

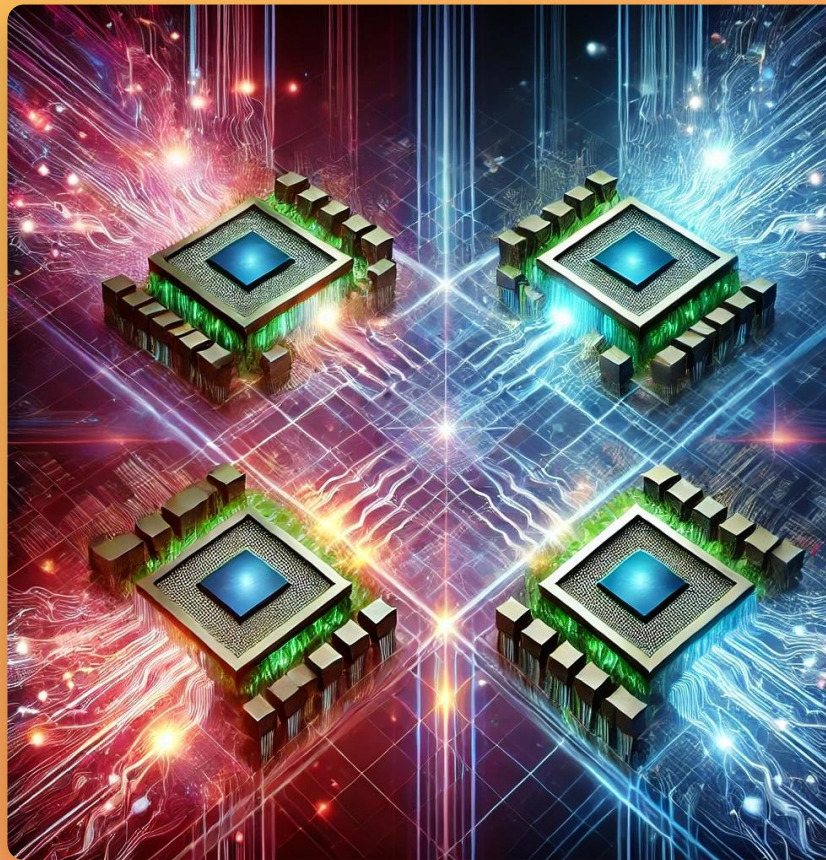
Never Underestimate Memory Architecture



Bryan Boreham



[@bboreham.bsky.social](https://bboreham.bsky.social)



Never Underestimate
Memory Architecture

又々

Presenter



Bryan Boreham

Distinguished Engineer
Grafana Labs



[@bboreham.bsky.social](https://bsky.app/profile/bboreham.bsky.social)



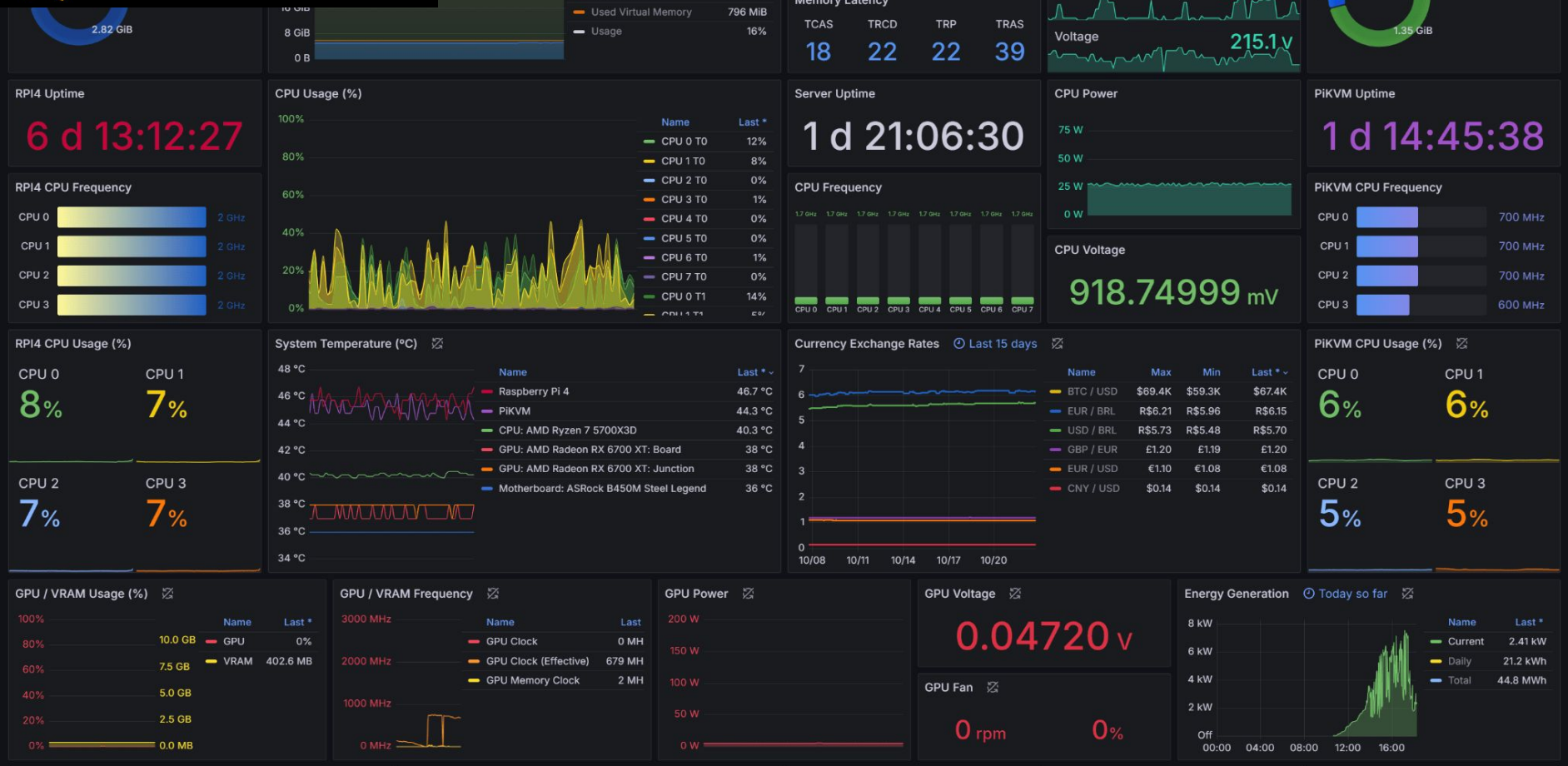
Outline

Why am I interested in NUMA?

What does NUMA mean?

What can you do about it?

More memory architecture?



