

Administração de Sistemas (ASIST)

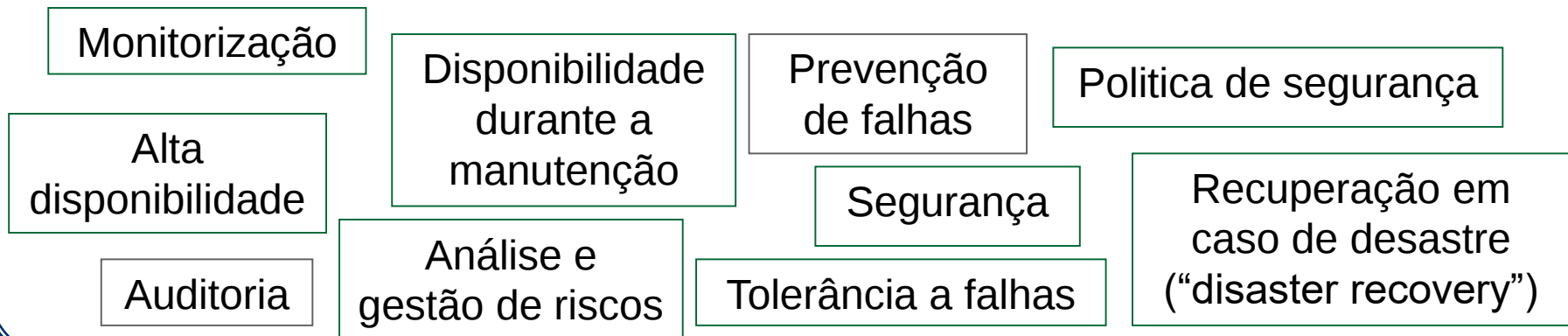
Continuidade de Negócio (“business continuity”).
Disponibilidade.
Prevenção de falhas. Tolerância a falhas. Detecção de falhas.

Setembro de 2014

Continuidade de negócio (“business continuity”)

A “continuidade de negócio” está dependente de muitos fatores, no domínio da administração de sistemas estamos preocupados com o impacto que a infraestrutura tecnológica tem sobre o negócio.

Continuidade de negócio (“business continuity”)

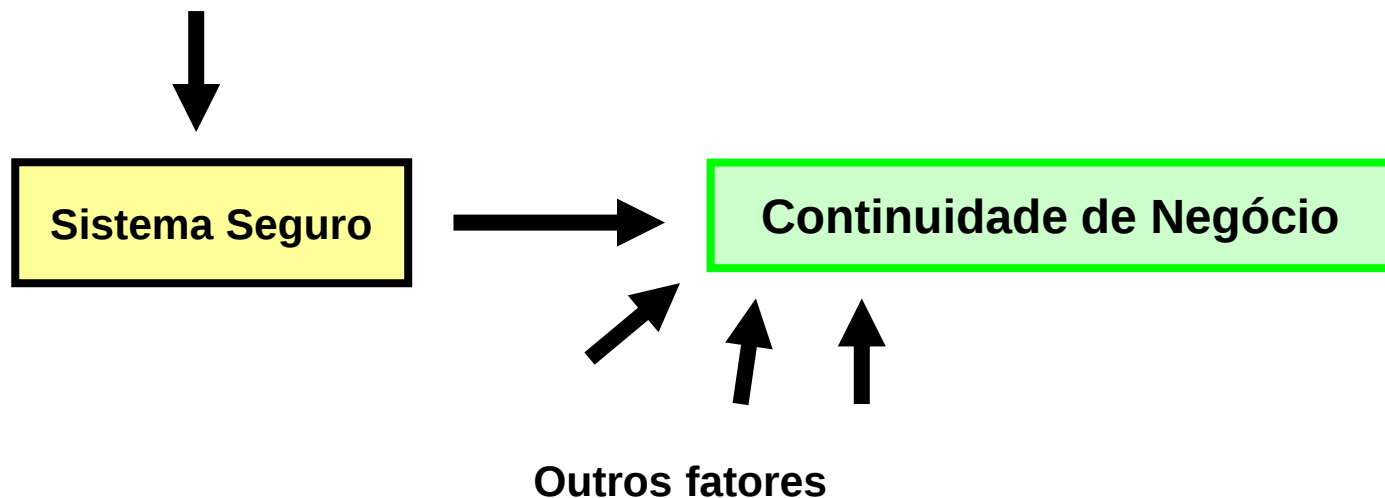


A infra-estrutura tecnológica deve assegurar a “continuidade de negócio”, para isso deve manter-se em funcionamento sem interrupções, dentro dos parâmetros previstos para o negócio suportado por essa infra-estrutura.

