

Lecture 13

CS695

- #2 due today

- Midsem exam. 1st March

- 2:30 pm

LA ???

- all content till today
(and some more)

⊗ [memgmt]

[livemigration]

- snapshots

- records & replay

Memory Resource Management
in VMware ESX Server.

OSDI 2002

operating system
design &

implementation

↳ Carl Waldspurger (vmware)

Context:

~ overcommit & server consolidation

pack multiple VMs
on a single server
PM

↳ committed resources greater than physical resources.

v p m

$\sum \text{P}_{\text{ranges}} > m_{\text{range}}$
of vms

~ vmware ESX server

- "management" actions:

~ allocate memory / generate mapping

⊗ - swap in / swap out (revoke / reclaim)

⊗ - vmmap / de-allocate.

- handling faults — revocation

- free page state

- permissions handling

- sharing / cow

fully virtualized
hypervisor

Shadow pages
tables
for memory virt.

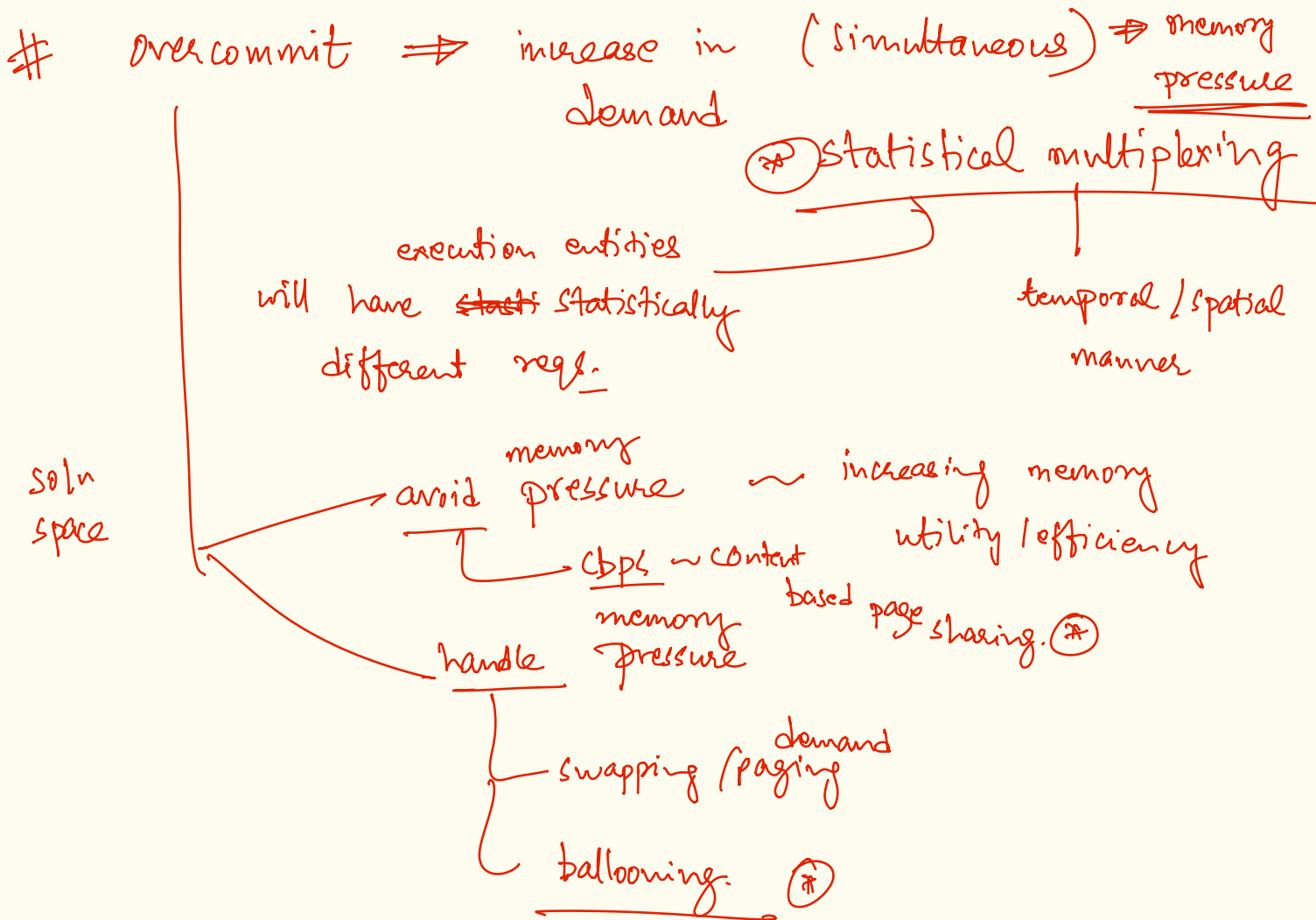
- 2 page tables

v2p

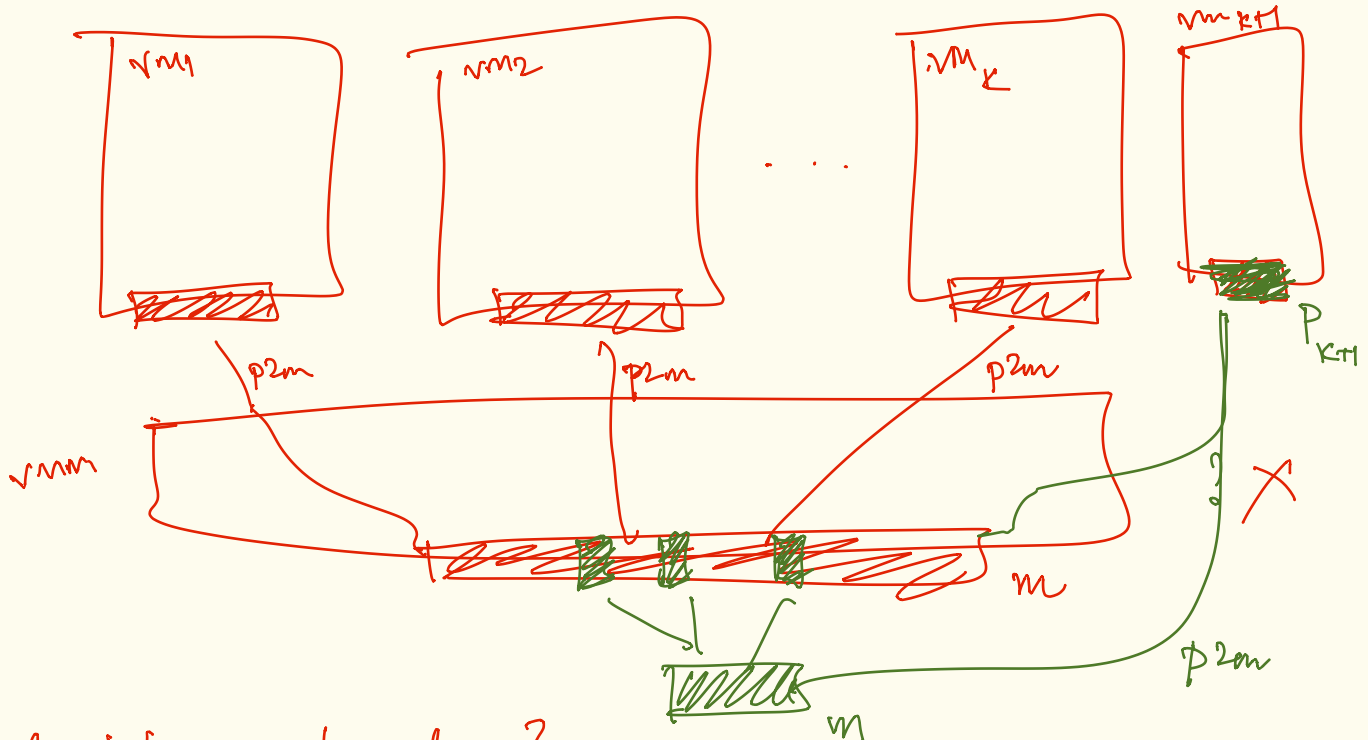
v2m

⊗ pmap

p2m mappings
per VM.



