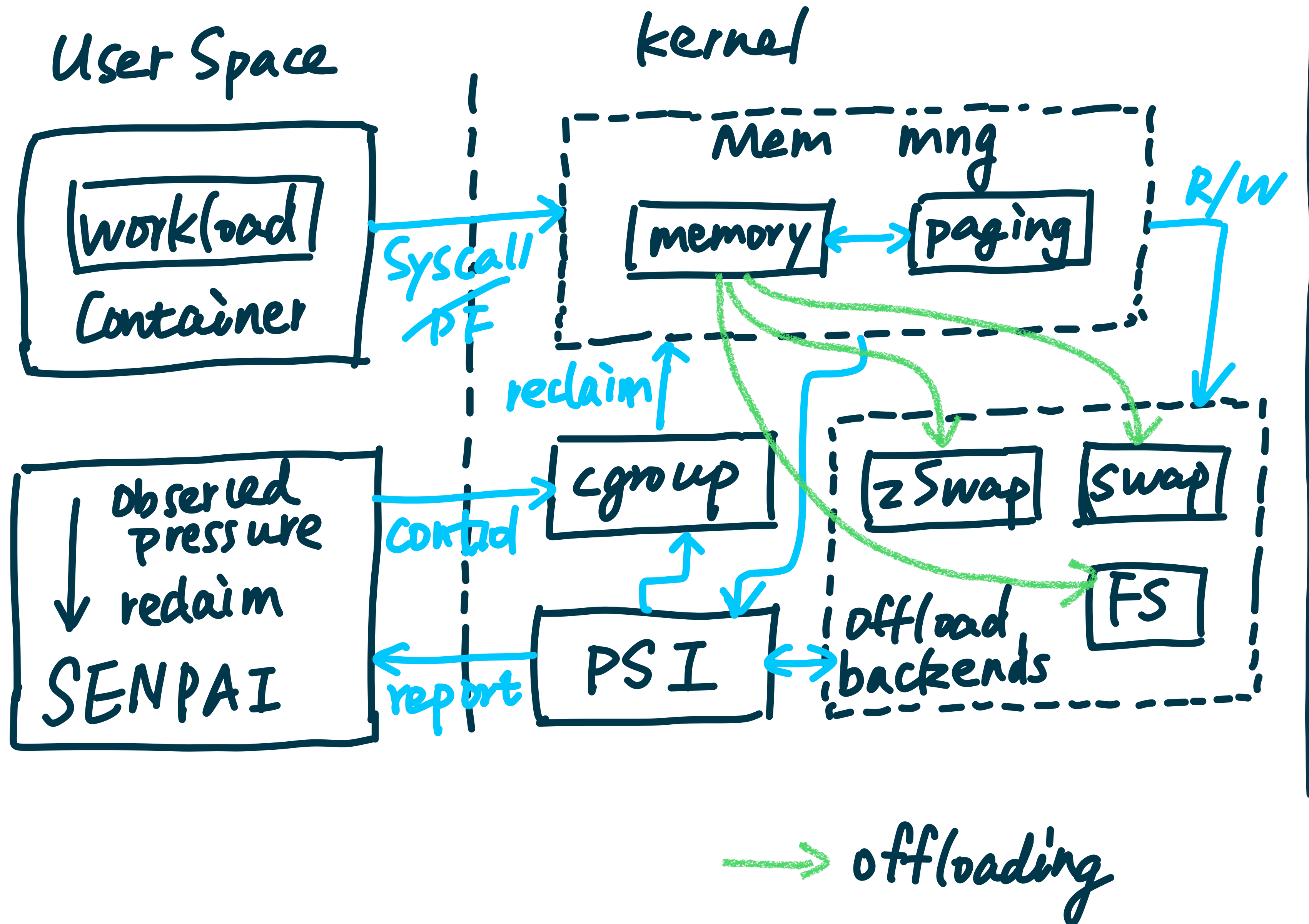
kernel



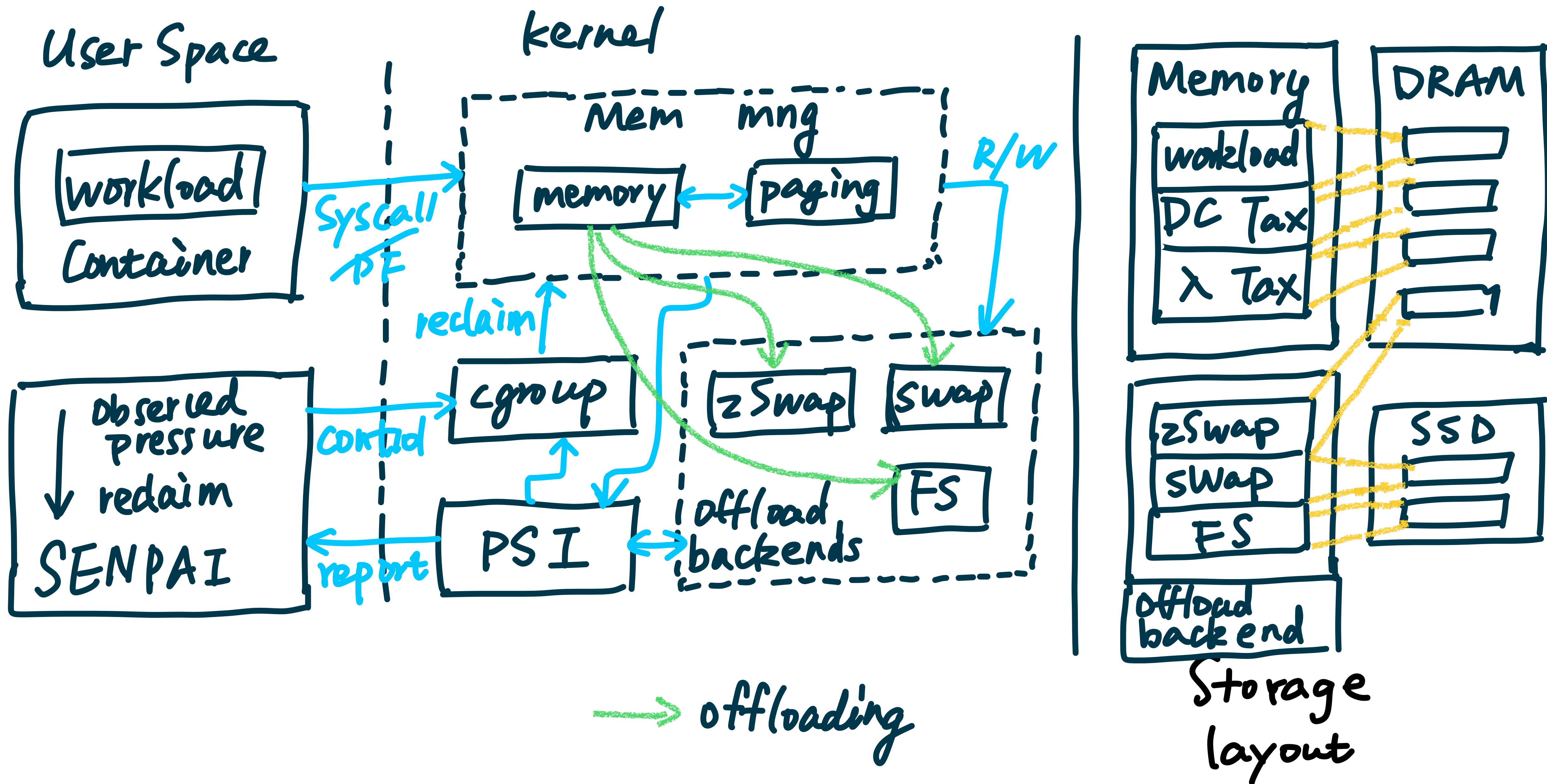
TMD Design



TMD Design



TMO Design



Resource Pressure: Pressure Stall Info (PSI)

How much memory to offload?

metric: predict the resulting performance after offloading

why hard? { kernel's information is indirect $\Rightarrow$ mechanism support $\oplus$