Sistema seguro

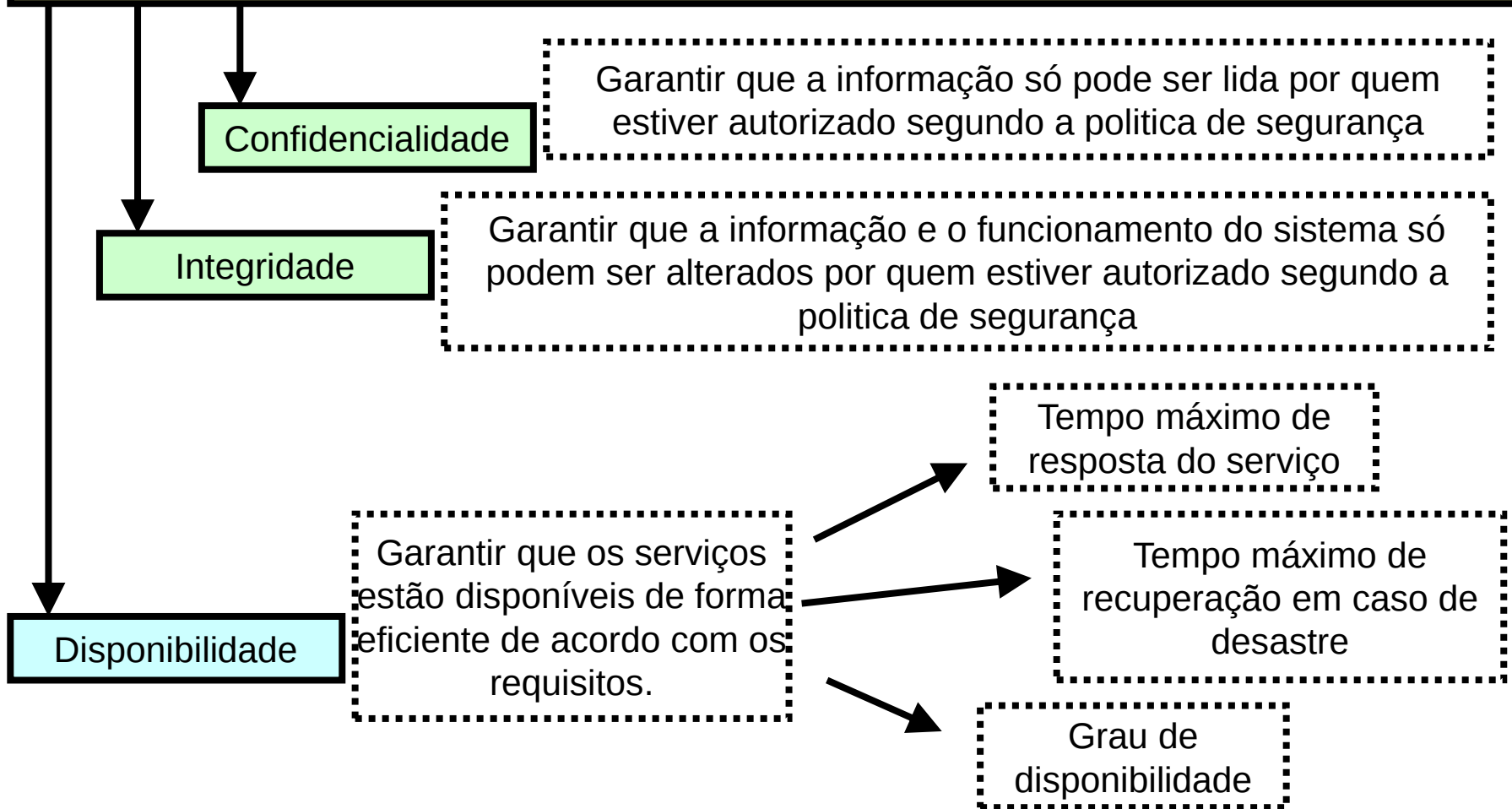
Considera-se um sistema seguro aquele que se mantém em funcionamento dentro dos parâmetros previstos. Qualquer desvio nesses parâmetros é considerado uma falha.