Memory Clock

4091 MHz

Memory Latency

TCAS

TRCD

TRP

TRAS

18

22

22

39

System Power Metrics

Power

Current

Voltage

5.3 w

59.0 mA

215.1 v

PIKVM Memory Usage (%)

256 MiB

307 MiB

1.35 GiB

Available

Used

Reserved

71%

16%

13%

RPI4 Uptime

6 d 13:12:27

RPI4 CPU Frequency

CPU 0

CPU 1

CPU 2

CPU 3

2 GHz

2 GHz

2 GHz

2 GHz

CPU Usage (%)

CPU 0 T0

CPU 1 T0

CPU 2 T0

CPU 3 T0

CPU 4 T0

CPU 5 T0

CPU 6 T0

CPU 7 T0

CPU 0 T1

CPU 1 T1

12%

8%

0%

1%

0%

0%

0%

0%

14%

0%

Server Uptime

1 d 21:06:30

CPU Frequency

CPU 0

CPU 1

CPU 2

CPU 3

CPU 4

CPU 5

CPU 6

CPU 7

1.7 GHz

1.7 GHz

1.7 GHz

1.7 GHz

1.7 GHz

1.7 GHz

1.7 GHz

1.7 GHz

CPU Power

75 W

50 W

25 W

0 W

CPU Voltage

918.74999 mV

PIKVM Uptime

1 d 14:45:38

PIKVM CPU Frequency

CPU 0

CPU 1

CPU 2

CPU 3

700 MHz

700 MHz

700 MHz

600 MHz

RPI4 CPU Usage (%)

CPU 0

CPU 1

CPU 2

CPU 3

8%

7%

7%

7%

System Temperature (°C)

Raspberry Pi 4

PIKVM

CPU: AMD Ryzen 7 5700X3D

GPU: AMD Radeon RX 6700 XT: Board

GPU: AMD Radeon RX 6700 XT: Junction

Motherboard: ASRock B450M Steel Legend

46.7 °C

44.3 °C

40.3 °C

38 °C

38 °C

36 °C

Currency Exchange Rates

Last 15 days

BTC / USD

EUR / BRL

USD / BRL

GBP / EUR

EUR / USD

CNY / USD

\$69.4K

R\$6.21

R\$5.73

£1.20

€1.10

\$0.14

\$59.3K

R\$5.96

R\$5.70

£1.19

€1.08

\$0.14

\$67.4K

R\$6.15

R\$5.70

£1.20

€1.08

\$0.14

PIKVM CPU Usage (%)

CPU 0

CPU 1

CPU 2

CPU 3

6%

6%

5%

5%

GPU / VRAM Usage (%)

GPU

VRAM

10.0 GB

7.5 GB

5.0 GB

2.5 GB

0.0 MB

0%

GPU / VRAM Frequency

GPU Clock

GPU Clock (Effective)

GPU Memory Clock

3000 MHz

2000 MHz

1000 MHz

0 MHz

0 MH

679 MH

2 MH

GPU Power

200 W

150 W

100 W

50 W

0 W

GPU Voltage

0.04720 v

GPU Fan

0 rpm

0%

Energy Generation

Today so far

Current

Daily

Total

2.41 kW

21.2 kWh

44.8 MWh



Data visualization



Testing



Observability
solutions



Incident response
management



Key capabilities



Building blocks & telemetry databases (LGTM+)

Caution, hot material!

There is a lot of detail in this talk.

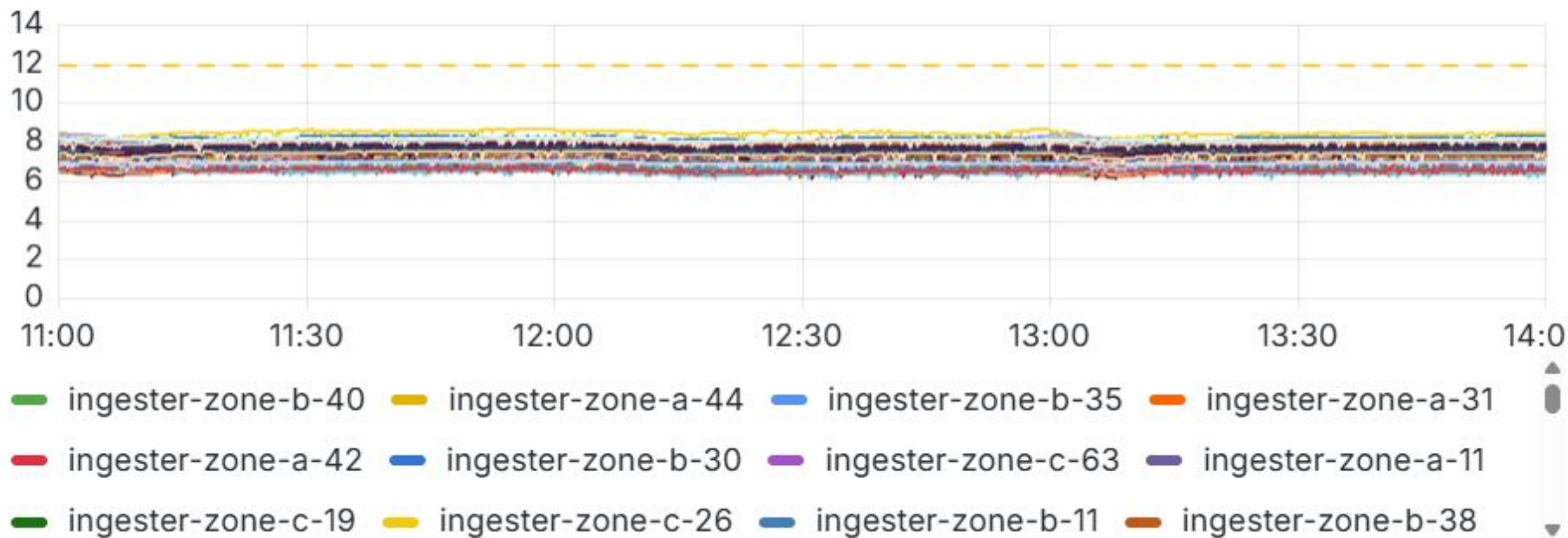
I settled on this sign, for the parts where you really have to pay attention.