{ single metric is hard

accurately attribute performance penalty to offloading

eg. workload start up

workload working set transition

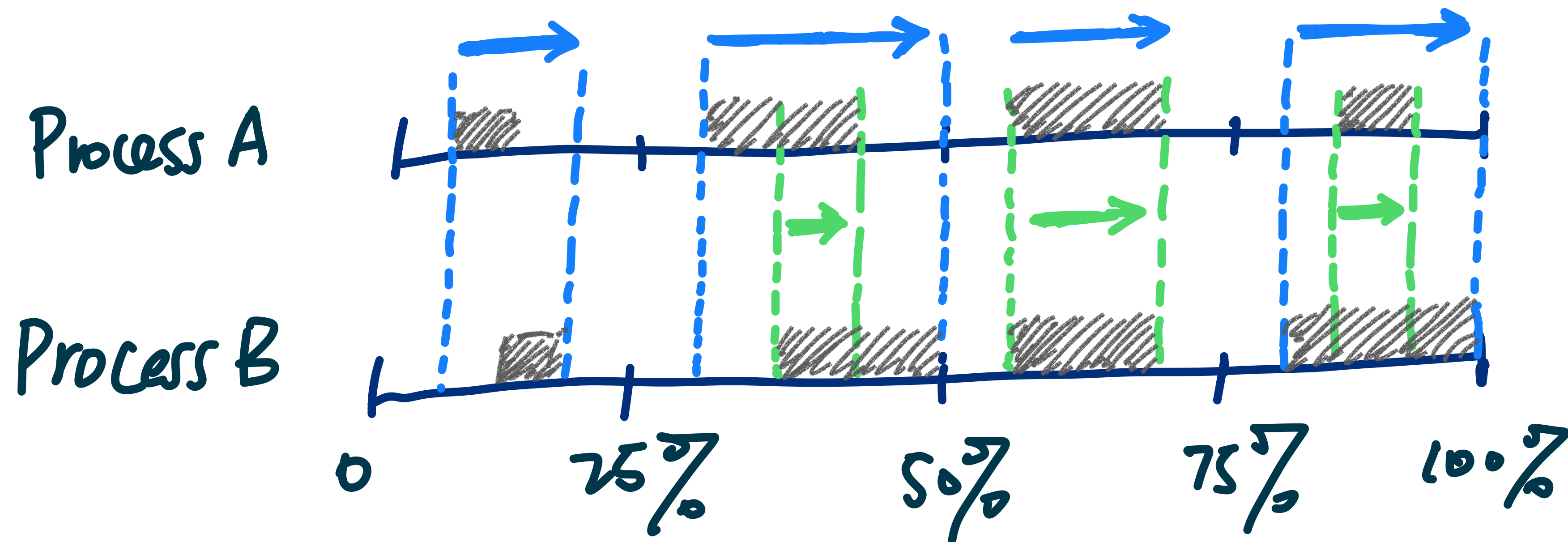
...

mechanism to provide information & accounting to container (in kernel)

Resource Pressure: Pressure Stall Info (PSI)

How much memory to offload?

