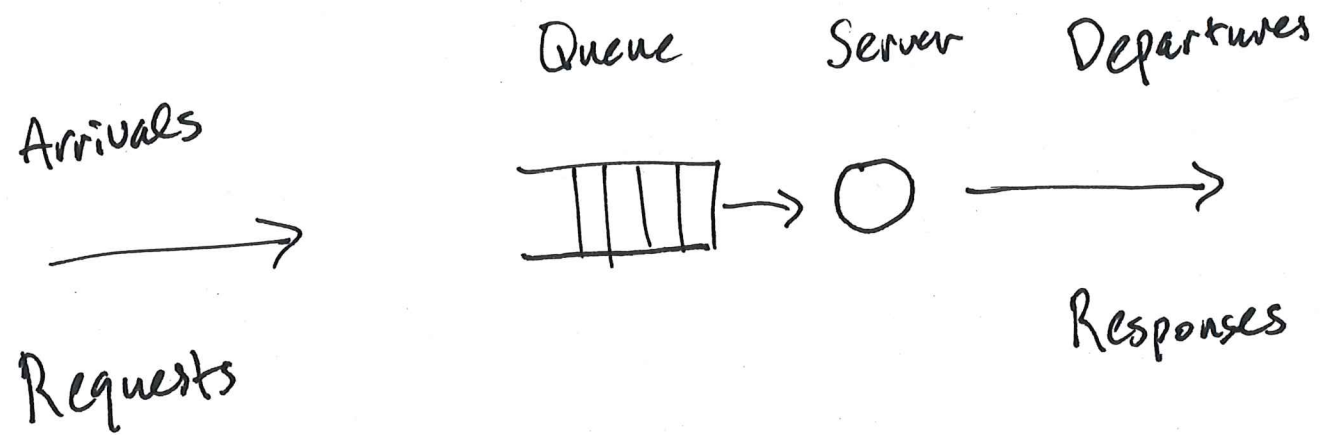


Lec 12

Queueing Theory



Web Server example

- processes 1 request at a time
- takes 1 ms to process a request
- if requests arrive exactly every 10ms,
what is the throughput? latency?
- every 2 ms?
- 1 ms?
- one extra request arrives...
- every 0.99 ms?

Definitions

$$\begin{array}{lcl} \text{response time} & = & \text{waiting time} + \text{service time} \\ R & & W \quad S \end{array}$$

arrival rate λ (avg, could be drawn from distribution)

service rate $\mu = 1/S$ how many requests can be served per unit time

$$\text{throughput } X = \min(\lambda, \mu)$$

number of tasks in system N

Little's Law

In a stable system ($\lambda < \mu$):

$$N = X R$$

$$\begin{array}{ccccc} \text{avg} & & \text{(avg)} & & \text{avg} \\ \text{number of} & = & \text{throughput} & \cdot & \text{response time} \\ \text{tasks} & & & & \\ \text{in system} & & & & \\ & & \frac{\text{req}}{\text{s}} & & \frac{\text{s}}{\text{req}} \end{array}$$

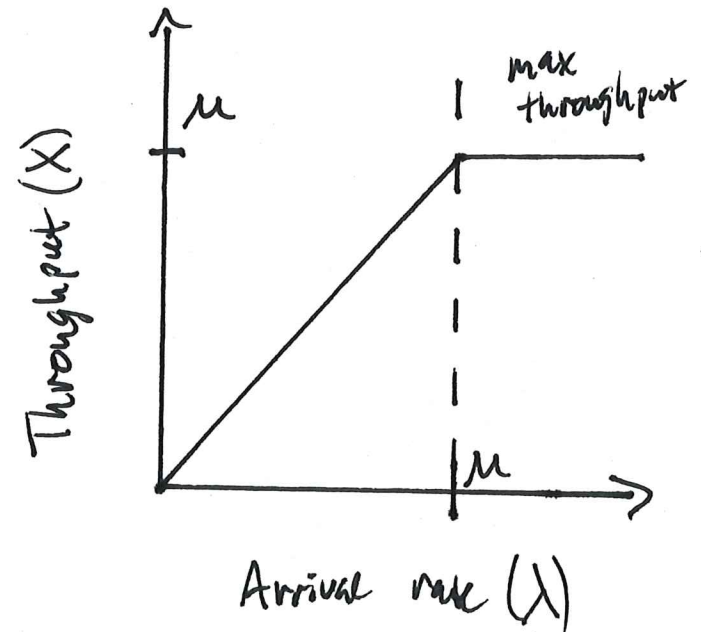
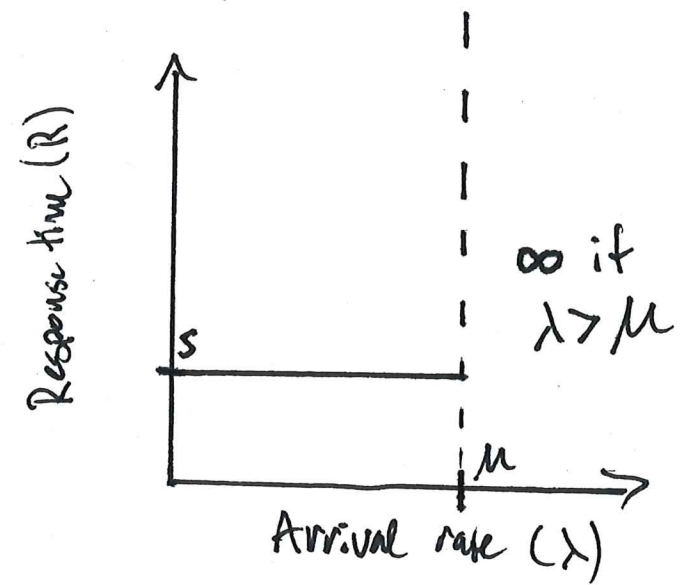
Example: if $X = 100 \text{ req/s}$ $R = \frac{1}{1000} \text{ s/req}$

