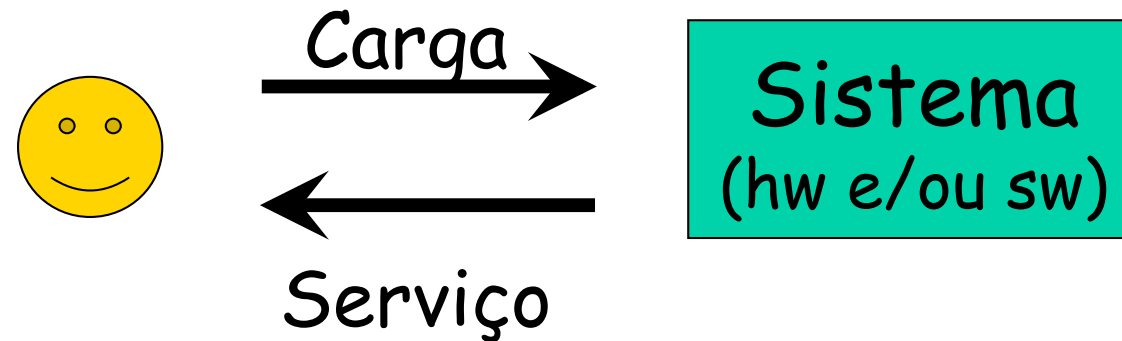


Análise e Modelagem de Desempenho de Sistemas de Computação

Profa. Jussara M. Almeida
1º Semestre de 2014

Modelo de Sistema



- Modelo: representação do comportamento do desempenho do sistema
- Etapas principais:
 - Caracterização do comportamento do usuário / carga: taxa de chegada, tempo de serviço,....
 - Medições no sistema: utilização de recursos, nível de multiprogramação
 - Medições no serviço: tempo de resposta, taxa de sucesso

Exemplos de Aplicação

- **Modelagem analítica** responde:
 - Como o tempo de resposta de um banco de dados de transações varia com a taxa de transações?
 - Qual o impacto no tempo de resposta de um upgrade de CPU? De disco?
 - Em média, qual o número de processos que ficam bloqueados no semáforo X da aplicação Y?
- **Simulação** responde:
 - Qual a política de replicação de conteúdo que resulta em maior byte hit ratio?
- **Experimentação** responde:
 - Quais os principais componentes do tempo de resposta em um servidor Web?
 - Qual o impacto do mecanismo de slow-start do protocolo TCP no desempenho de servidores Web tradicionais?

Seleção da Técnica de Avaliação de Desempenho

- Modelagem analítica
 - Pode ser bastante precisa e simples
 - Baixo custo: fornece resultados rápidos
 - Facilita projeto e configuração do sistema: melhora conhecimento sobre ele
 - avaliação dos compromissos entre vários parâmetros
 - impacto de cada parâmetro
 - Responde perguntas do tipo *what if*
 - Captura aspectos mais gerais do funcionamento do sistema
 - não captura alguns aspectos mais específicos

Seleção da Técnica de Avaliação de Desempenho

- Simulação

- Custo mais elevado: simulação deve cobrir estado estacionário (as vezes difícil), várias execuções (quantas?)
- Captura detalhes do funcionamento do sistema
- Responde perguntas do tipo *what if*

- Experimentação

- Alta complexidade, muitas variáveis: alto custo
- Difícil avaliar impacto de fatores isolados: falta de controle, ruído de fatores externos
- Alta precisão *se e somente se* experimentação realizada corretamente (difícil)

Seleção da Técnica de Avaliação de Desempenho

Critério	Modelagem Analítica	Simulação	Experimentação
Estágio	Qualquer	Qualquer	Após protótipo
Tempo necessário	Pequeno/ Médio	Médio	Variável
Ferramenta	Análise	Linguagem de Programação	Instrumentação e Monitoração
Precisão	Variável		
Avaliação de Compromissos	Fácil	Médio	Difícil
Custo	Pequeno	Médio	Grande

Abordagem Sistemática para Avaliação de Desempenho

- Definir objetivos e escopo (sistema)
- Listar serviços e saídas
- Selecionar métricas de desempenho
- Listar parâmetros
- Selecionar *fatores* para estudo
- Selecionar técnica de avaliação
- Selecionar carga de trabalho
- Projetar experimentos
- Analisar e interpretar dados (resultados)
- Apresentar resultados