(i) Ballooning



- which 'in' pages to reclaim?

- how many 'in' pages per vm?

$\Rightarrow$ collect meta-information about accesses to estimate per-page utility.

high overhead

incomplete! / wrong decisions.

⊗ double-paging problem — (i) VMM swaps out a m -page invalidates $p2m$ mapping
(ii) VM decides to swap out the same p page

⊗ solⁿ — memory ballooning.

⇒ rebuild $p2m$
⇒ swap back p to virtual disk.

idea: guest knows state of ^{its} memory utility.

approach ~ co-operative guest OS & balloon inflate/deflate.

~ balloon driver (LKM) in the guest

~ interacts w/ a virtual/pseudo device

~ device IO is a priv. channel for the VMM ↔ balloon driver.

↑ communicate actions
↓ transfer 'p' info

increase/decrease
memory pressure
in the guest.

balloon driver

inflate

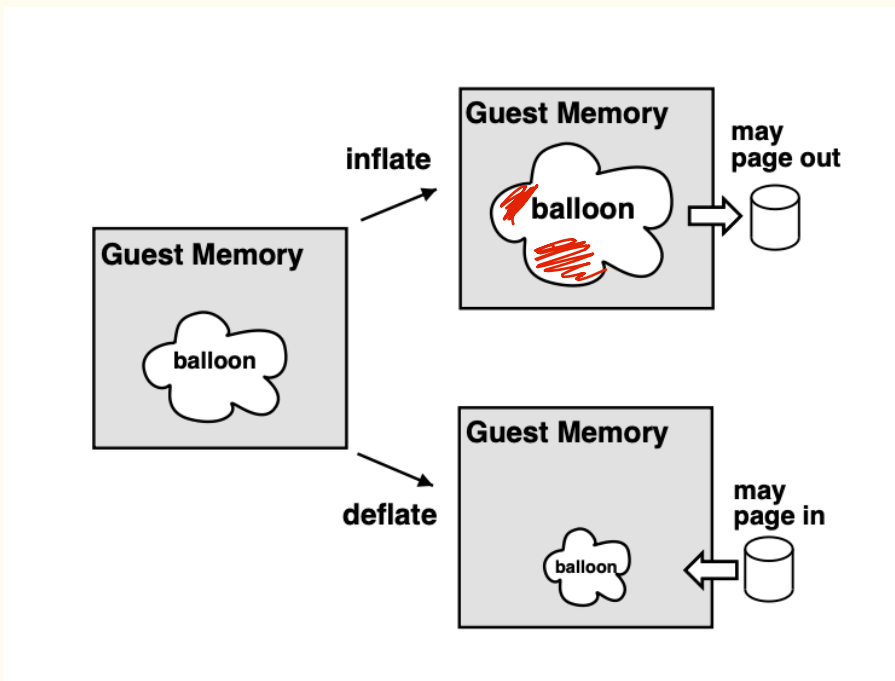
for
- alloc memory balloon (kernel) guest OS