O administrador do sistema tem como missão garantir que o sistema é seguro



Sistema seguro - caracterização

A caracterização do sistema seguro usa um conjunto de parâmetros qualitativos e quantitativos que constituem o SLA (“Service-level agreement”)

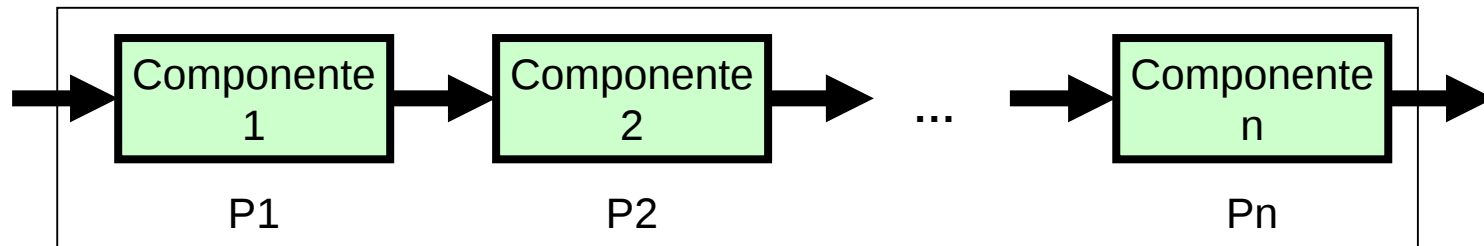


Sistema Seguro – Probabilidade de falha

Planear um “sistema seguro” que garanta a “continuidade de negócio” envolve uma ponderação entre custos e benefícios de modo a obter uma probabilidade de falha aceitável.

Não existem sistemas totalmente seguros ao ponto de haver garantias totais de que nunca ocorre uma falha (0% de probabilidade de falha).

A probabilidade de falha num serviço dependente de vários componentes é igual à somatório das probabilidades de falha em cada um dos componentes. Reduzir o número de dependências é por isso um bom caminho.



$$P = P1 + P2 + \dots + Pn$$

Mean Time Between Failures (MTBF)

Embora a probabilidade de falha seja um útil, é mais usado o MTBF que indica o tempo médio decorrido entre falhas. Normalmente expresso em horas.

A taxa de falhas (λ) é o número médio de falhas por unidade de tempo (por hora). Normalmente é um valor muito reduzido.

$$\lambda = \frac{1}{MTBF}$$

Tal como acontece com a probabilidade de falha, a taxa de falhas num serviço dependente de vários componentes é igual à somatório das taxas de falhas dos vários componentes.

A probabilidade de falha num período de tempo T é dada por:

$$P = 1 - e^{-\lambda T}$$

O parâmetro FIT (“Failures In Time”) indica o número médio de falhas num período de mil milhões de horas (10^9 horas).

$$FIT = 1.000.000.000 \times \lambda$$

$$MTBF = \frac{1.000.000.000}{FIT}$$

Probabilidade de falha - Redundância

A probabilidade de falha de disponibilidade pode ser reduzida com o investimento em redundância.

Considerando um sistema N redundante onde existem N componentes totalmente autónomos que disponibilizam um mesmo serviço, cada um com probabilidade de falha P.

Então a probabilidade de falha do serviço redundante passa a ser: $Pr = (P)^N$

Pressupõe-se uma total eficiência do sistema de gestão de redundância, sem falhas.

Pressupõe-se não se procede à recuperação em caso de falha num componente.