Objetivos e Escopo

- Definir objetivos do estudo é essencial para definir escopo
- Definir escopo é chave para as demais escolhas de métricas, cargas, técnica de avaliação
- Exemplos: Dadas 2 CPUs
 - Objetivo 1: estimar impacto no tempo de resposta de usuários interativos
 - Escopo: sistema de time sharing, resultado depende de outros fatores externos a CPU
 - Objetivo 2: CPU são similares com exceção da ALU
 - Escopo: somente componentes internos da CPU

Métricas de Desempenho

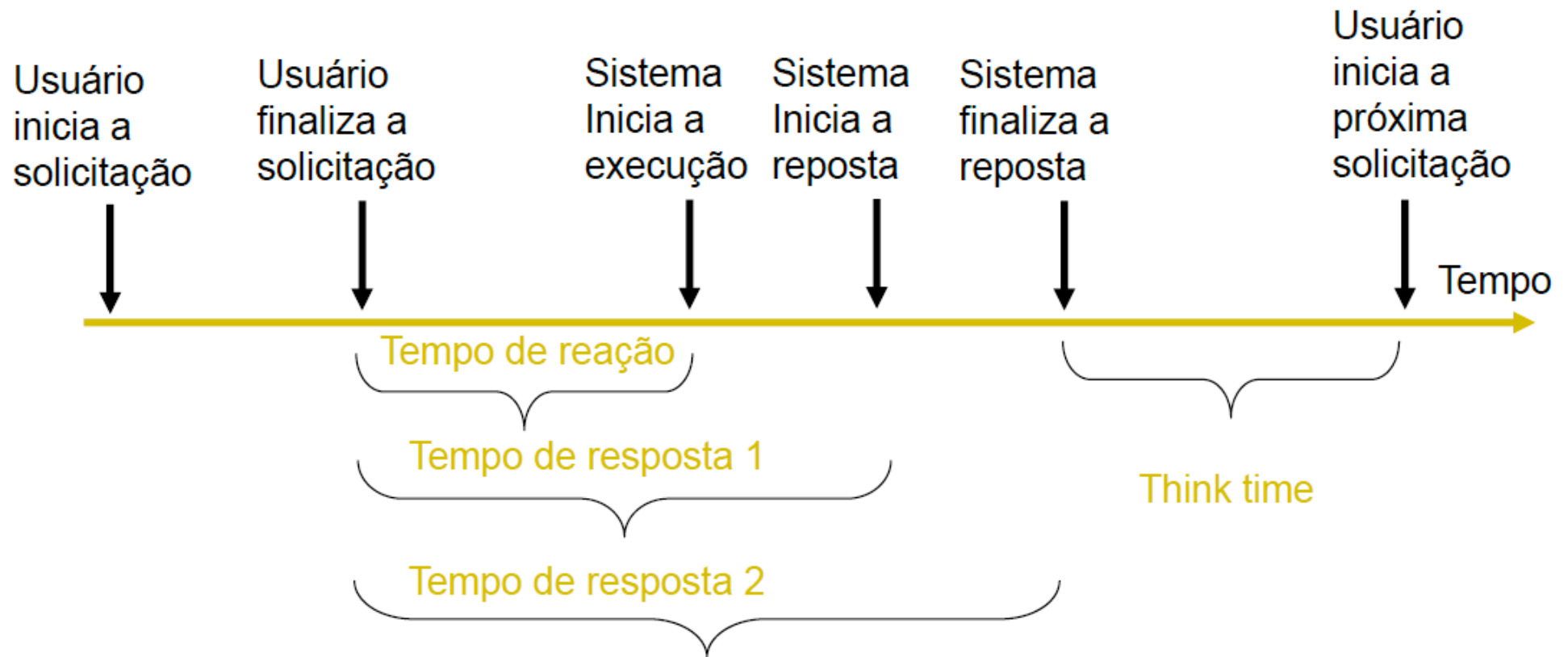
- Escolha específica para estudo, a partir da lista de serviços e possíveis saídas
 - Execução correta: desempenho, escalabilidade
 - tempo de resposta, taxa de processamento (serviço), utilização de recursos
 - Execução incorreta: confiabilidade
 - identificação das classes de erros
 - probabilidade de cada tipo de erro, tempo entre erros
 - Não execução: disponibilidade
 - Identificação das possíveis causas
 - Uptime (% tempo disponível), prob. de downtime, tempo entre falhas (MTTF = *Mean Time To Failure*)