OK, let's get on with the talk

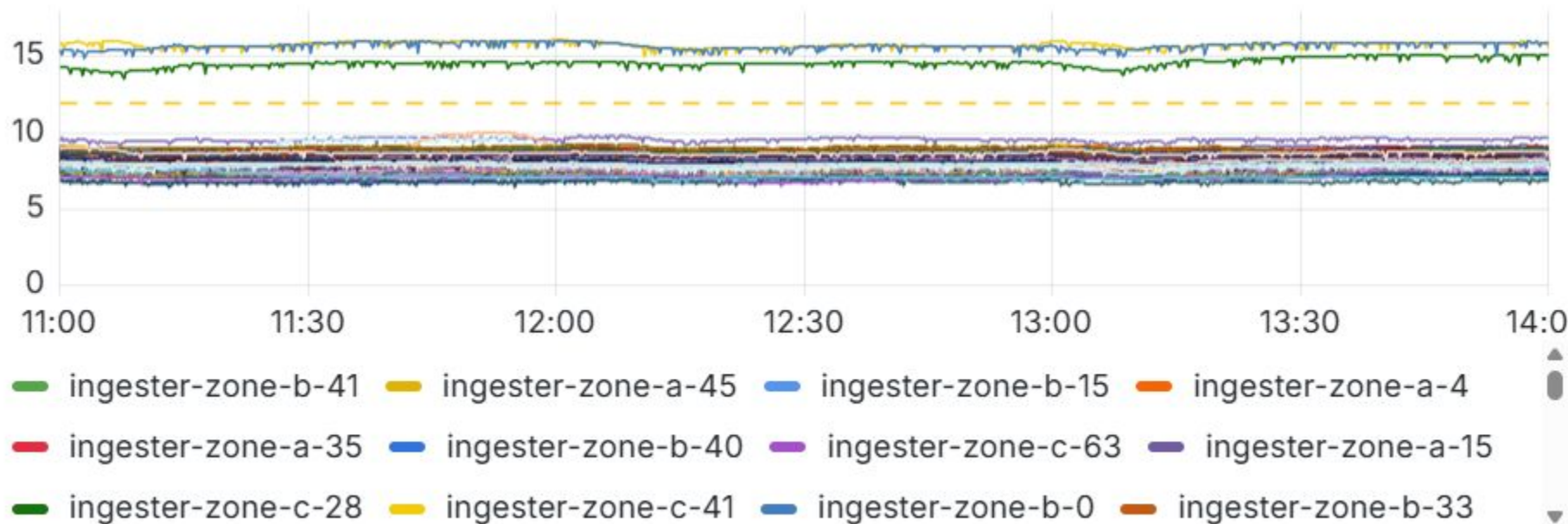
Story begins here

CPU



One day, some pods are using more CPU

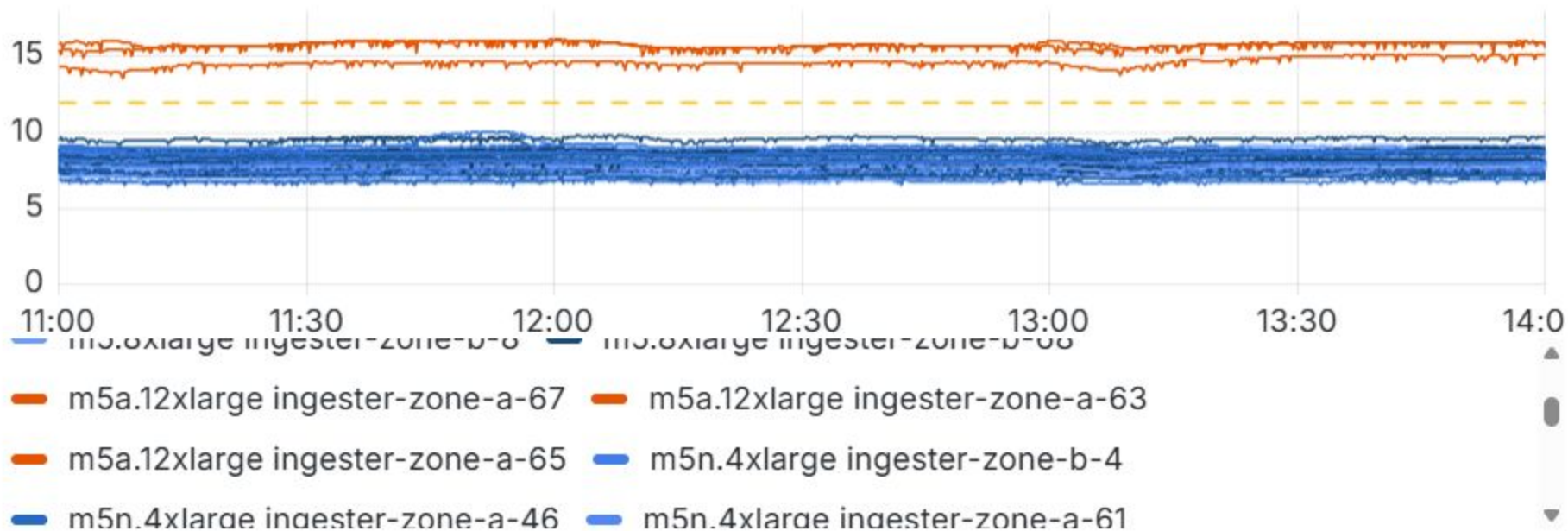
CPU



What's going on?

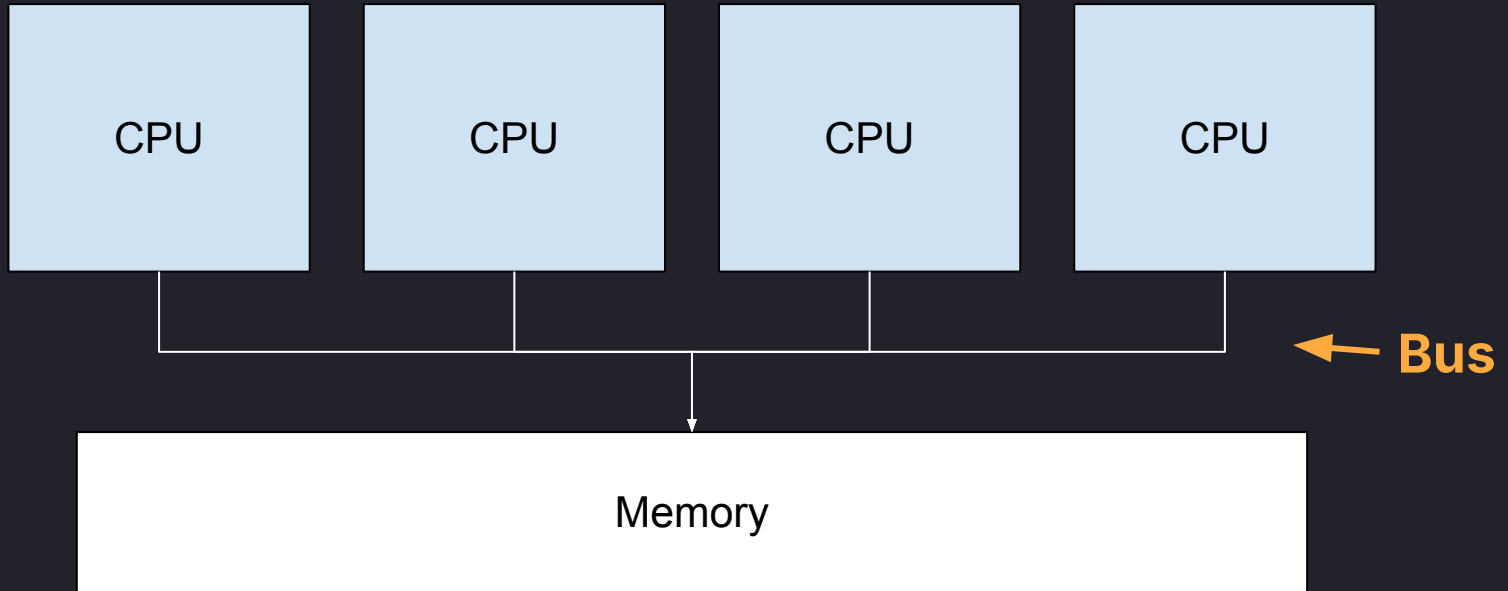
Showing instance type

CPU



How do computers work?