then $N = 1/10 \text{ tasks in system}$

→ queue is usually empty

Evenly spaced arrivals (best case)

- tasks arrive evenly spaced with rate λ
- if $\lambda < \mu$, no queuing, $W = 0$
 $R = W + S = 0 + S = S$
 $X = \lambda$
- if $\lambda = \mu$, queue size does not change
empty initially $\rightarrow$ always empty
initially 10 in queue $\rightarrow$ always 10
- if $\lambda > \mu$, unbounded queuing
 $R = \infty$
 $X = \mu$



Bursty Arrivals (worst case)

- tasks arrive in groups
say 1 group per second
group size = K ($= \lambda$)

R = avg response time

$$= (S + 2S + \dots + KS) / K$$

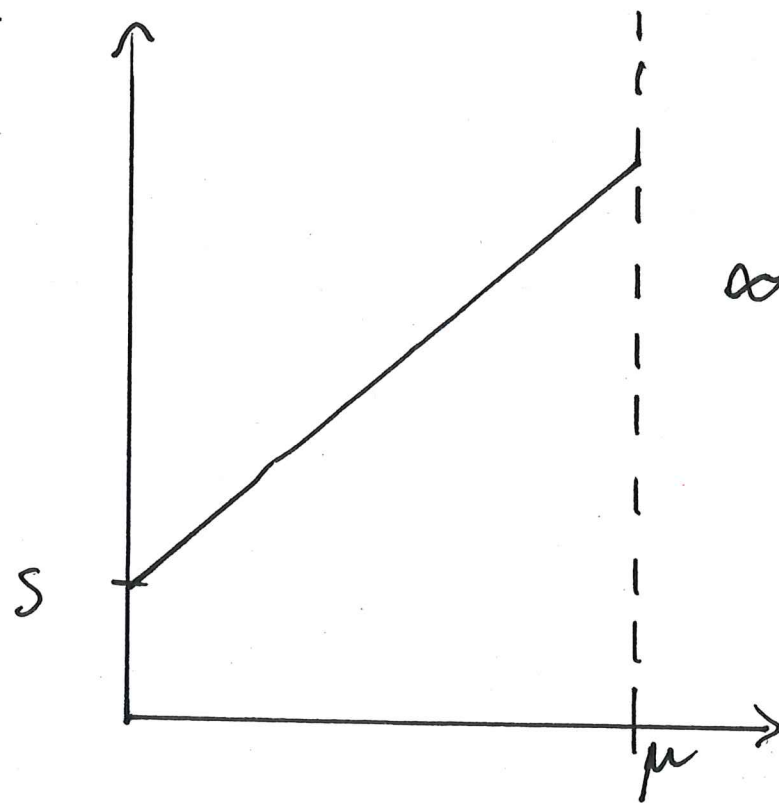
$$= S(1 + 2 + \dots + K) / K$$

$$= \frac{SK(K+1)}{2K}$$

$$= \frac{S}{2} \cdot (K+1) = S \cdot \frac{K+1}{2}$$

$$= O(K) \cdot S$$

Response time (R)



Arrival rate (λ)

Worse by a factor of
half the group size

groups of 10
 $\hookrightarrow$ 5x worse!