Tempo de Resposta

- Intervalo de tempo entre requisição do usuário e a resposta do sistema
- Definição do intervalo tem que ser clara:
 - Inclui tempo entre momento que usuário termina comando e sistema inicia execução?
 - Inclui tempo entre início e término da geração da resposta?
- Pode conter vários componentes, com influência de vários subsistemas e da carga durante execução
- Ex: `time <programa> (Unix)`
3.5 real 0.2 user 0.9 sys

$\text{real} - (\text{user} + \text{sys}) = 2.4 \text{ segundos gastos ONDE???$

Tempo de Resposta



Slowdown

- Causado por:
 - operações de I/O (leituras, escritas, paginação)
 - tempo de rede
 - tempo gasto em outros programas (escalonamento)
 - contenção por recursos: filas dos recursos

Tempo de resposta = tempo de serviço + tempo nas filas

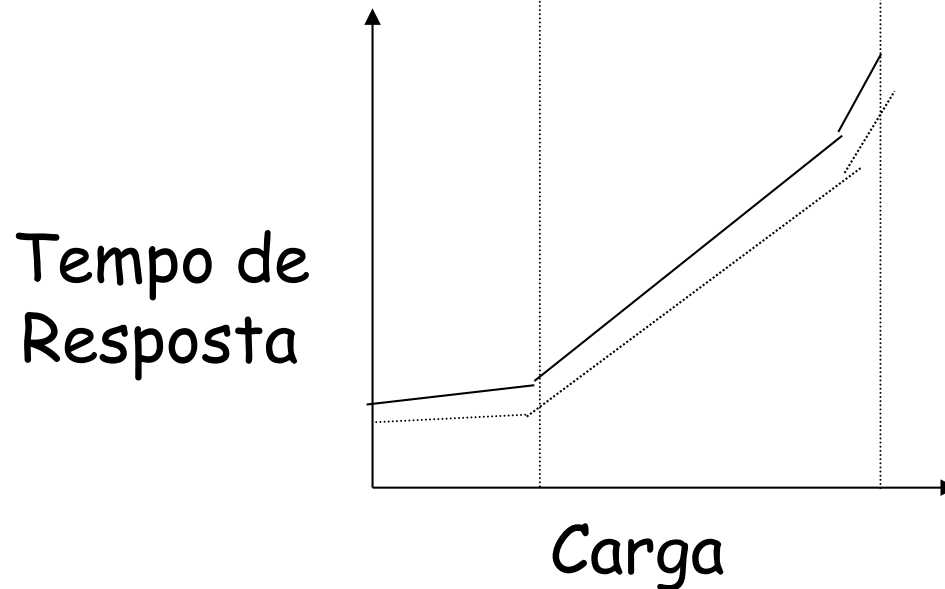
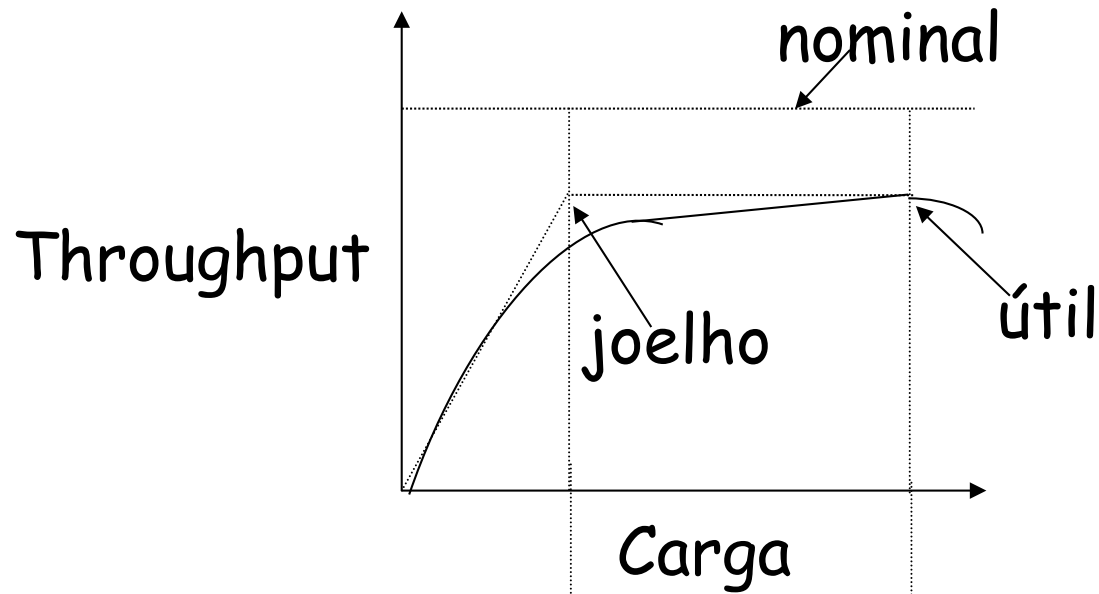
$$TR = TS + TF$$