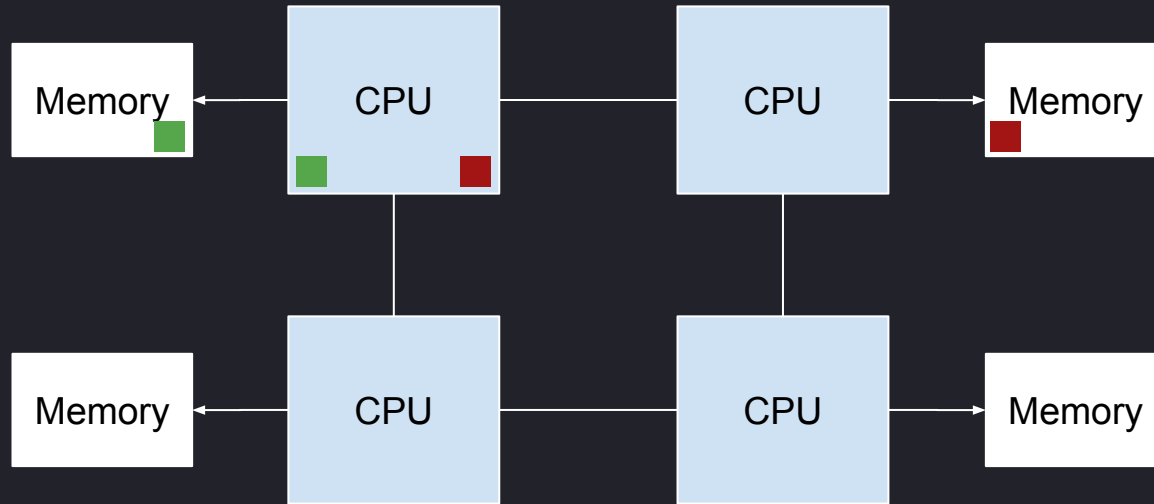
Conceptually



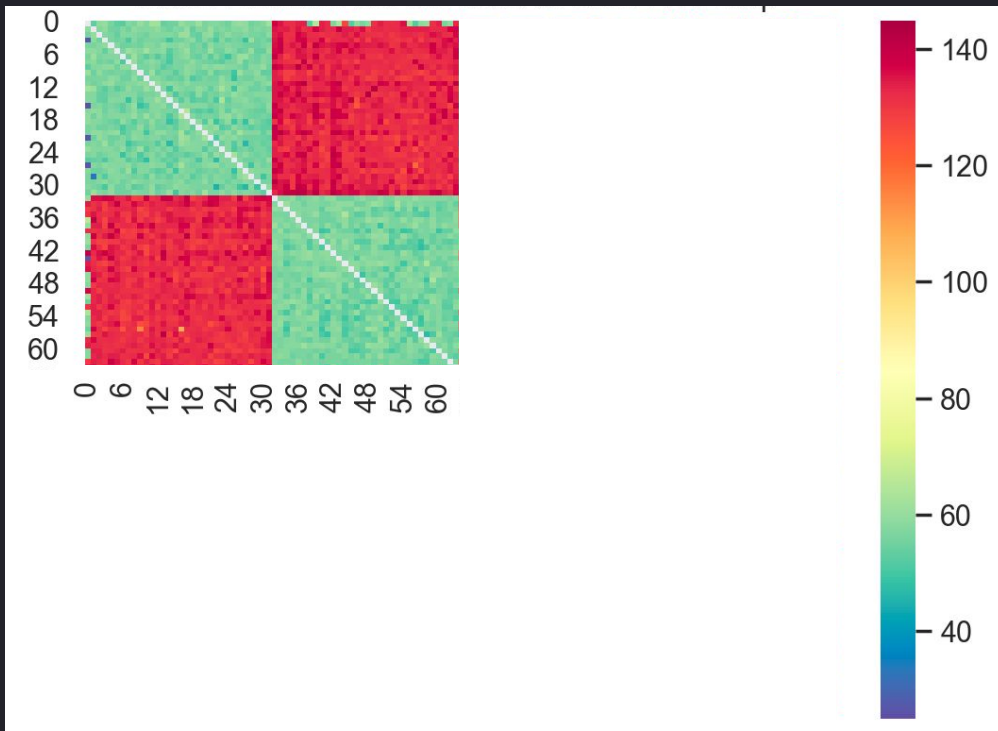
Reality



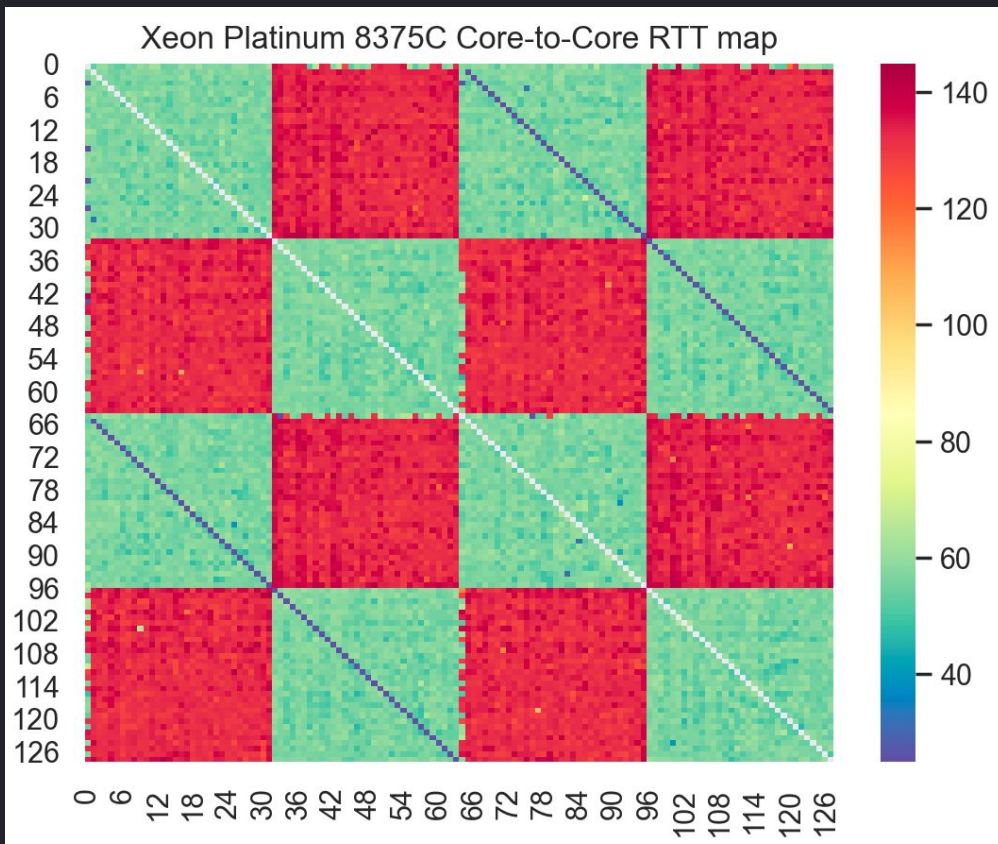
Non-Uniform Memory Access (NUMA)



Memory latency



Memory latency, with hyperthreading



Sidenote: workloads vary

How do I know if I'm suffering from NUMA?