metric: predict the resulting performance after offloading

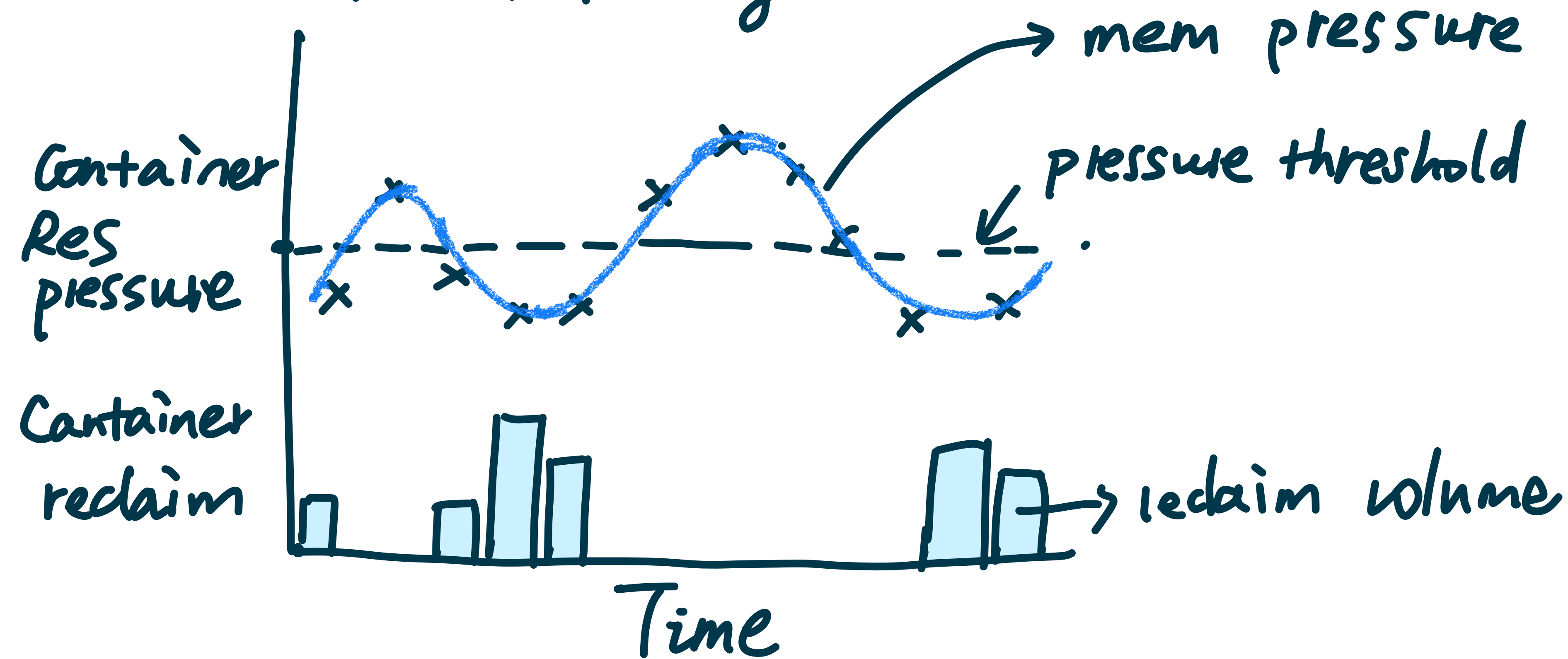


→ Some: One process stall for resource
→ Full: All processes stall for resource
Stall

Determine Memory Requirement of Containers

How much memory to offload?

x PSI tracking



Senpai: userspace process
PSI as feedback

cgroup
config

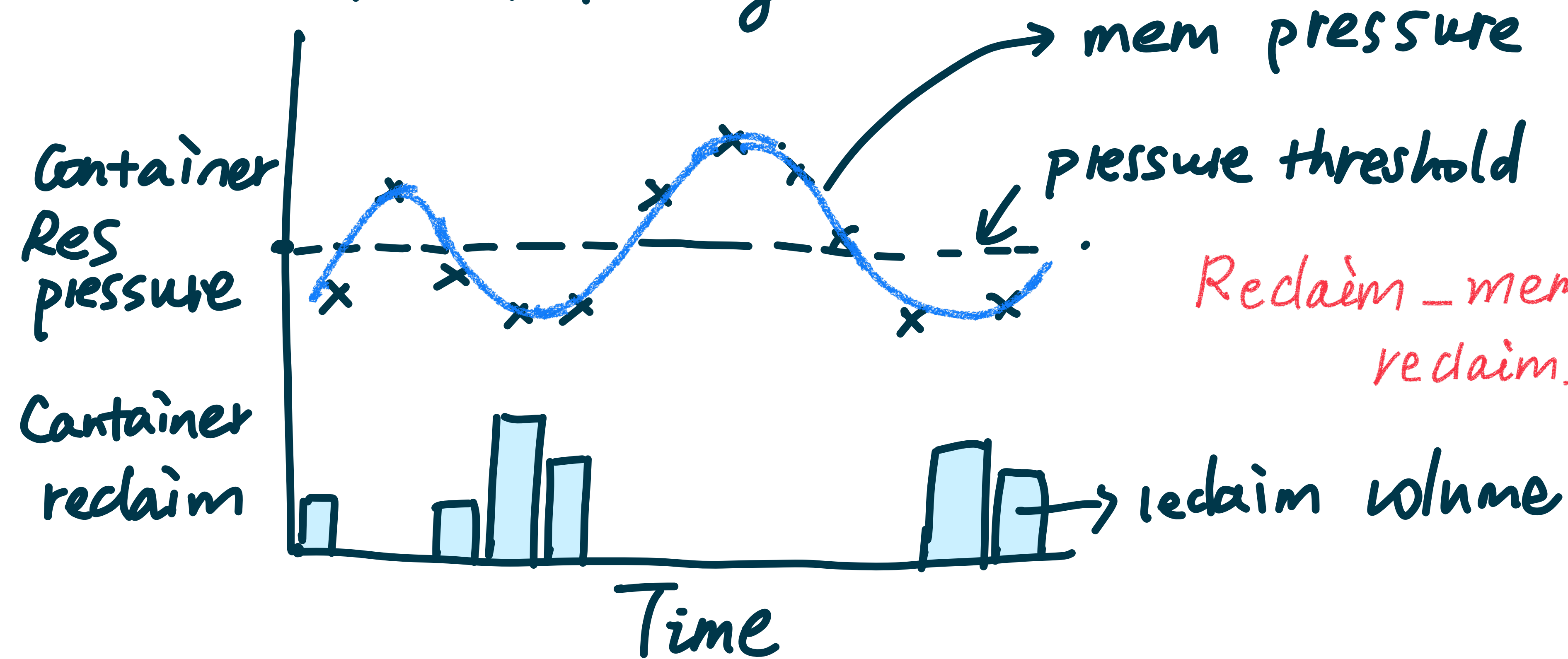
kernel

mem PSI
IO PSI
SSD endurance
CPU PSI
not significant

Determine Memory Requirement of Containers

How much memory to offload?