Slowdown = TF/TS : impacto do tempo de fila

Taxa de Processamento

- Ou taxa de serviço: quantidade de serviço executado por unidade de tempo (throughput)
- Capacidade nominal: capacidade especificada pelo fabricante
 - Ethernet de 100 Mbps
 - Disco com 40Mbps
- Capacidade útil: throughput máximo alcançável
 - Ethernet 100: 70-80 Mbps

Tempo de Resposta x Taxa de Serviço



Joelho da curva =
ponto ótimo de operação

Outras Métricas

- Eficiência: capacidade útil / capacidade nominal
- Utilização : % tempo que recurso está executando serviço
 - Tempo ocioso (*idle time*)
- Custo-benefício = custo / desempenho
 - custo por taxa de serviço
 - US\$/consultas/s, US\$/hits/s
- Métricas específicas
 - % Perda de pacotes, tamanho das rajadas de perdas
 - Qualidade do sinal

Escolha das Métricas

- Incluir métricas para
 - Execução correta, incorreta e não execução
- Avaliar
 - Média, mediana, percentis
 - Variância, coeficiente de variabilidade (CV)
 - Distribuições
 - Medidas individuais, agregadas, por classes

Especificação de Requisitos de Desempenho

- Especificação deve ser precisa e realista
- Problemas:
 - Falta de especificação numérica
o sistema deve ser eficiente
 - Métricas difíceis de avaliar
 - Especificação não realista
o sistema não deve produzir respostas com erros

Especificação de Requisito de Desempenho

- Como fazer:
 1. Escolha um serviço S
 2. Escolha uma métrica M
 3. Escolha um valor máximo X para a métrica M

Opções:

1. média entre valores observados para M para o serviço S deve ser menor que X : arriscado (variabilidade)
3. $Y\%$ (grande) dos valores observados devem ser menores do que X : SIM!!!

Acordo de Nível de Serviço (SLA)

- Exemplos:
 - $RTT < 100$ ms para conexões dentro dos EUA
 - Sistema deve estar disponível X% do tempo
- Exemplos de SLAs de disponibilidade (Os 5 9's)
 - AT&T switches: 2hs de downtime em 40 anos
 - Cisco, HP, MS, Sun: garantem 99.999% de disponibilidade (5 min /ano downtime)

0s 5 9's

In Search of Five 9s – Calculating Availability of Complex Systems

Filed under: Data Center Design,General,Systems — bill @ 5:30 am

I've spent the past few days trying to develop a simple mathematical model to predict the expected availability of complex systems. In IT, we are often asked to develop and commit to service level agreements (SLAs). If the points of failure of the system are not analyzed, and then the system availability calculated, the SLA is flawed from the beginning. To complicate matters further, different people have different definitions of availability. For instance, does scheduled downtime for maintenance count against your system availability calculation?