Prometheus metrics

From node-exporter (with zoneinfo collector enabled):

`node_zoneinfo_numa_other_totalvs`

`node_zoneinfo_numa_local_total`



How do I find a server's NUMA characteristics?

On the CSP Website?

Instance Size	vCPU	Memory (GiB)	Instance Storage (GB)	Network Bandwidth (Gbps)	EBS Bandwidth (Gbps)
m7i.large	2	8	EBS-Only	Up to 12.5	Up to 10
m7i.xlarge	4	16	EBS-Only	Up to 12.5	Up to 10
m7i.2xlarge	8	32	EBS-Only	Up to 12.5	Up to 10
m7i.4xlarge	16	64	EBS-Only	Up to 12.5	Up to 10
m7i.8xlarge	32	128	EBS-Only	12.5	10
m7i.12xlarge	48	192	EBS-Only	18.75	15



lscpu command - m5a.12xlarge

```
Architecture:          x86_64
  Model name:          AMD EPYC 7571
    Thread(s) per core: 2
    Core(s) per socket: 24
    Socket(s):         1

NUMA:
  NUMA node(s):        3
  NUMA node0 CPU(s):   0-7,24-31
  NUMA node1 CPU(s):   8-15,32-39
  NUMA node2 CPU(s):   16-23,40-47
```



lscpu command - m7g.16xlarge

```
Architecture:          aarch64
  Model name:          Neoverse-V1
    Thread(s) per core: 1
    Core(s) per cluster: 64
    Cluster(s):         1
NUMA:
  NUMA node(s):        1
  NUMA node0 CPU(s):   0-63
```

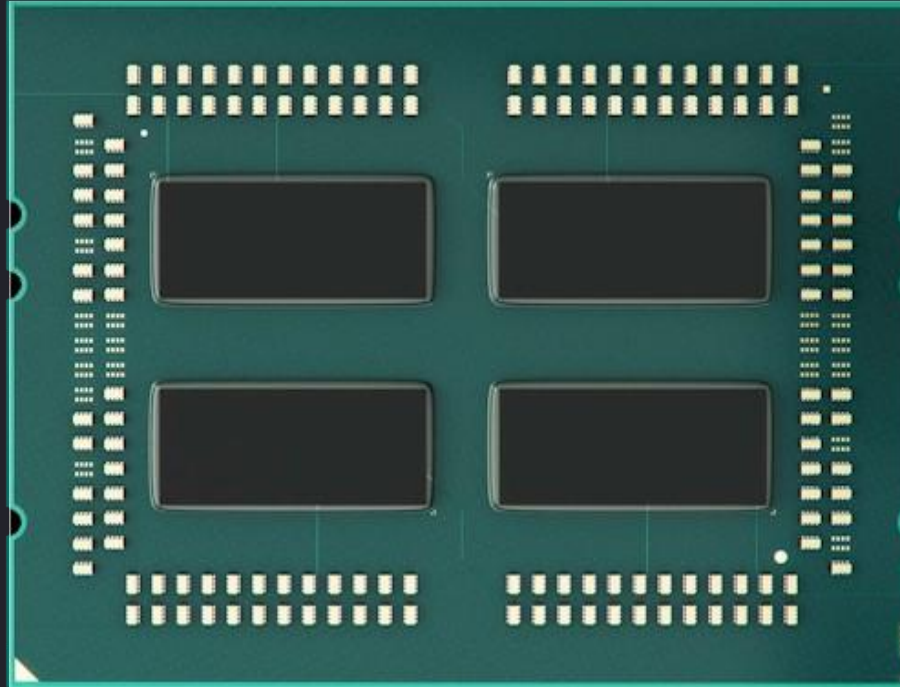


lscpu command - m8g.metal-48xl

```
Architecture:          aarch64
  Model name:           Neoverse-V2
    Thread(s) per core: 1
    Core(s) per cluster: 96
    Cluster(s):         2
NUMA:
  NUMA node(s):         2
  NUMA node0 CPU(s):    0-95
  NUMA node1 CPU(s):    96-191
```



AMD Chiplets



AMD EPYC 7xx1 32-core (64 with SMT)