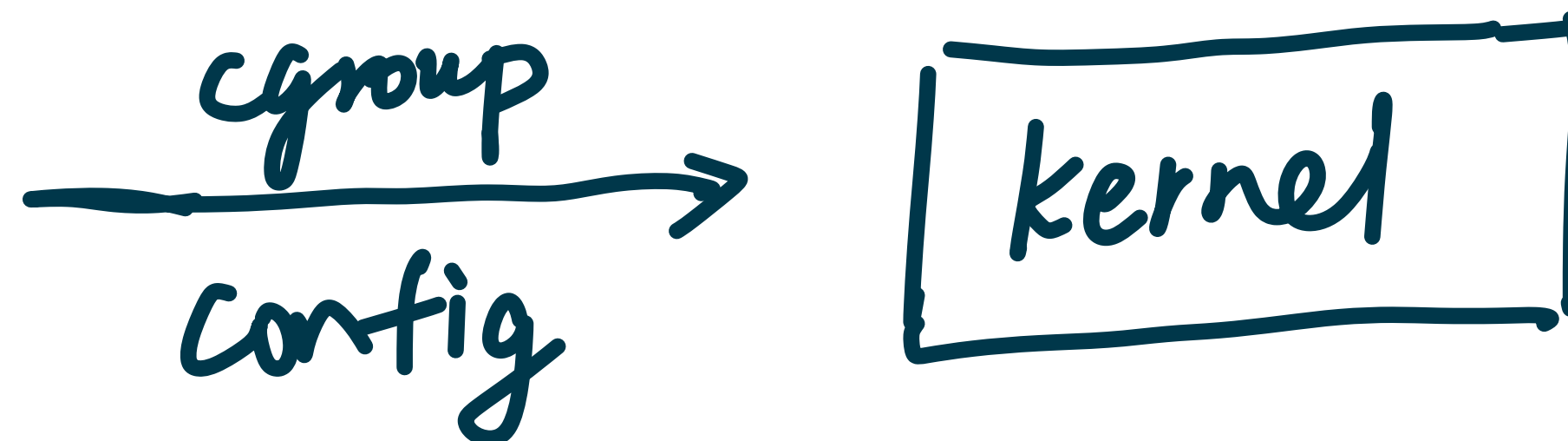
x PSI tracking



$$\text{Reclaim_mem} = \text{Current_mem} \times \text{reclaim_ratio}$$

$$\text{reclaim_ratio} = \max\left(0, 1 - \frac{\text{PSI}_{\text{some}}}{\text{PSI}_{\text{thr}}}\right)$$

Senpai: userspace process
PSI as feedback



Kernel Optimization for memory offloading

What memory to offload?

mechanism: kernel's reclaim

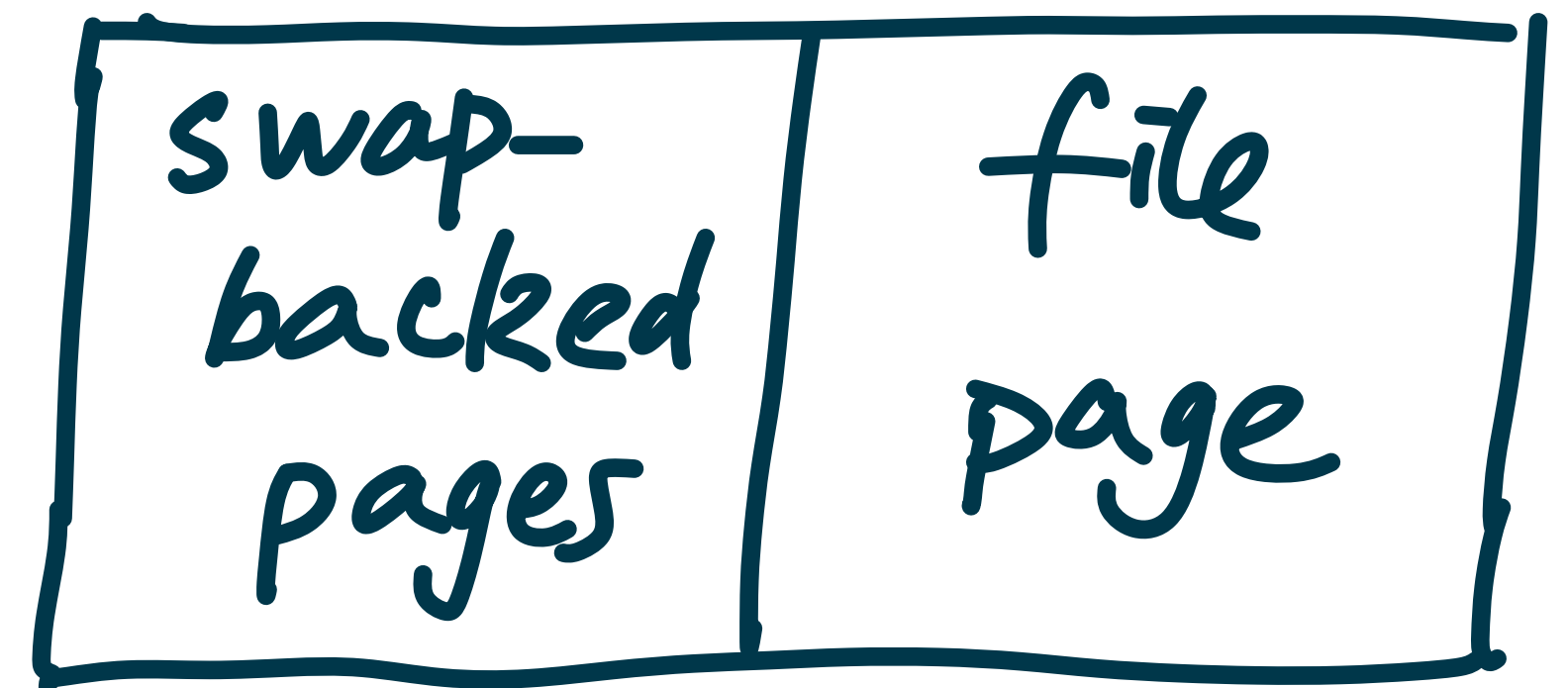
issue: kernel's swapping is very conservative
skewed heavily towards file cache

+ refault-detect mechanism

non-resident cache tracking

⇒ reuse-distance

⇒ exclude first-time-accessed
file cache



+ if refault happens $\left\{ \begin{array}{l} \text{refault rate} \\ \text{swap-in rate} \end{array} \right\}$ balance