Common Availability Definitions:

1. Availability = $MTBF / (MTTR + MTBF)$ (Mean Time Between Failure, Mean Time To Recover). This is a classic definition of availability and is often used by hardware manufacturers when they publish an availability metric for a given server.
2. Availability = $(Uptime + Scheduled Maintenance) / (Unscheduled Downtime + Uptime + Scheduled Maintenance)$. This is an IT centric availability metric where the business can support scheduled downtime after hours. This model works for some types of systems, such as a file server that isn't needed at night, but it doesn't work as well for websites, even though many web companies still use this for their SLAs.
3. Availability = $Uptime / (Uptime + Downtime)$. This metric best applies to systems that are needed 24x7 such as e-commerce sites.

Availability is most often expressed as a percentage. Sometimes, people will refer to "four nines" (99.99%) or "five nines" (99.999%). To simplify things, the following table shows the minutes of downtime allowed per year for a given availability level:

Availability	Min Downtime/Year	Hours Downtime/Year
95.000%	26,298	438
98.000%	10,519	175
98.500%	7,889	131
99.000%	5,260	88
99.500%	2,630	44
99.900%	526	8.8
99.990%	52.6	.88
99.999%	5.26	.088

Disponibilidade Custos Downtime

Downtime Cost Per Year

More According to Dunn & Bradstreet, 59% of Fortune 500 companies experience a minimum of 1.6 hours of downtime per week. This means that if you take the average Fortune 500 company (at least 10,000 employees) paid an average of \$56 per hour, including benefits (\$40 per hour salary + \$16 per hour in benefits). The labor part of downtime costs for an organization this size would be \$896,000 weekly, translating into more than \$46 million per year. (*Assessing The Financial Impact Of Downtime*).

\$46 MM
lost from outages per year



Average Outage Period

When an outage occurs, it's a race against time to handle it before it spirals out of control. According to the *IT Process Institute*, resolution time per outage is around 200 minutes. It's really interesting to see just how much time is being put in to resolve outages, when you consider what is happening to the customer experience and company reputation in this time. The average reported incident length was 90 minutes, resulting in an average cost per incident of approximately \$505,500. (*Unplanned IT Outages Cost More than \$5,000 per Minute*)

200min

MTTR per medium outage

IT Process Institute

Downtime Cost Per Hour

On average, the businesses surveyed said they suffered 14 hours of IT downtime per year. Half of those said IT outages damage their reputation and 18% described the impact on their reputation as 'very damaging.' Headlines about IT failures certainly don't help. (*IT Downtime Costs \$26.5 Billion In Lost Revenue*)

InformationWeek
THE BUSINESS VALUE OF TECHNOLOGY

IT Downtime Costs \$26.5 Billion In Lost Revenue

The average downtime costs vary considerably across industries, from approximately \$90,000 per hour in the media sector to about \$6.48 million per hour for large online brokerages. (*How Much Does Downtime Really Cost?*). According to a survey of IT managers, companies are becoming more aware of the direct financial costs of computer downtime. The survey results showed that one in five businesses lose £10,000 an hour through systems downtime. (*Companies count the cost of IT failure*)

IT downtime costs businesses, collectively, more than 127 million person-hours per year—or an average of 545 person-hours per company—in employee productivity. (*IT Downtime Carries a High Pricetag*)

35 percent of survey respondents believe that one hour of downtime for their most business critical applications will cost their company \$25,000 or less, potentially underestimating the adverse impact that IT downtime can have on their entire businesses. Just 10 percent of survey respondents report financial impact of \$150,000 or greater per hour, which is closer to most industry cost estimates — e.g., the Aberdeen Group's estimate of \$110,000 an hour for the average company. (*Working in the dark: financial impact of IT downtime*)

\$42,000
hourly cost of downtime
Gartner

A conservative estimate from Gartner pegs the hourly cost of downtime for computer networks at **\$42,000**, so a company that suffers from worse than average downtime of 175 hours a year can lose more than \$7 million per year. But the cost of each outage affects each company differently, so it's important to know how to calculate the precise financial impact. (*How to quantify downtime*).