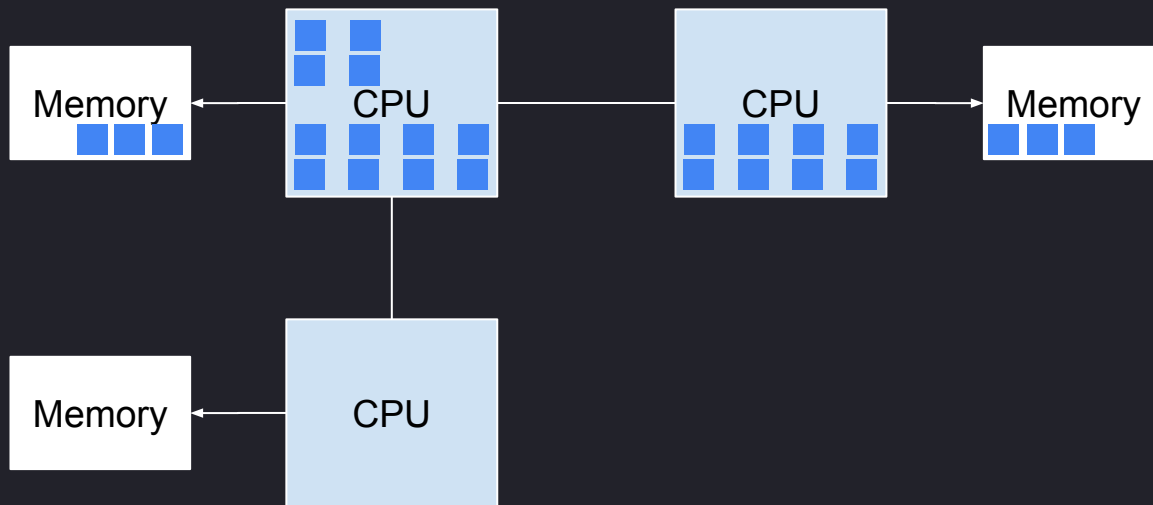


Memory latency, m5a.12xlarge model

C → C	CPU0	CPU1	CPU2	CPU3	CPU4	CPU5	CPU6	CPU7	CPU8	CPU9	CPU10	CPU11	CPU12	CPU13	CPU14	CPU15	CPU16	CPU17	CPU18	CPU19	CPU20	CPU21	CPU22	CPU23
CPU0	X	32	33	34	182	182	183	182	220	220	220	220	220	220	220	219	222	222	223	223	222	223	223	223
CPU1	32	X	32	33	182	181	183	182	219	218	219	219	219	218	219	219	221	225	222	225	221	224	222	225
CPU2	33	32	X	33	183	183	183	183	220	220	221	220	220	220	221	220	223	223	223	223	223	223	223	223
CPU3	34	33	33	X	182	182	183	182	219	219	220	219	219	219	220	219	222	222	223	223	222	222	222	223
CPU4	182	182	183	182	X	32	33	34	220	219	220	220	220	220	220	220	222	222	223	222	222	222	223	223
CPU5	182	181	183	182	32	X	32	33	218	218	219	219	219	218	219	219	221	224	222	225	221	224	222	225
CPU6	183	183	183	183	33	32	X	33	221	220	221	220	220	220	221	220	223	223	223	223	223	223	224	223
CPU7	182	182	183	182	34	33	33	X	220	219	220	220	220	219	220	220	222	222	222	222	222	222	222	223
CPU8	220	219	220	219	220	218	221	220	X	32	33	34	182	182	183	182	222	224	222	225	222	224	223	224
CPU9	220	218	220	219	219	218	220	219	32	X	32	33	182	181	183	182	221	221	222	221	221	221	222	222
CPU10	220	219	221	220	220	219	221	220	33	32	X	33	183	182	183	183	222	226	223	227	223	226	223	227
CPU11	220	219	220	219	220	219	220	220	34	33	33	X	182	182	183	182	222	224	222	224	222	224	222	224
CPU12	220	219	220	219	220	219	220	220	182	182	183	182	X	32	33	34	222	224	222	224	222	224	223	224
CPU13	220	218	220	219	220	218	220	219	182	181	182	182	32	X	32	33	221	221	222	222	221	221	222	222
CPU14	220	219	221	220	220	219	221	220	183	183	183	183	33	32	X	33	222	226	223	226	223	226	223	227
CPU15	219	219	220	219	220	219	220	220	182	182	183	182	34	33	33	X	222	224	222	225	222	224	222	224
CPU16	222	221	223	222	222	221	223	222	222	221	222	222	222	221	222	222	X	32	33	34	192	193	192	194
CPU17	222	225	223	222	222	224	223	222	224	221	226	224	224	221	226	224	32	X	32	33	193	191	196	193
CPU18	223	222	223	223	223	222	223	222	222	222	223	222	222	222	223	222	33	32	X	33	193	196	194	197
CPU19	223	225	223	223	222	225	223	222	225	221	227	224	224	222	226	225	34	33	33	X	193	193	196	195
CPU20	222	221	223	222	222	221	223	222	222	221	223	222	222	221	223	222	192	193	193	193	X	32	33	34
CPU21	223	224	223	222	222	224	223	222	224	221	226	224	224	221	226	224	193	191	196	193	32	X	32	33
CPU22	223	222	223	222	223	222	224	222	223	222	223	222	223	222	223	222	192	196	194	196	33	32	X	33
CPU23	223	225	223	223	223	225	223	223	224	222	227	224	224	222	227	224	194	193	197	195	34	33	33	X

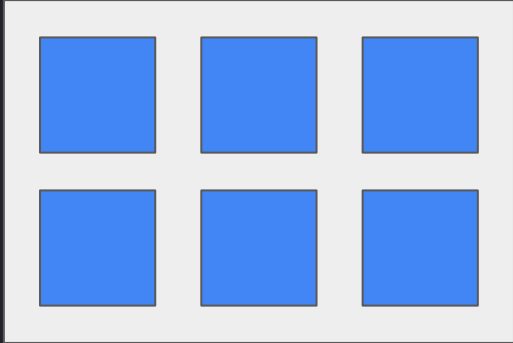


Situation in the incident (probably)

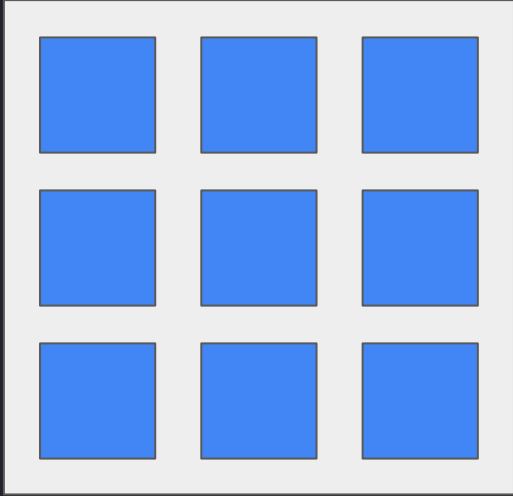


How did we resolve the incident?

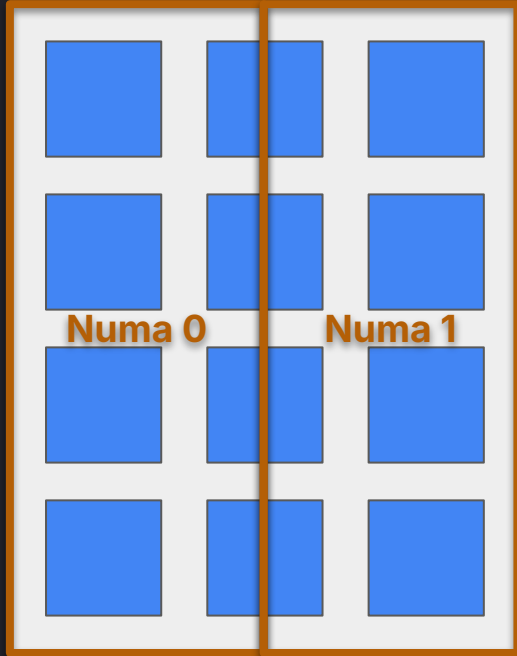
Say you start with small instances



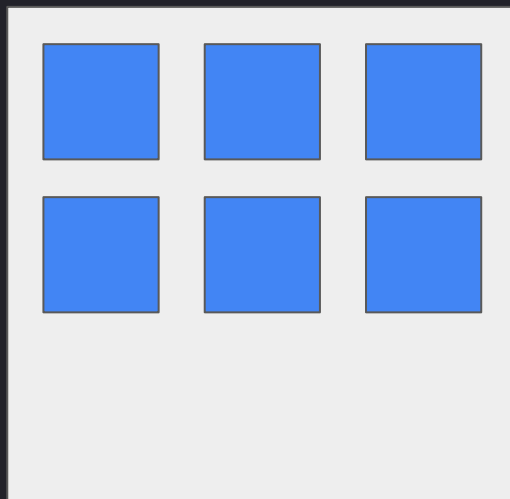
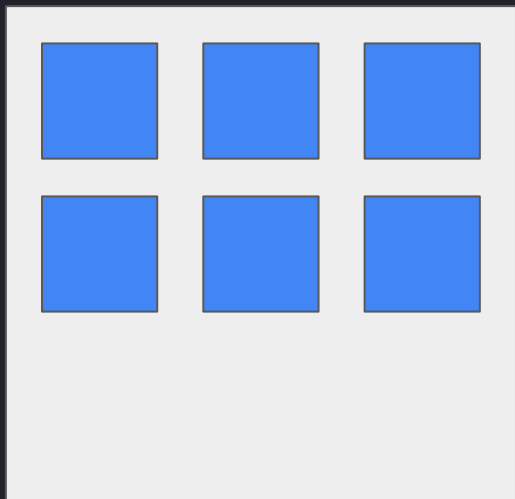
Then you move to bigger instances