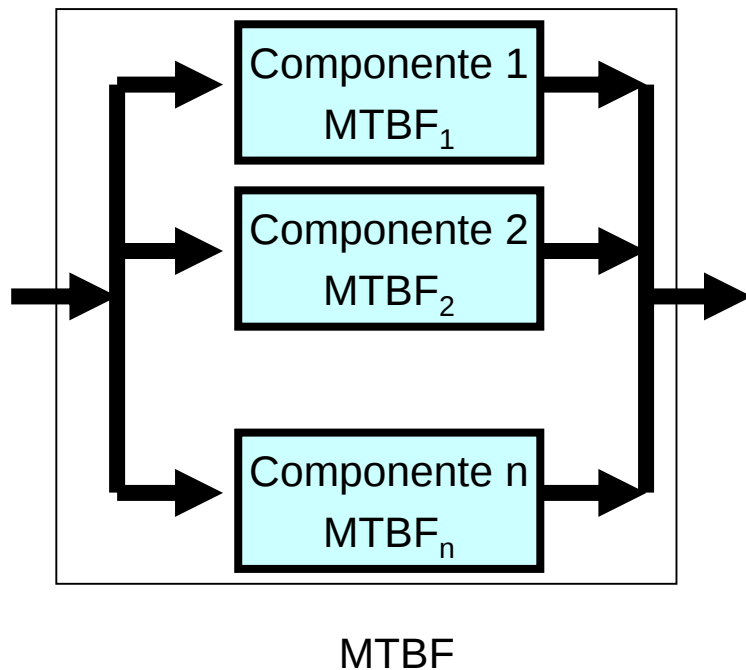
Exemplo: $P = 0,001\%$; $N=3$, então $Pr = 0,000000001 \%$

Para aplicar esta expressão é necessário que a probabilidade de falha (P) do componente não inclua factores comuns tais como falha de alimentação eléctrica, falha nas comunicações e desastres naturais.

Essa independência pode ser garantida por exemplo através de uma razoável separação geográfica.

MTBF - Redundância

O calculo formal do MTBF em sistemas redundantes não é simples, mas considera-se uma boa estimativa considerar que o MTBF de um sistema redundante é igual somatório dos MTBF dos seus componentes.



$$MTBF = MTBF_1 + MTBF_2 + \dots + MTBF_n$$

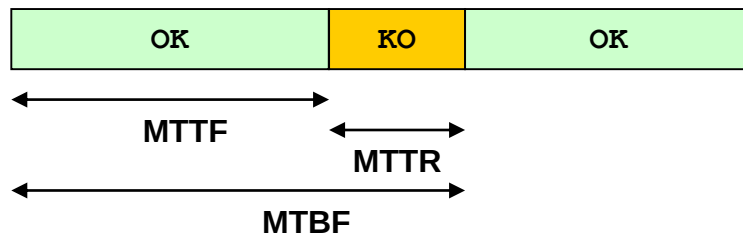
Pressupostos

- Os valores dos MTBF dos vários componentes são independentes entre si.
- A gestão da redundância é 100% eficiente.
- Não há recuperação de elementos em falha

Grau de disponibilidade (“Availability”)

A disponibilidade de um sistema é o rácio entre o somatório dos períodos de tempo em que o sistema opera sem falhas e o tempo total considerado (normalmente um ano, ou um mês).

$$\textit{Disponibilidade} = \frac{\textit{tempo de operação sem falhas}}{\textit{tempo total}}$$



$$\textbf{MTBF} = \textbf{MTTF} + \textbf{MTTR}$$

A disponibilidade pode também ser calculada através do MTBF e do MTTR (“mean time to repair”) ou do MTTF (“mean time to fail”):

$$\textit{Disponibilidade} = \frac{\textit{MTTF}}{\textit{MTTF} + \textit{MTTR}} = \frac{\textit{MTTF}}{\textit{MTBF}} = \frac{\textit{MTBF} - \textit{MTTR}}{\textit{MTBF}}$$

Alta disponibilidade

A alta disponibilidade tem como objectivo aproximar o grau de disponibilidade dos 100%, isso consegue-se aumentando o MTTF e reduzindo o MTTR.

O aumento do MTTF consegue-se com a prevenção de falhas e com a tolerância a falhas.

$$Disponibilidade = \frac{MTTF}{MTTF + MTTR}$$

Reduzir o MTTR, resultado direto do RTO (“Recovery Time Objective”) é uma questão de bom planeamento, nomeadamente através de um DRP (“Disaster Recovery Plan”) bem concebido.