<http://www.evolver.com/blog/downtime-outages-and-failures-understanding-their-true-costs.html>

Disponibilidade

Custos Downtime (US\$/hora)

Typical Hourly Cost of Downtime by Industry (in US Dollars)

Brokerage Service	6.48 million
Energy	2.8 million
Telecom	2.0 million
Manufacturing	1.6 million
Retail	1.1 million
Health Care	636,000
Media	90,000

*Sources: Network Computing, the Meta Group and Contingency Planning Research.
All figures in U.S. dollars.*

Parâmetros

- Listar parâmetros que afetam desempenho
 - Parâmetros do sistema: software e hardware
 - CPU, memória, disco, controladora, tamanho de buffer (cache), políticas de escalonamento
 - Parâmetros de carga: usuário (imprevisível-?)
 - Tamanho, tipo e frequência das requisições a serviços
 - Eliminar parâmetros redundantes e/ou normalizar

Ex: servidor de vídeo:

Taxa de chegada λ , Tamanho do arquivo T (minutos)

Impacto no sistema: $N = \lambda T$

Não precisa variar λ e T isoladamente,
mas apenas o produto N

Fatores

- Parâmetros que vão variar no estudo
 - Variação = nível
 - Escolha: parâmetros com maior impacto e controlável
 - Começar com poucos parâmetros e níveis e estender a partir da avaliação dos resultados
 - Preferível maior número de parâmetros e poucos níveis
 - Avaliação inicial do impacto relativo de cada um
 - Refinamentos posteriores

Técnica de Avaliação

- A escolha depende:
 - Escopo (aspectos gerais x detalhes) e estágio
 - Tempo e recursos disponíveis
 - Precisão desejada
 - Seja qual for a escolha:
 - Não acredite nos resultados de simulação até que sejam validados por análises ou experimentos
 - Não acredite nos resultados de modelos analíticos até que sejam validados por simulação ou experimentos
 - Não acredite nos resultados de experimentos até que sejam validados por modelos analíticos ou simulação
- É NECESSÁRIO VALIDAR OS RESULTADOS!!!

Carga de Trabalho

- Carga baseada na lista de serviços do sistema
- Deve ser representativa do sistema real
 - Caracterização das cargas
- Cargas sintéticas vs. cargas reais
 - Cargas reais: traces
 - Cargas sintéticas: modelo baseado em distribuições estatísticas
- É importante definir nível de agregação (classes)

Agregação: Exemplo 1

- Várias tarefas, demanda por CPU

Tarefa	Tempo médio de utilização de CPU (s)
T1	10
T2	0.7
T3	0.04
T4	12
T5	0.8
T6	1
T7	0.5
T8	0.01

Média = 3.13

Agregação: Exemplo 1

- Várias tarefas, demanda por CPU

Tarefa	Tempo médio de utilização de CPU (s)	Classe
T1	10	A
T2	0.7	B
T3	0.04	C
T4	12	A
T5	0.8	B
T6	1	B
T7	0.5	B
T8	0.01	C

Média A = 11 Média B = 0.67 Média C = 0.025

Agregação: Exemplo 2

- Tempo entre chegada de requisições em um servidor de vídeo em um dia típico

Período	Taxa de chegadas (#reqs/min)	Tempo médio entre chegadas (min)
3:00-6:00	0.067	15
7:00-12:00	1.67	0.6
12:00-20:00	3.33	0.3
20:00-3:00	0.033	30

- Agregação também pode ocorrer no tempo

Agregação: Exemplo 3

- Tempo entre chegada de e-mails no servidor central da UFMG em um dia típico

Período	Taxa de chegadas (#e-mails/hora)
1:00-6:00	550
7:00-18:00	1800
19:00-24:00	550

- Agregação também pode ocorrer no tempo

Agregação: Exemplo 4

- % do vídeo assistido em cada interação de um usuário em um vídeo educacional (utilização de banda do servidor)

Tamanho do Vídeo (min)	% Vídeo assistido por interação
0 - 5	57%
5 - 20	17%
20 - 55	7%

Experimentos e Resultados

- Projeto dos experimentos a partir da definição dos fatores e níveis
 - Análise de sensibilidade:
 - E se premissas feitas não forem verdadeiras?
- Análise dos resultados (a luz das premissas)
 - Tratamento estatístico
 - Duração da simulação e/ou experimento suficiente
 - Número de repetições com sementes diferentes para capturar e/ou filtrar variabilidade
 - Importante transformar números em conclusões
- Apresentação: gráficos significativos
- Reavaliar decisões tomadas e possivelmente refazer estudo: novo ciclo