- "pin" memory of balloon driver

- convey p -range of balloon to VMM

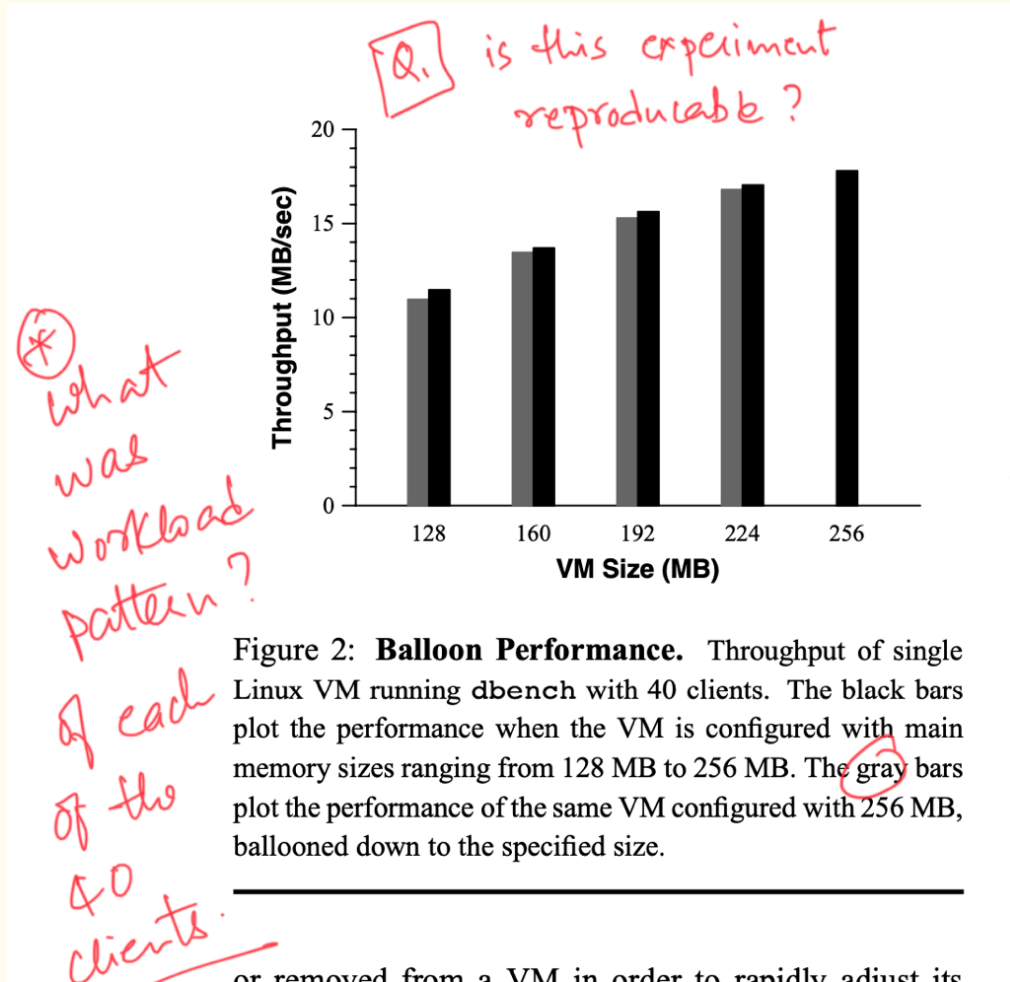
- invalidate $p2m$ mappings at VMM

~ ~~is~~ move free pages in the pmap



Ⓚ how much to balloon per VM?

↳ effective memory size of a VM.



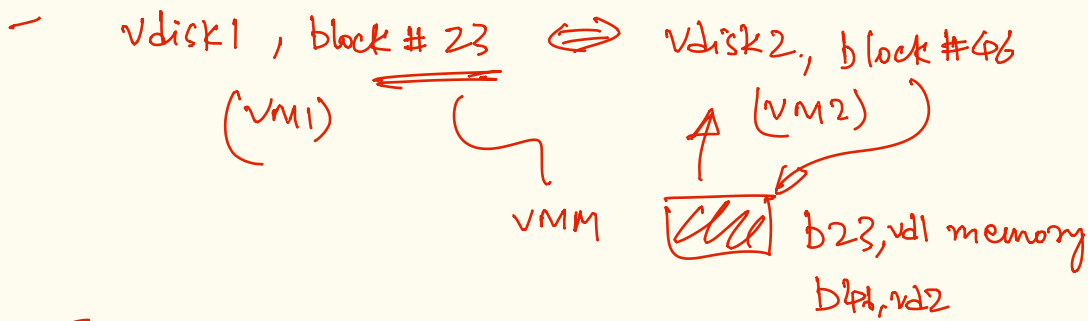
② Sharing memory (avoid pressure)

~ VMs — some similar guest OS, applications, libraries, data.