TMO - Eval: Performance + memory saving

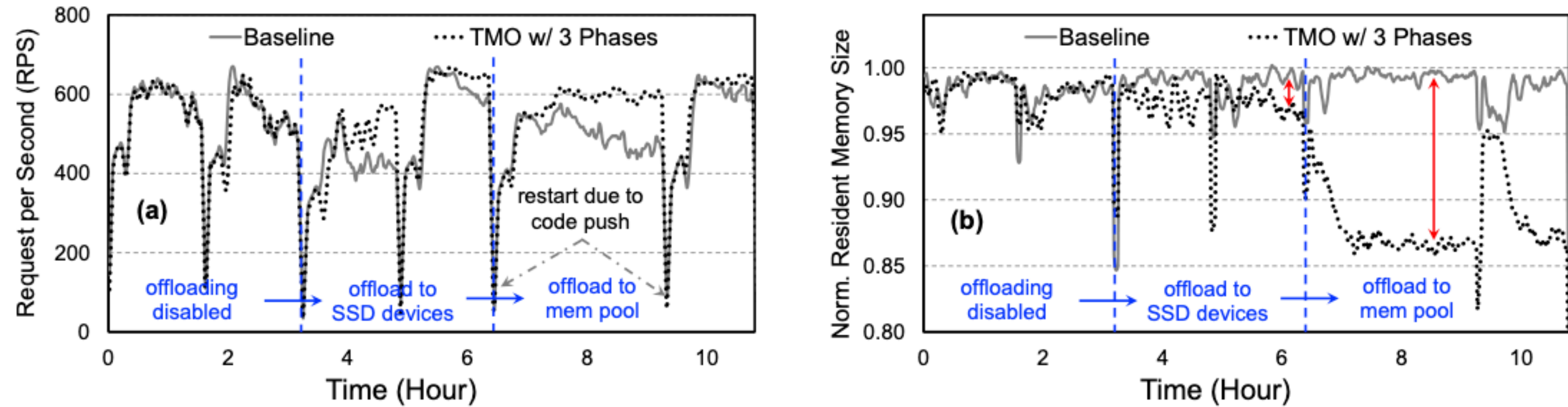


Figure 11: Performance improvement and memory saving for the Web application on memory-bound hosts. We set up two tiers: started with identical configuration without any swap enabled and later enabled SSD & then compressed memory offloading on the 2nd tier.

TMO - Eval: PSI

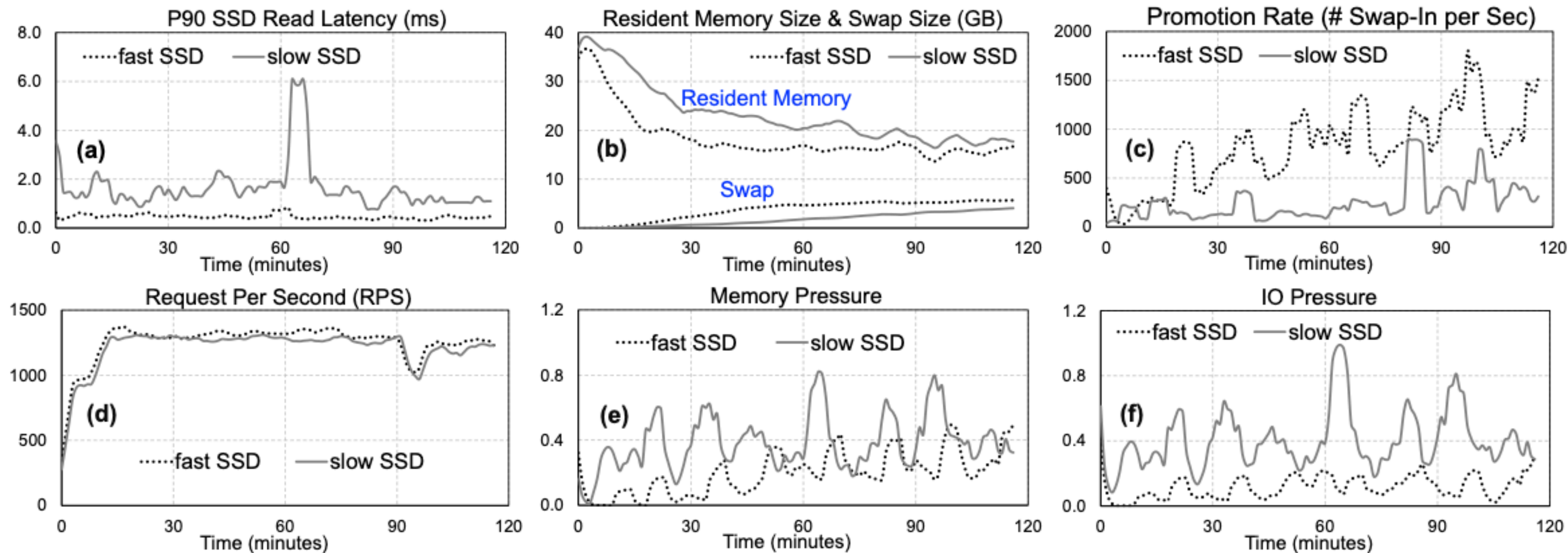


Figure 12: The Web application under the control of TMO with fast and slow SSDs. The solid-gray configuration has lower RPS/performance but actually lower promotion/swap-in rate; while its PSI metrics are indeed higher.