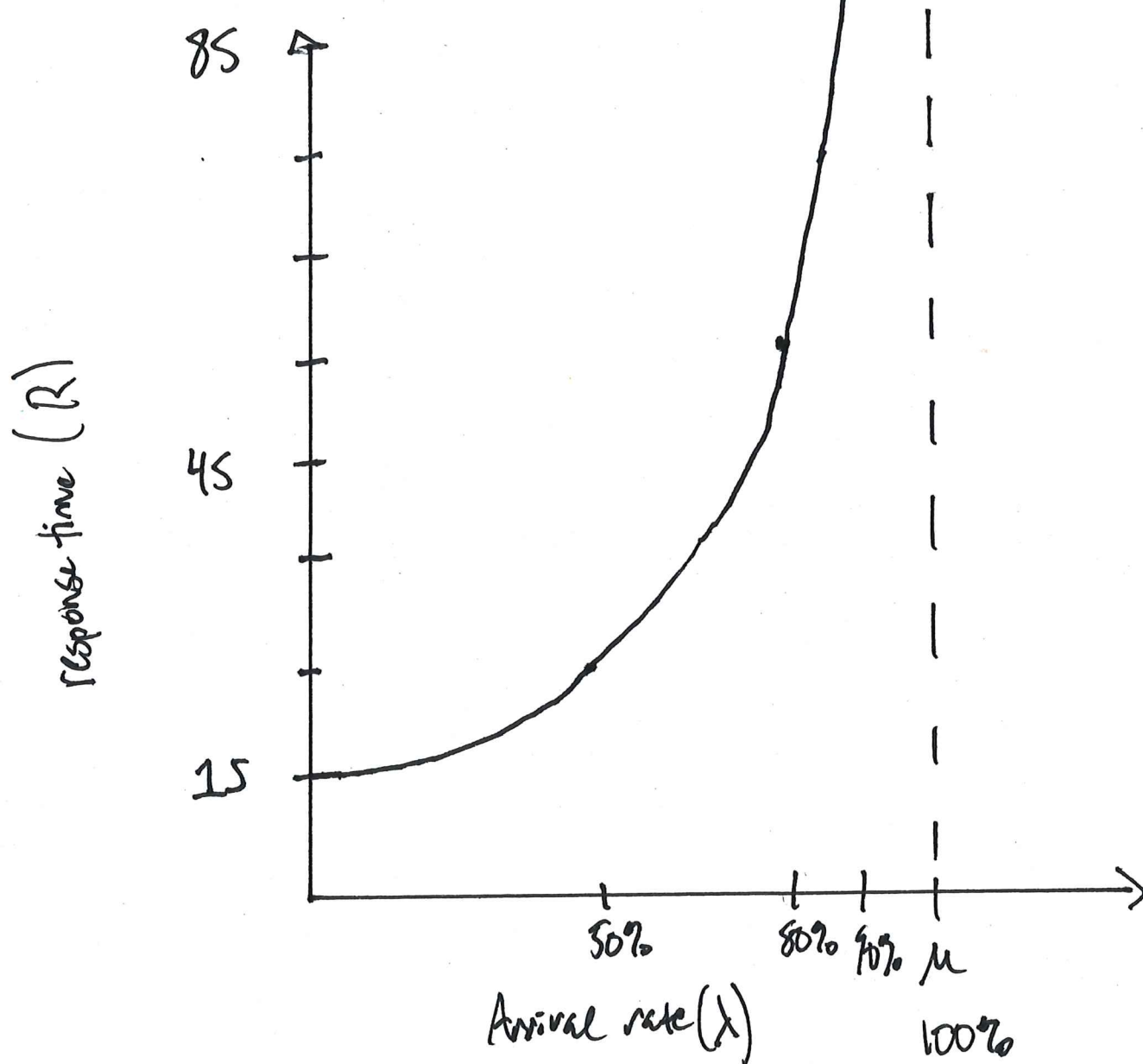
12.6

Independent Arrivals

- evenly spaced $\rightarrow$ perfectly anti correlated
- grouped $\rightarrow$ perfectly correlated
- reality: different users send requests independently
 $\hookrightarrow$ somewhere between worst + best case
- good model: exponentially distributed interarrival times
(Poisson process) $\rightarrow$ memoryless (Markov)

$$\hookrightarrow R = \frac{S}{1-u} \quad \text{where} \quad \text{utilization } u = \frac{\lambda}{\mu} (\leq 1)$$

(exponentially distributed
interarrival times)



$$R = \frac{S}{1 - \rho}$$
$$= \frac{S}{1 - \lambda/\mu}$$

12.8

Overload Management

- if $\lambda > \mu$, response time unbounded: "overload"
- how to get out of overload?
 - reject some requests
 - make requests cheaper somehow (dyn $\rightarrow$ static)
- very common problem: systems try to do more work per request when overloaded
 - e.g. the more concurrent requests, the more lock contention
 - e.g. drop a request $\rightarrow$ sender retransmits

Why do buses cluster?

- Suppose bus schedule has one bus every 10 minutes
- they leave the first stop on time (10 mins apart)
- why do they tend to get to last stop at around the same time? (instead of 10 mins apart)

12.10