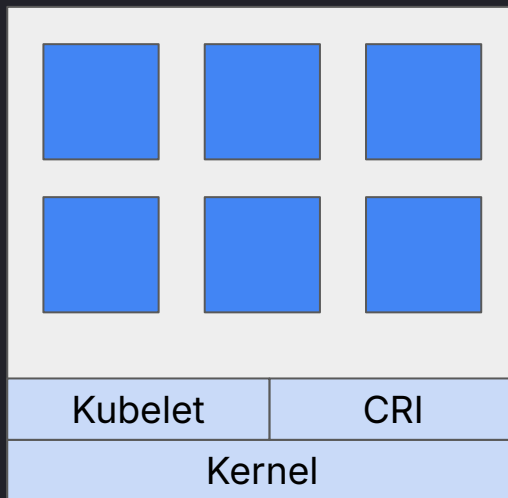
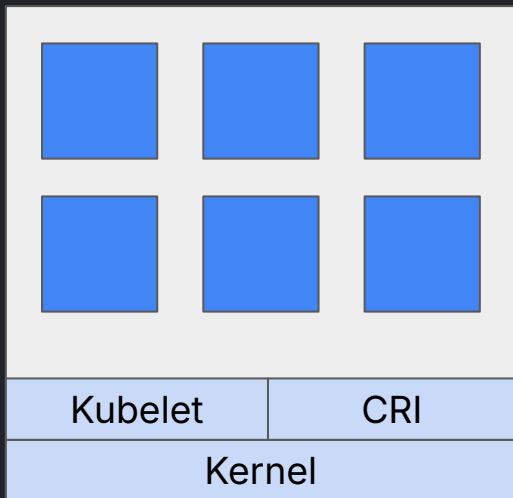
And bigger



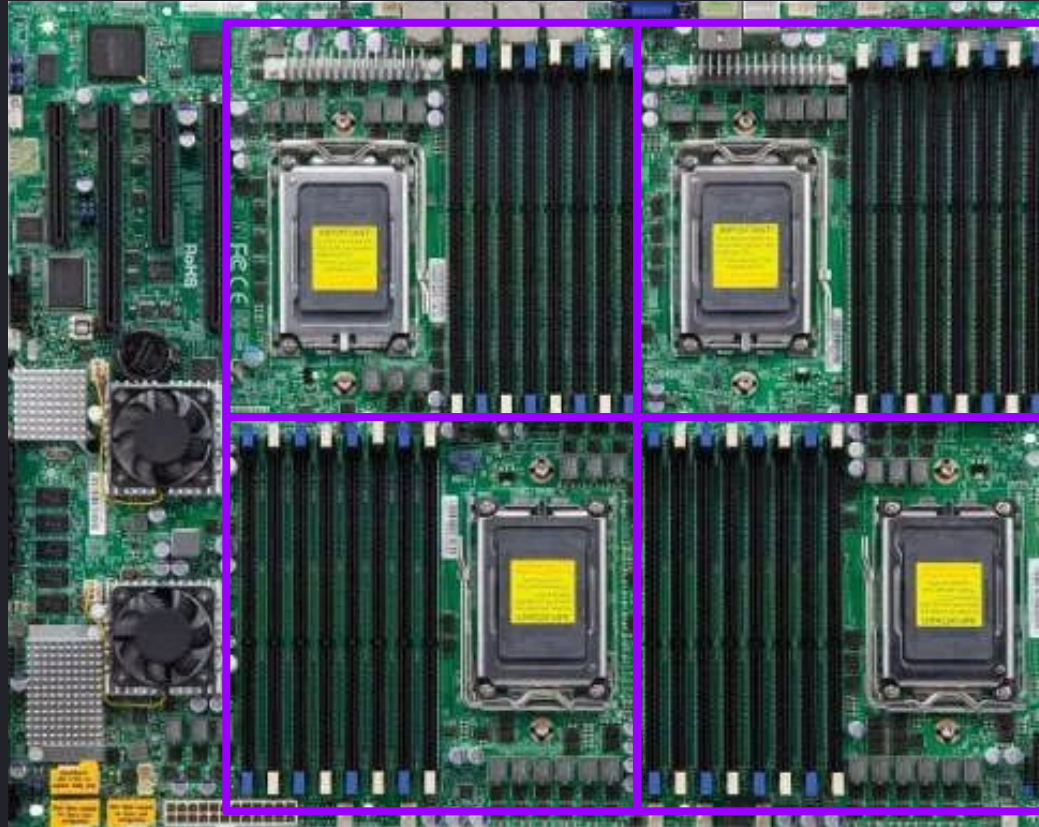
You might be better off with multiple smaller instances



But remember there are per-instance overheads



Physical view





Kubernetes

CPU Manager



`--cpu-manager-policy` kubelet flag

- **none**: the default policy.
- **static**: for containers in Guaranteed pods with integer CPU requests

"CPUManager aligns CPU allocations at the NUMA boundary"



<https://kubernetes.io/docs/concepts/policy/node-resource-managers/>

Topology Manager

scope=pod - to locate a set of containers together on the same NUMA nodes.

policy=single-numa-node - will reject a pod if it can't all run on one node.