TMO - Eval: Device Endurance

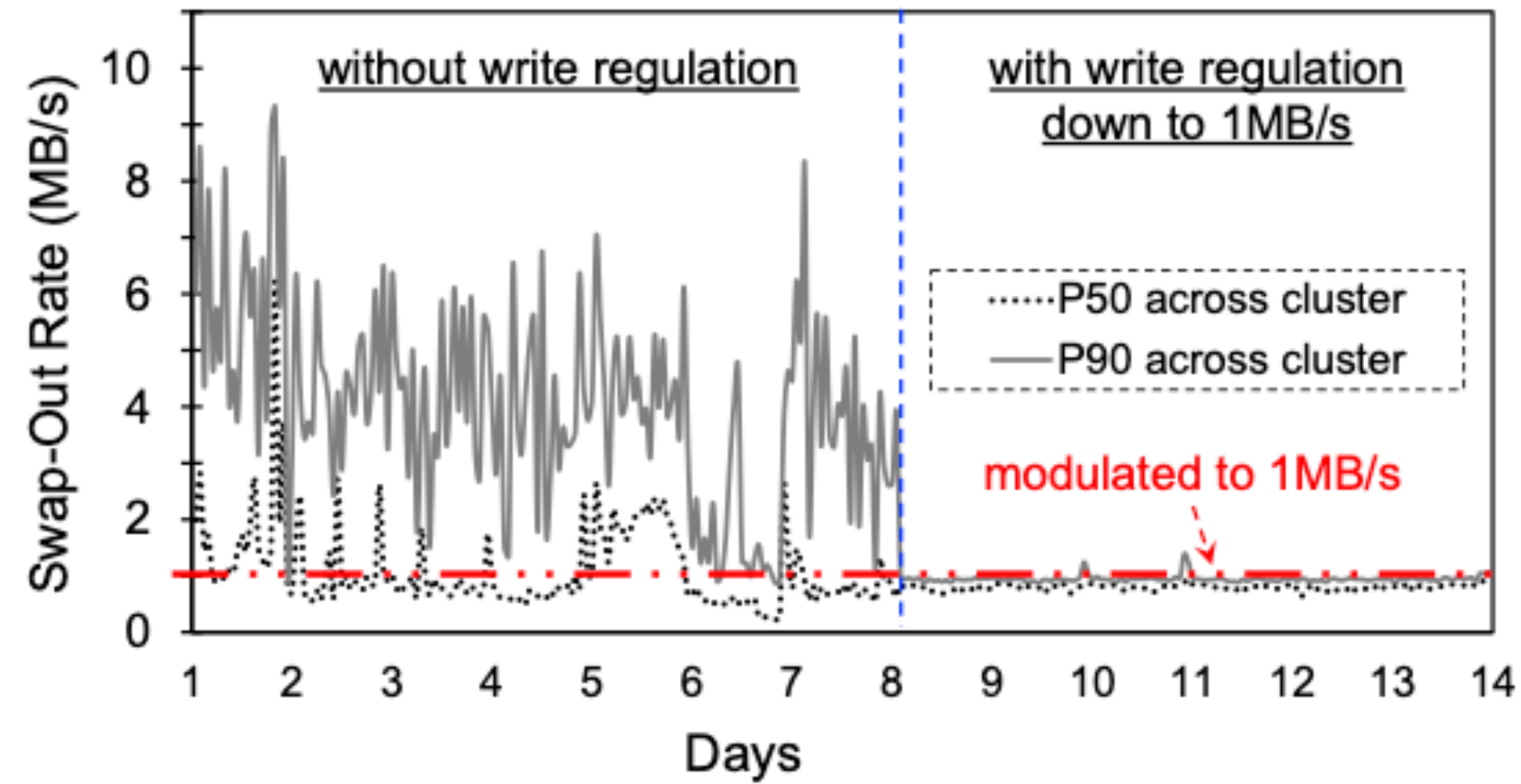
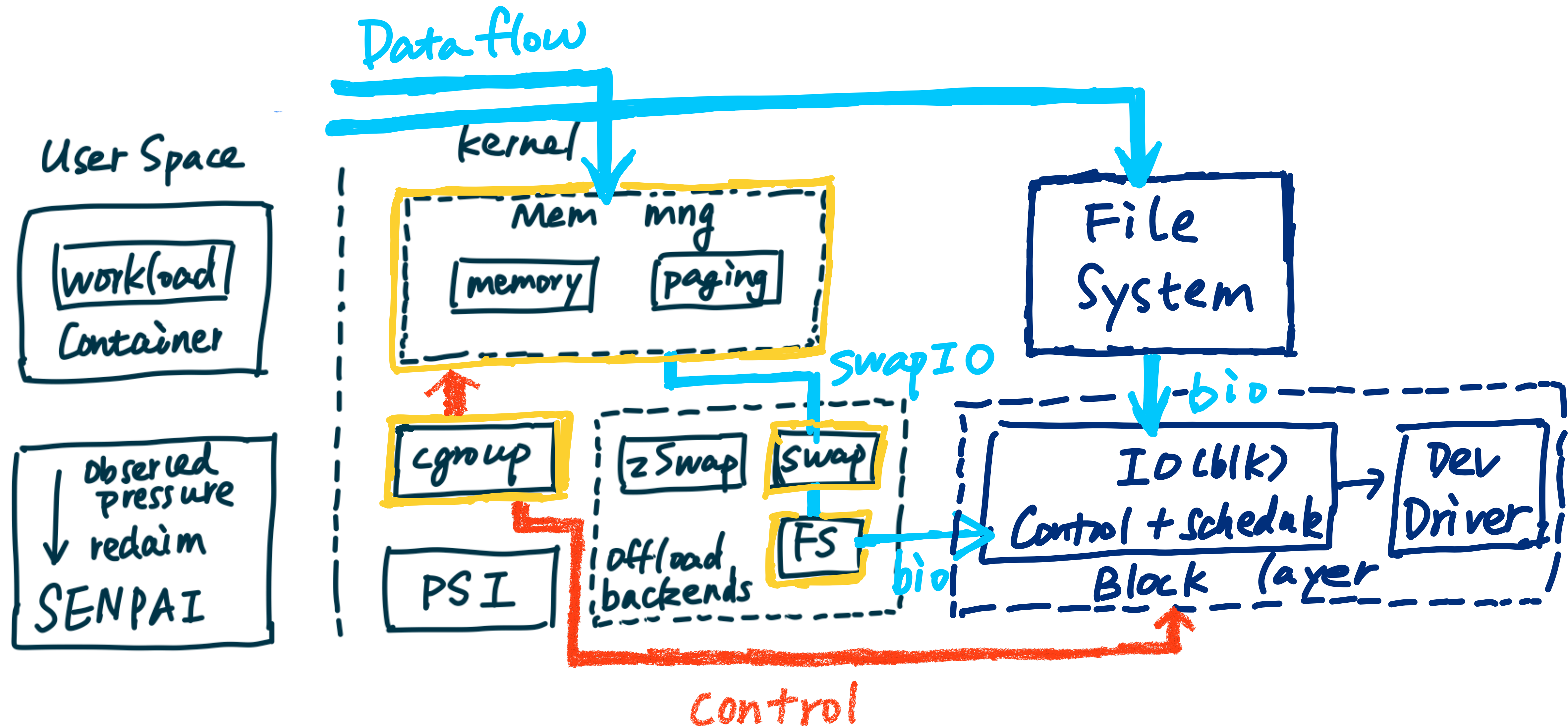


Figure 14: Swap rate with and without write regulation.