~ explicit tagging of shareable content.

~ copy — block IO ~ hook in VM. ~ falls into an unified disk cache at the VMM. block

read/write



Cbps: content based page sharing

~ page identity $\Leftrightarrow$ content.

~ scan all pages
 ~ share identical pages
 } out-of-band
 ~ $n \times n \sim O(n^2)$ — $\times$ per page scan
 - hash based technique

(iv) if entry does not exist

L add hint entry
 to the hash

(iii) if entry exists

+ rehash the entry
 + byte-by-byte comparison to confirm match
 + if match
 L share
 L cow

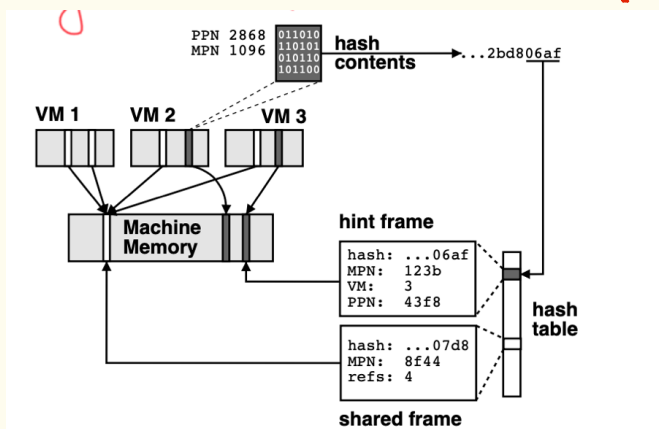


Figure 3: **Content-Based Page Sharing.** ESX Server scans for sharing opportunities, hashing the contents of candidate PPN 0x2868 in VM 2. The hash is used to index into a table containing other scanned pages, where a match is found with a hint frame associated with PPN 0x43f8 in VM 3. If a full comparison confirms the pages are identical, the PPN-to-MPN mapping for PPN 0x2868 in VM 2 is changed from MPN 0x1096 to MPN 0x123b, both PPNs are marked COW, and the redundant MPN is reclaimed.

Q ~ scan rate? MADNESS

~ where to scan?