É habitual classificar os sistemas de alta disponibilidade pelo número de dígitos 9 no grau de disponibilidade:

90% - “um nove”
99% - “dois nove”
99,9% - “três nove”
99,99% - “quatro nove”
99,999% - “cinco nove”
99,9999% - “seis nove”
99,99999% - “sete nove”

Prevenção de falhas (“fault avoidance”)

A prevenção de falhas tem como objectivo evitar que as falhas ocorram. Baseia-se em várias medidas de bom senso, entre elas:

- Utilização de componentes de qualidade comprovada
- Controlo ambiental (temperatura; humidade; poeira)
- Controlo de alimentação eléctrica (estabilidade e filtragem)
- Controlo de acesso físico, incluindo linhas de comunicação
- Controlo de acesso remoto (“firewall”, autenticação)
- Prevenção e combate a incêndio
- Execução de testes exaustivos antes de colocar componentes em operação
- Simplificar administração dos sistemas, por exemplo com virtualização
- Controlo de permissões e privilégios de administração
- Divulgação da “Politica de Segurança” e formação dos utilizadores e operadores
- Aplicação de todas as actualizações ao “software”
- Garantias de autenticidade (mecanismos de autenticação sólidos)
- Monitorização (permite detetar potenciais pontos de falha)
- Controlo de utilização de recursos (limitação / reserva). Ex.: CPU; RAM; DISCO; REDE

Tolerância a falhas (“fault tolerance”)

A tolerância a falhas total (“full fault tolerance”) garante que a falha de um componente não tem qualquer impacto nos parâmetros de funcionamento, nomeadamente na disponibilidade.

Exemplos: RAID1; servidores UDP contactados por “broadcast”

A tolerância a falhas é habitualmente conseguida através da redundância de componentes, mas nem sempre é possível uma substituição perfeita e instantânea do componente em falha.

Nesse caso assiste-se a uma degradação temporária dos parâmetros de funcionamento (“graceful degradation”).

Se essa degradação for significativa ou prolongada, o sistema passa a designar-se “fail soft” e não “fault tolerant”.

Exemplos: “server farm / server cluster”; “spanning tree”; “dynamic routing”; “sparing redundancy”

Um sistema designa-se “fail safe” se a falha provoca a indisponibilidade, mas sem comprometer a sua integridade.

Exemplo: UPS sem gerador

Ponto único de falha (SPOF – “single point of failure”)

Aumentar a disponibilidade através da redundância de recursos é dispendioso e por isso a eficácia deve ser cuidadosamente avaliada. Os alvos prioritários são:

- componentes que dão suporte aos serviços mais importantes para o negócio
- componentes com MTBF menor, pois o MTBF de um sistema dependente de vários componentes é igual ao menor MTBF dos seus componentes.

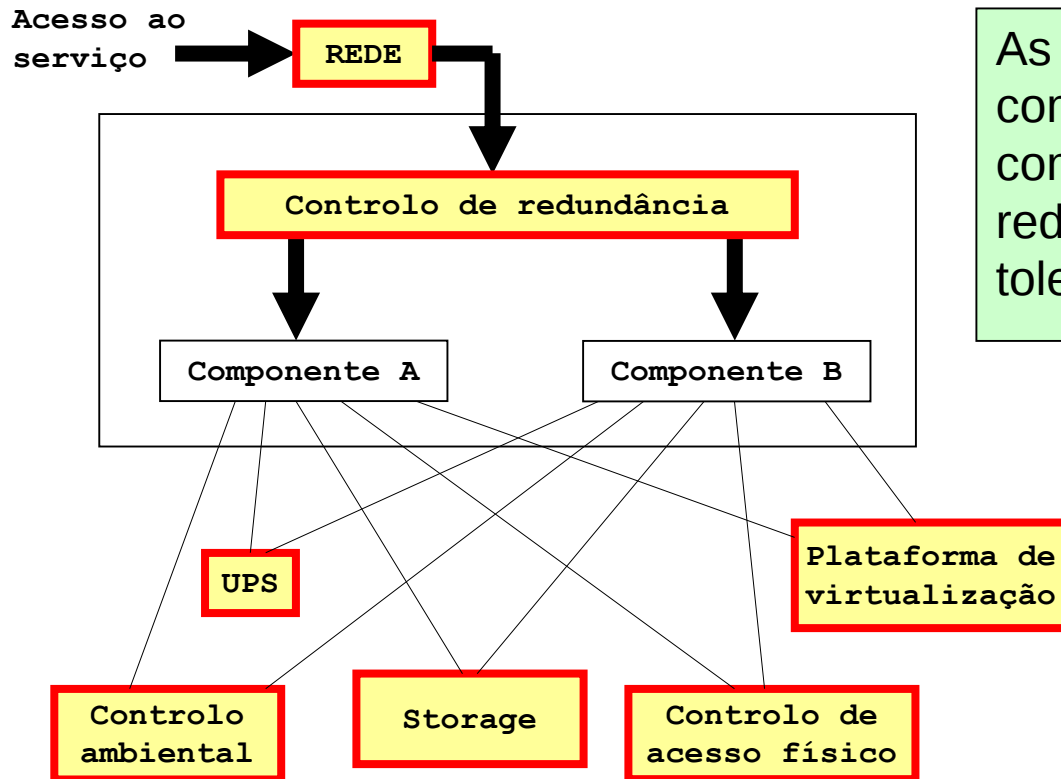
A redundância pode ser totalmente comprometida se os componentes redundantes forem afectados por uma dependência comum. Esta dependência designa-se “ponto único de falha” (SPOF).