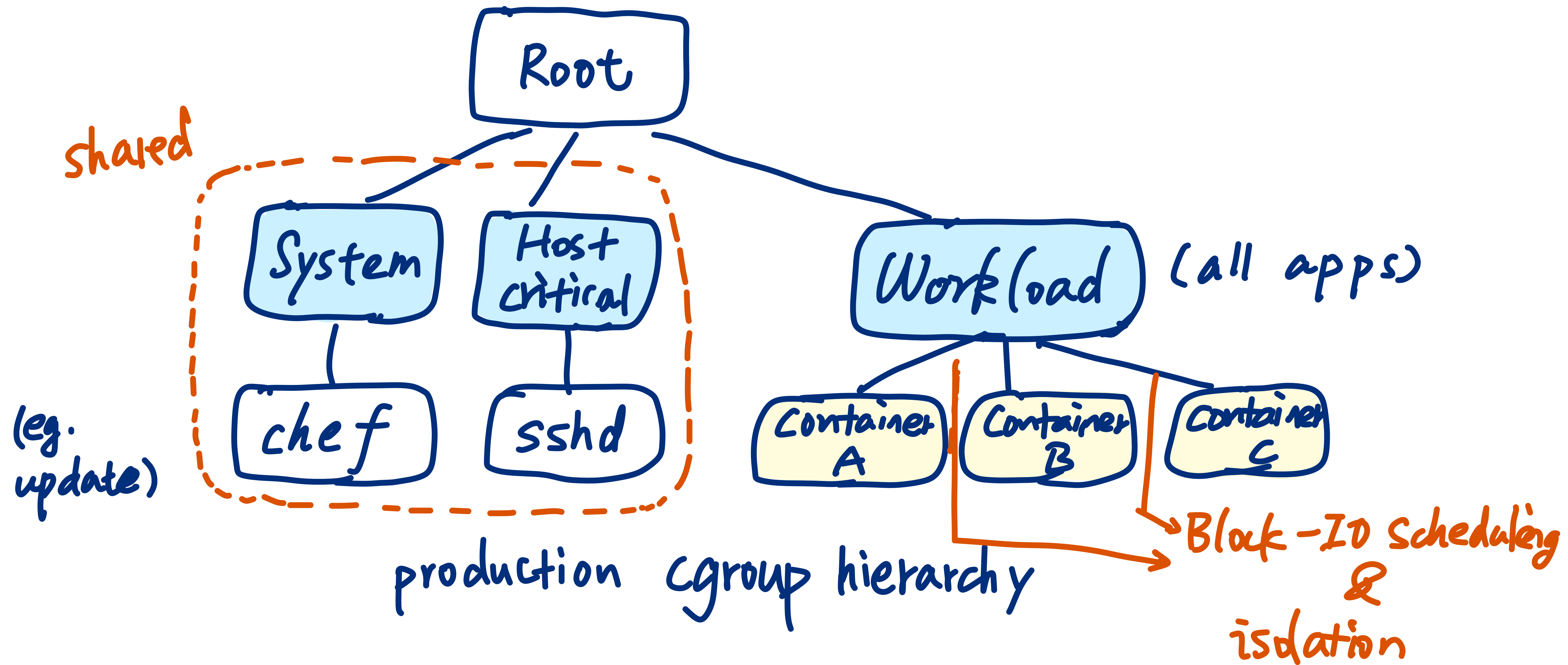


IO cost: Block IO control

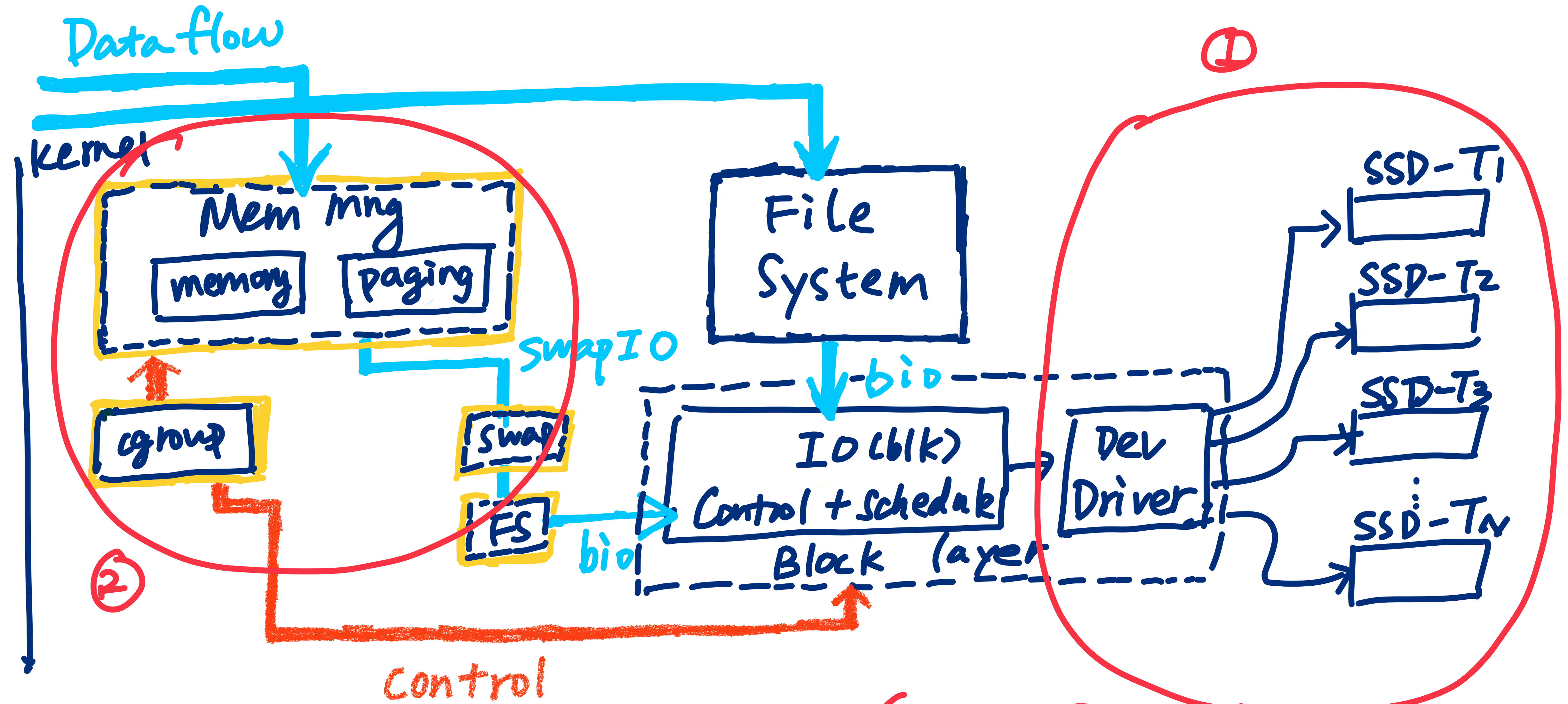
DC containers

heterogeneity in Devices

Data center Setting / Deployment ⊕ Challenges

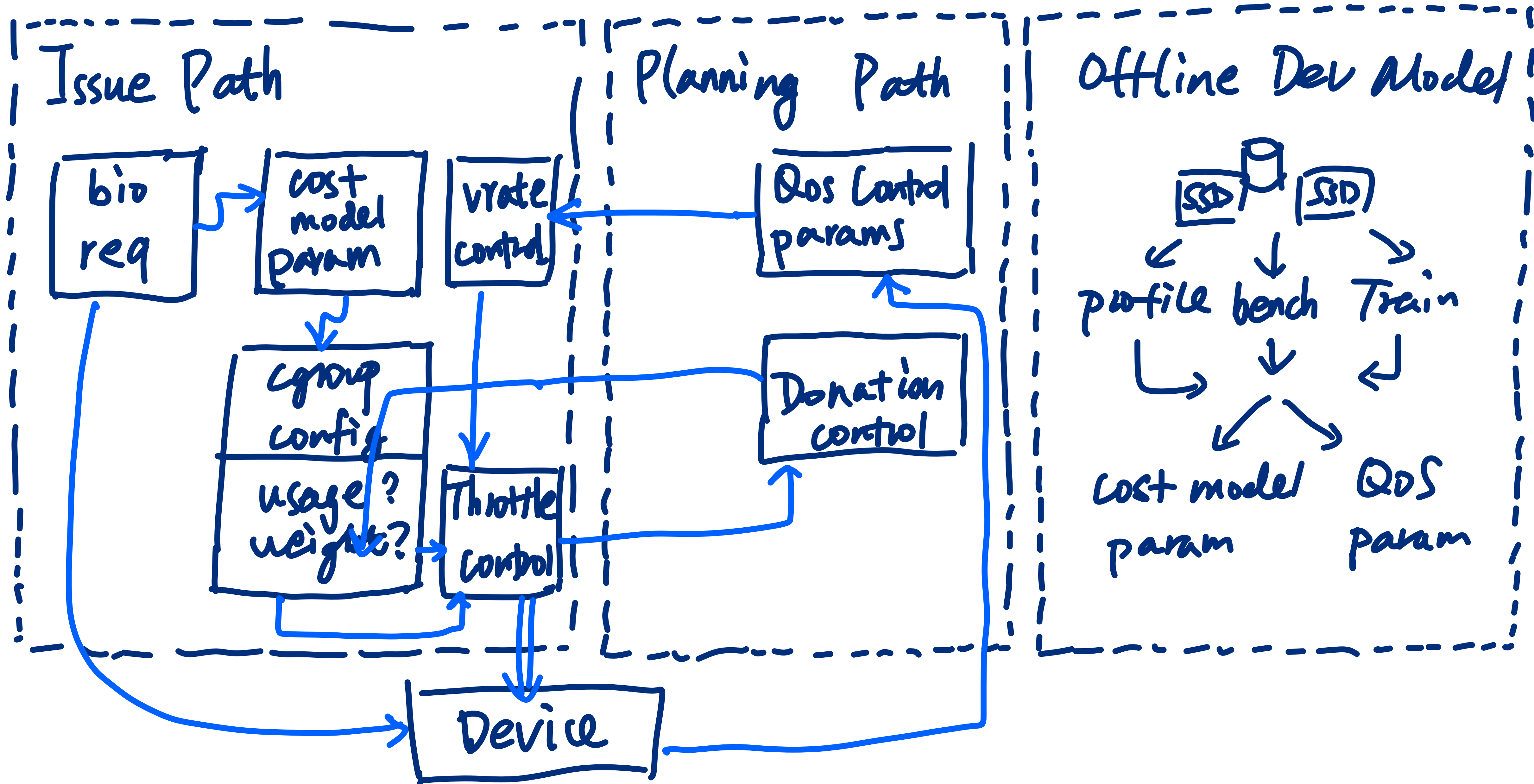


Data center Setting / Deployment ⊕ Challenges



(2) Interaction with memory subsystem (mem + IO control make coherent decision)

IO Cost Overview



Summary of IO Cost

⇒ Offline performance model of device (heavyweight model)

⊕ runtime block IO control (live model) under dynamic workload ← ⊕

{ fast path : io operation / policy enforcement
slow path : policy decision

Summary: TMO + IO cost

⇒ Mechanism support for more control in kernel/
+ cross component coordination

⇒ userspace process for policy level decision

⇒ Modeling the device from different aspects
+ online approximate + query

⇒ more direct app-performance metrics after certain
resource control: $\Delta \text{res} \Rightarrow \Delta \text{perf}^{\text{APP}}$

↓
levels, a lot of components, shared infra
eg. sidecar container