KSM — Linux Kernel Same Page Merging

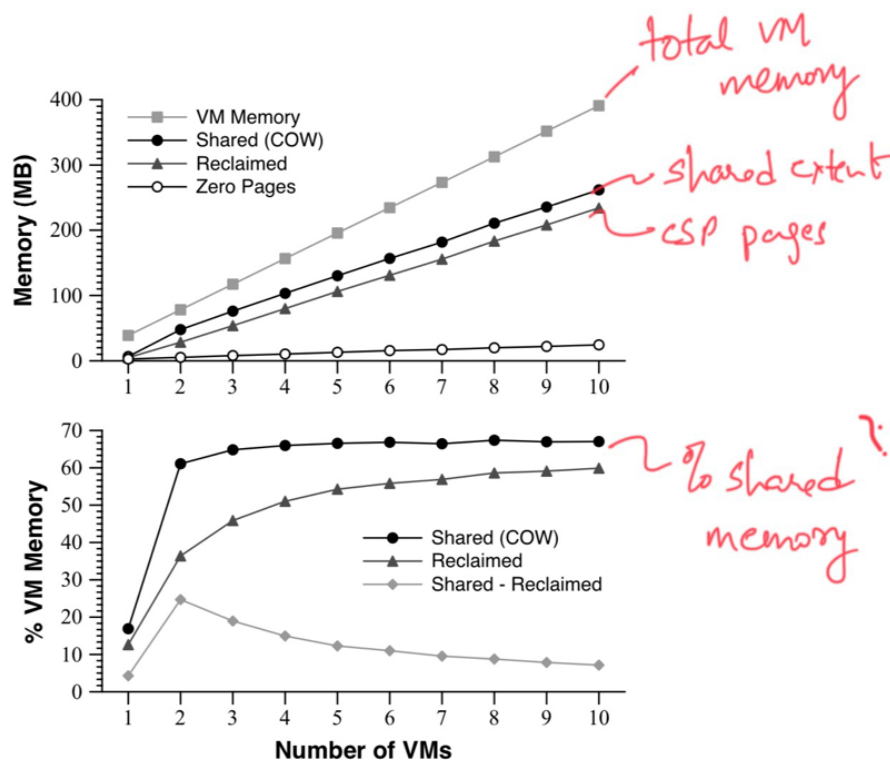


Figure 4: **Page Sharing Performance.** Sharing metrics for a series of experiments consisting of identical Linux VMs running SPEC95 benchmarks. The top graph indicates the absolute amounts of memory shared and saved increase smoothly with the number of concurrent VMs. The bottom graph plots these metrics as a percentage of aggregate VM memory. For large numbers of VMs, sharing approaches 67% and nearly 60% of all VM memory is reclaimed.

	Guest Types	Total	Shared		Reclaimed	
		MB	MB	%	MB	%
A	10 WinNT	2048	880	42.9	673	32.9
B	9 Linux	1846	539	29.2	345	18.7
C	5 Linux	1658	165	10.0	120	7.2

Figure 5: **Real-World Page Sharing.** Sharing metrics from production deployments of ESX Server. (a) Ten Windows NT VMs serving users at a Fortune 50 company, running a variety of database (Oracle, SQL Server), web (IIS, Websphere), development (Java, VB), and other applications. (b) Nine Linux VMs serving a large user community for a nonprofit organization, executing a mix of web (Apache), mail (Major-domo, Postfix, POP/IMAP, MailArmor), and other servers. (c) Five Linux VMs providing web proxy (Squid), mail (Postfix, RAV), and remote access (ssh) services to VMware employees.