A maioria dos sistemas redundantes possuem um SPOF intrínseco que é o próprio sistema de controlo de redundância.

Os SPOF devem ser procurados e eliminados, em especial aqueles com baixo MTBF. A forma de eliminar o SPOF depende dos casos concretos, mas no limite o próprio SPOF pode ser tornado redundante e desta forma eliminado.

“Common-mode fault”

Uma “common-mode fault” é uma falha que afecta simultaneamente mais do que um dos componentes de um sistema redundante. Pode ser causada directa ou indirectamente por um SPOF, mas também pode ter outras causas tais como “bugs” de “software” ou ações maliciosas.



As dependências comuns aos componentes do sistema redundante comprometem as vantagens da redundância sob o ponto de vista de tolerância a falhas.

A solução que dá mais garantias é distribuir geograficamente os vários componentes.

RAID (redundant array of independent disks)

O RAID é um exemplo comum de tolerância a falhas baseada em redundância.

A configuração mais directa e flexível é o RAID 1 (“mirroring”), que usa um conjunto de “n” discos idênticos (no mínimo 2) todos contendo a mesma informação. É capaz de suportar a falha simultânea de “n-1” discos. O MTBF do “array” é superior ao somatório dos MTBF dos vários discos que o compõem.

Tipicamente os discos do “array” estão instalados num sistema “hot-swap” que permite a troca imediata de um disco em falha sem qualquer interrupção de funcionamento.

A maioria dos sistemas operativos suporta RAID por “software”, ou seja é o sistema operativo que simula e gere o “array” baseando-se em discos ligados a um controlador de disco normal.

O RAID por “hardware” utiliza um controlador RAID que gere autonomamente o “array” de uma forma mais eficiente e transparente para o sistema operativo. As falhas são resolvidas pelo controlador sem que o sistema operativo se aperceba disso e sem qualquer degradação de performance.

Principais níveis RAID

RAID 0

Bloco 1	Bloco 2
Bloco 3	Bloco 4
Bloco 5	Bloco 6
Disco A	Disco B

RAID 1

Bloco 1	Bloco 1
Bloco 2	Bloco 2
Bloco 3	Bloco 3
Disco A	Disco B

Embora o RAID 1 seja simples e flexível permitindo obter valores de MTBF muito elevados, tem dois inconvenientes:

- não melhora a velocidade, em especial a de escrita;
- o espaço disponível no “array” é sempre igual ao de um único disco

RAID 2

Bit 1	Bit 2	Paridade
Bit 3	Bit 4	Paridade
Bit 5	Bit 6	Paridade
Disco A	Disco B	Disco C

RAID 3

Byte 1	Byte 2	Paridade
Byte 3	Byte 4	Paridade
Byte 5	Byte 6	Paridade
Disco A	Disco B	Disco C

RAID 4

Bloco 1	Bloco 2	Paridade
Bloco 3	Bloco 4	Paridade
Bloco 5	Bloco 6	Paridade
Disco A	Disco B	Disco C