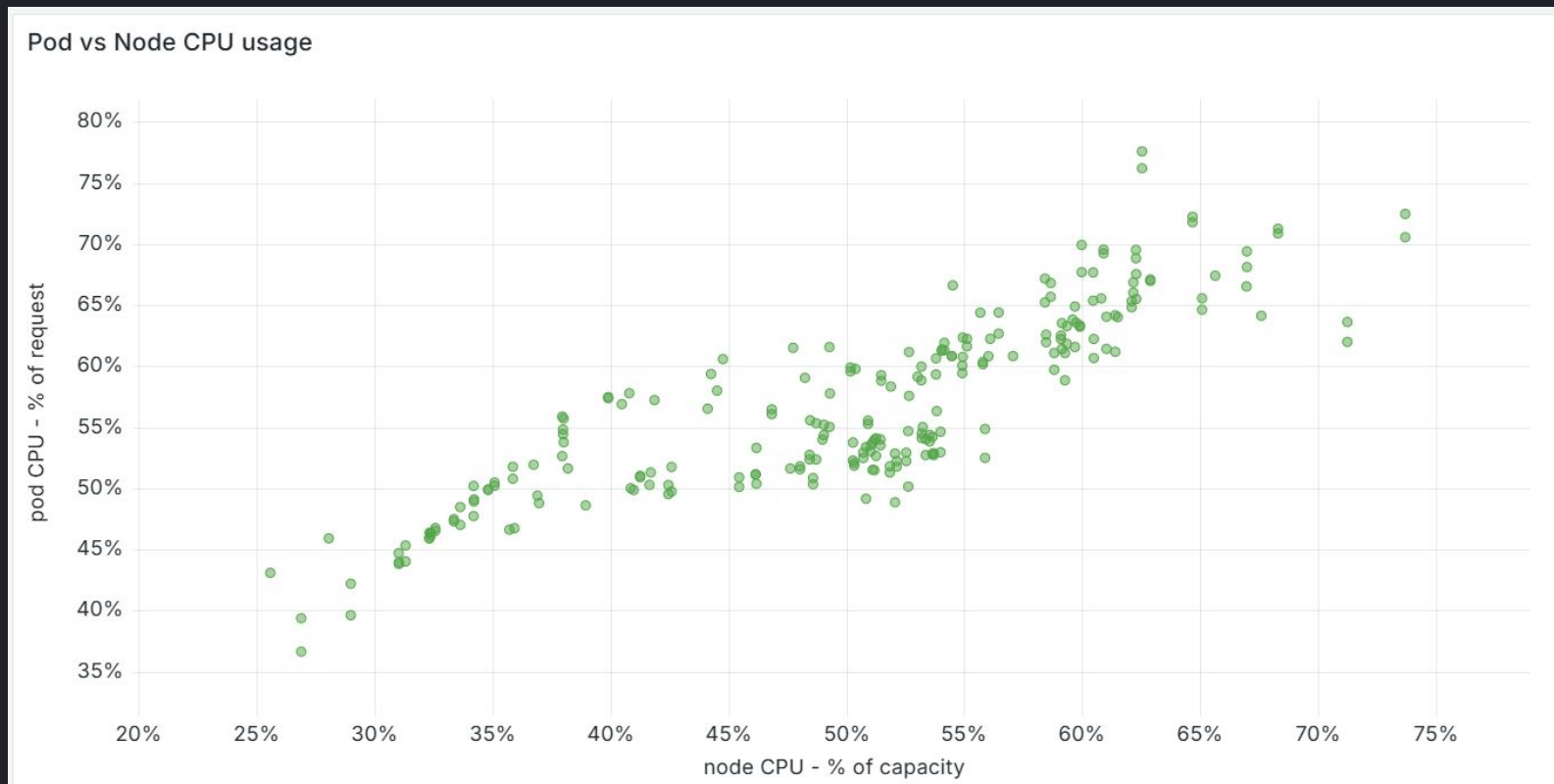


<https://kubernetes.io/docs/tasks/administer-cluster/topology-manager/>

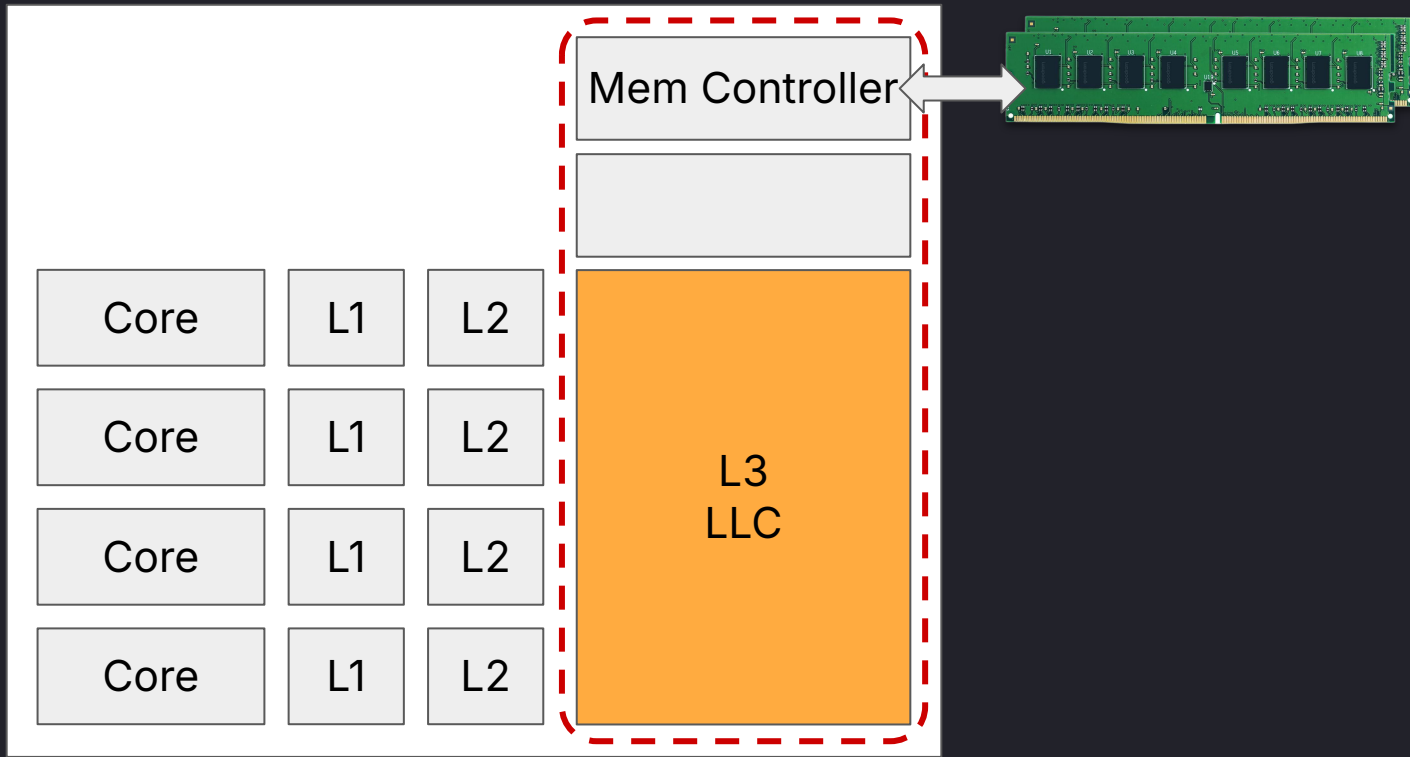


"Uncore"

Pods vs instances CPU plot



"Uncore" = the parts outside of cores



CPU Manager



`--cpu-manager-policy=static` kubelet flag

Option: `prefer-align-cpus-by-uncorecache`

New in Kubernetes version 1.32 (alpha, hidden by default).



<https://kubernetes.io/docs/tasks/administer-cluster/cpu-management-policies/>

Wrapping up

Benefit from a little understanding of underlying hardware.

Stay within NUMA zones if you possibly can.

Everything is non-uniform.



Thank you

Special thanks to:

JuanJo Ciarlante, Grafana Labs

Anton Kolesnikov, Grafana Labs

Jonathan Perry, Unvariance

Milind Chabbi, Uber

Links

<https://github.com/SoilRos/cpu-latency>

<https://www.intel.com/.../articles/technical/intel-vtune-amplifier-functionality-on-aws-instances>

<https://jprahman.substack.com/p/amd-cpu-topology-rome-milan>

<https://www.nextplatform.com/2019/08/15/a-deep-dive-into-amds-rome-epyc-architecture/>

<https://anarazel.de/talks/2024-10-23-pgconf-eu-numa-vs-postgresql/numa-vs-postgresql.pdf>

<https://frankdenneman.nl/2019/10/14/amd-epyc-naples-vs-rome-and-vsphere-cpu-scheduler-updates/>