- sharing

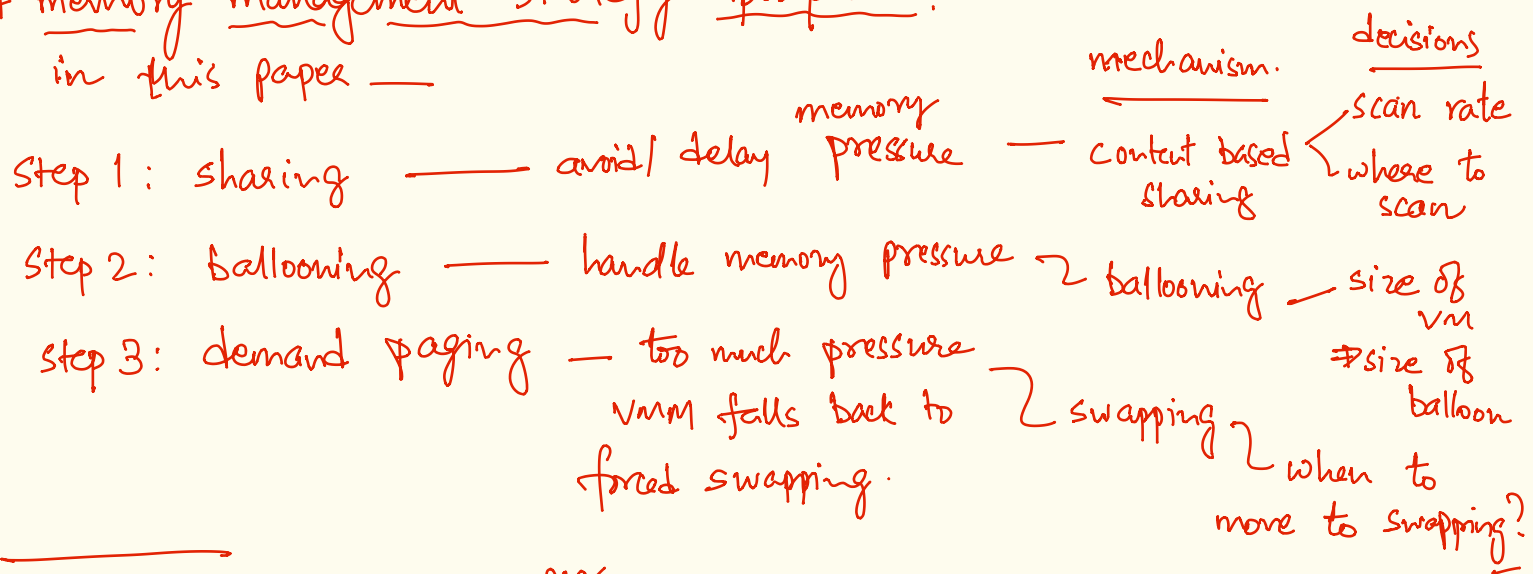
- ballooning

- demand paging.

WWS

working set size

memory management strategy proposed. in this paper —



Q what should the ^{memory} size of a VM? (by implication size of balloon).

- just enough to meet current memory usage
- demand on memory is dynamic/time-varying
- usage different from allocation (mapped memory).

Q One common strategy is share-based or proportionate sharing/allocation of memory.

eg: of VM₁, VM₂ & VM₃

a 1:2:3 ratio yields an allocation of 1GB, 2GB & 3GB resp. from a 6GB total memory.

paper refers to this as shares-per-VM.

and memory can be allocated to VMs that "pay" a higher price (shares/page) compared to a VM that pays a lower price.

BUT

allocation is not same as usage!
so, we really want effective shares-per-page.

i.e.: how allocated pages is the VM using.

$$f = \frac{S}{P \cdot (f + K \cdot (1-f))}$$

S ← shares of a VM.
 f ← effective / adjusted shares per page ratio
 P ← # allocated pages
 K ← idle page cost
 $f + K \cdot (1-f)$ ← fraction of active pages
 $P \cdot (f + K \cdot (1-f))$ ← referred as working set size in paper

$$K = \frac{1}{1-\tau}$$

① how to estimate 'f'?

- ① invalidate all p2m mappings.
 on each access.
 - page fault
 - add to active count
 - reset mapping.
 - redo access.

redo this loop periodically to get time varying ~~up~~ estimate of 'f'.

① invalidate ⇒ flip pte flag present bit to 0.
 do not remove mapping.

costly ~ each page access is a fault.

② redo 1 with a sample of 'n' pages.

- fraction of n approximates all pages
 n pages chosen at random.

eg: 100 pages every 30 second.

no cost ⇒ $f = \frac{S}{P}$

$\tau = 0$

no shares. $f = 0$
 ⇒ random allocation!

⇒ in practice

$0 \leq \tau < 1$

$\tau = 1$

⊛ estimate 'f' in every time window

- $n=100$ $T=30$ sec.

$f = \frac{a}{n}$ — # accessed pages out of n
in time t .

- moving average across time windows
to keep updating f.

from f estimate s ← effective shares
per page.
for each VM.

- use s_1, s_2, s_3 etc. for proportionate
sizing of memory of each VM.

- adjust memory sizes of each VM
via ballooning.

Questions answered via experiments —

- Does the sampling based heuristic track memory usage well/accurately?
- ~~Do~~ How does idle tax & active fraction change allocation levels of VM?