RAID 5

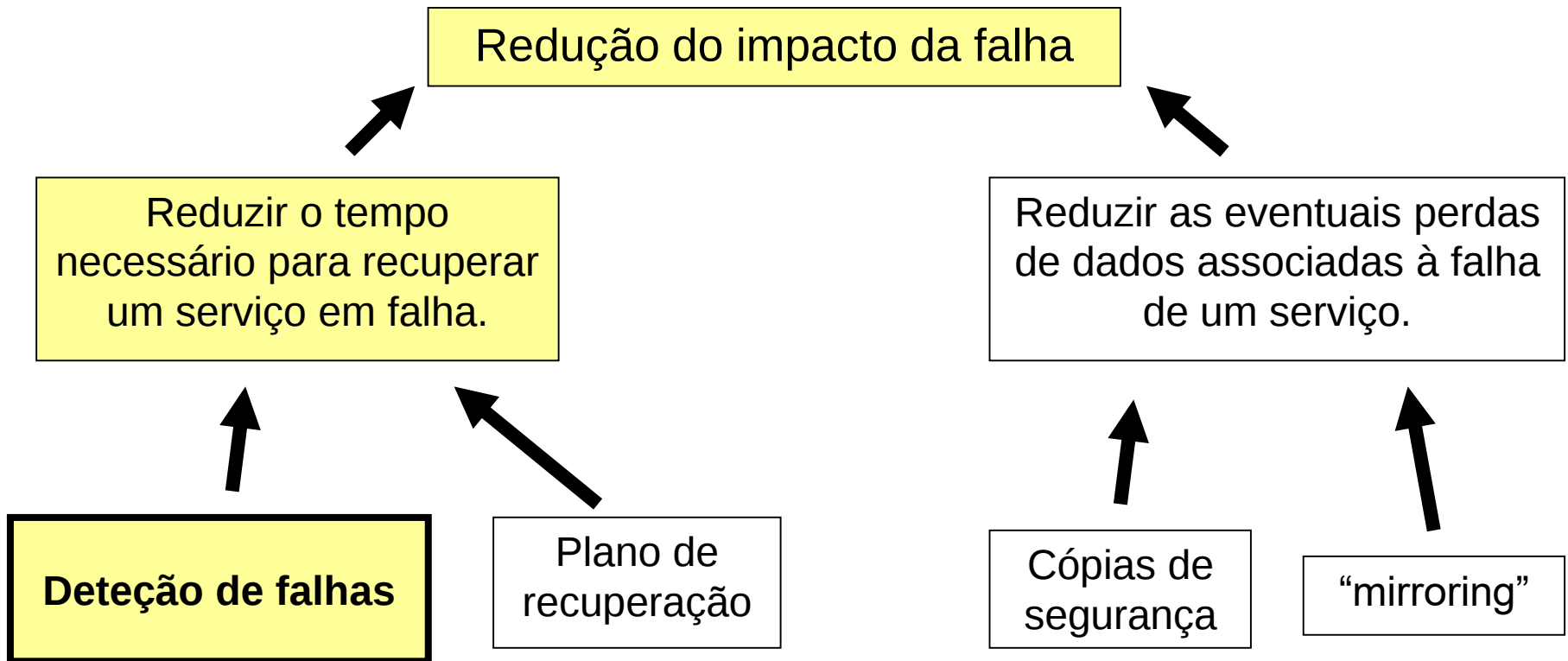
Bloco 1	Bloco 2	Paridade
Bloco 3	Paridade	Bloco 4
Paridade	Bloco 5	Bloco 6
Disco A	Disco B	Disco C

RAID 6

Bloco 1	Bloco 2	Paridade	Paridade
Bloco 3	Paridade	Paridade	Bloco 4
Paridade	Paridade	Bloco 5	Bloco 6
Disco A	Disco B	Disco C	Disco D

Deteção de falhas

Por muito cuidadas que sejam as medidas tomadas nos domínios da “prevenção de falhas” e da “tolerância a falhas”, as falhas não podem ser totalmente eliminadas, o último recurso é a “redução do impacto das falhas”.



Deteção de falhas - monitorização

A deteção de falhas deve funcionar de forma automatizada 24x7, consiste na execução periódica de testes sobre os componentes da infra-estrutura informática:

- tempos de resposta de serviços
- medições (temperaturas, etc.)
- volumes e tipos de tráfego de rede
- estado interno de dispositivos
- anomalias nos registos de atividades
- deteção de anomalias e intrusos

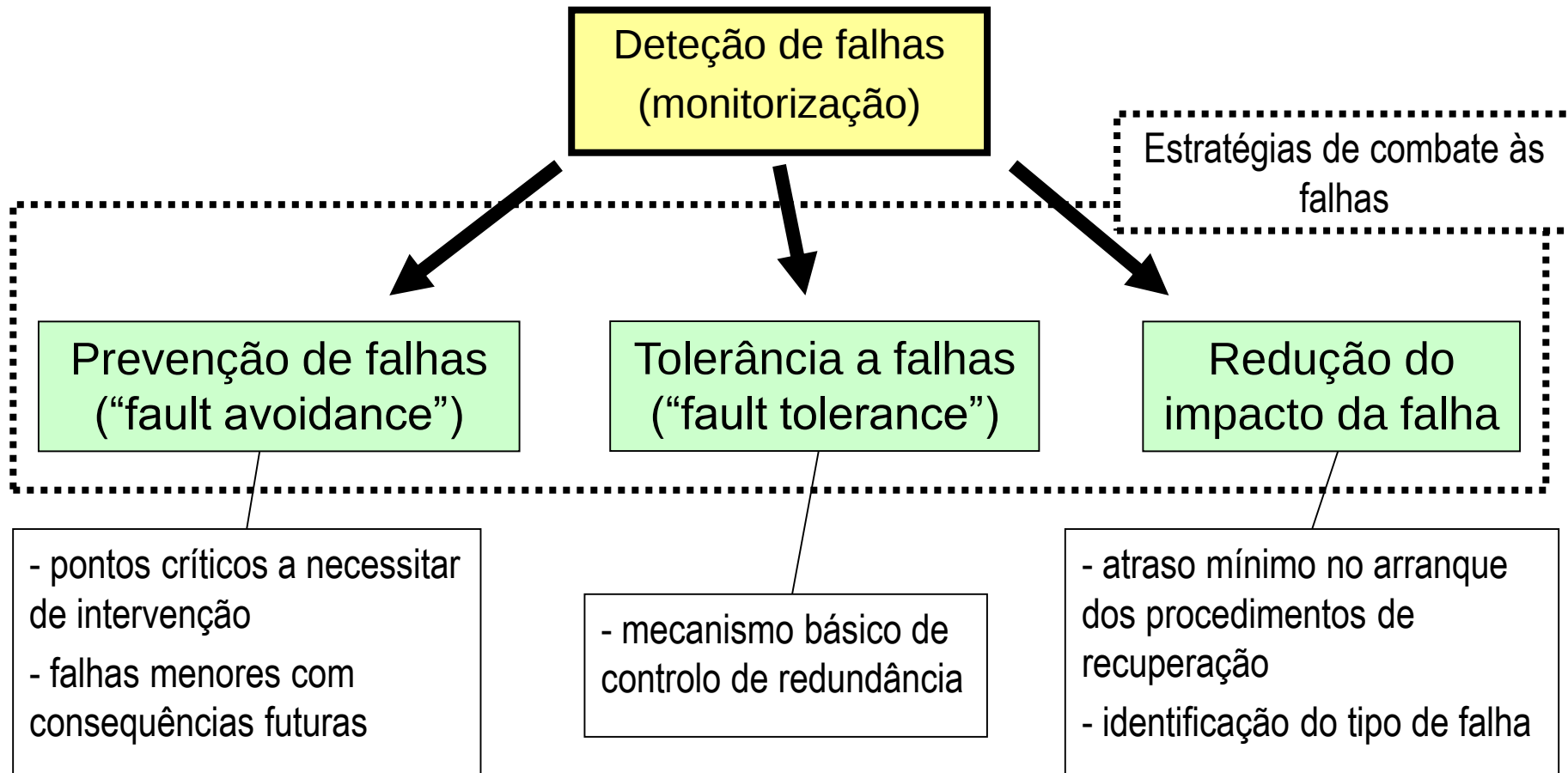
A periodicidade de cada teste deve ser ajustada ao nível de risco associado ao componente em causa.

Uma vez detectada uma anomalia o sistema de monitorização deve notificar os administradores com a maior brevidade possível para que seja desencadeado o processo de recuperação. Normalmente usa-se correio electrónico, mas é preferível complementar essa opção com uma forma de mensagens instantâneas.

Em alguns sistemas poderá ser possível definir mecanismos de recuperação automática para algumas anomalias.

Deteção de falhas

Por motivos diversos, a deteção de falhas está directamente envolvida nas 3 vertentes complementares de combate às falhas.



Referências bibliográficas

[1] Lhamas A. (2010-2011). Aulas Teóricas de ASIST.

(...)