

Administração de Sistemas (ASIST)

Segurança de redes

Novembro de 2014

Segurança em redes

As redes são por natureza um meio privilegiado para condução de ataques:

- sendo meios de transmissão de informação, podem ser usadas para atacar remotamente os sistemas que estão salvaguardados de acesso físico num CPD.
- são extensas logo é muito difícil controlar de forma eficiente o acesso físico, tornando-se mesmo uma missão impossível para as redes sem fios. Embora o controlo de acesso físico não ofereça garantias, nunca deve ser descurado.

A autenticação e encriptação são duas ferramentas fundamentais para contrariar muitos dos ataques, mas podem não ser suficientes.

A segmentação das redes em zonas de nível de segurança distintos é fundamental, tipicamente podem ser consideradas três zonas:

- Rede do CPD (onde se encontram os servidores)
- Rede interna dos utilizadores (“intranet”)
- Redes exteriores (“internet”)

A separação entre zonas é assegurada através da interligação por encaminhadores que analisam e filtram a informação, estes dispositivos são designados “firewalls”.

Rede – Acesso Físico

Nas redes sem fios o controlo de acesso físico é totalmente impossível. Embora atualmente as redes locais com cabos recorram à comutação de pacotes no nível 2 (ex.: Ethernet), esses comutadores não separam domínios de “broadcast” e o seu funcionamento pode ser comprometido. Sob o ponto de vista de segurança, este tipo de infra-estrutura deve ser sempre considerada equivalente a uma rede de meio de transmissão partilhado: qualquer pacote emitido num dado nó é entregue em todos os outros nós da rede.

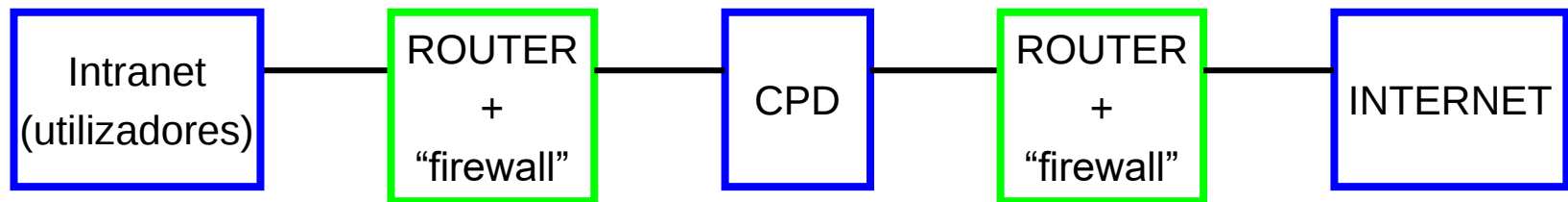


A separação entre redes assegurada pelas VLAN é meramente lógica, as VLAN apenas dão algumas garantias de segurança se apenas forem usadas exclusivamente nos comutadores da infra-estrutura com acesso físico vedado.

A comutação de pacotes de nível 3 (“routing”) garante que os pacotes apenas são transferidos quando pertencem ao protocolo correcto e o seu endereço de destino está de acordo com a configuração (“tabela de routing”). Os “routers” asseguram por isso algum grau de proteção sob o ponto de vista de segmentação. É neste tipo de dispositivo que são normalmente instalados os “firewalls”.

Controlo de acesso pela rede – “firewall”

Os controlo de acesso pela rede tem como objetivo condicionar as transferências de informação entre diferentes zonas das infra-estrutura de rede, por isso normalmente são implementados nos nós intermédios que asseguram a respetiva ligação, tipicamente os “routers”. O “software” que assegura este controlo de acesso é designado de “firewall”.



Os nós finais (clientes e servidores) também devem implementar um “firewall”, neste caso para controlar a saída de dados e especialmente a entrada de dados no nó.



Os nós intermédios (“routers”) que implementam a funcionalidade de “firewall” são eles próprios muitas vezes designados de “firewalls”.

Tipos de “firewall” – “Packet Filter”

Um “firewall” analisa o tráfego de rede e com base num conjunto de regras pré-estabelecidas pelo administrador determina se a informação deve passar ou ser bloqueada (descartada).

Designa-se “firewall” estático, “packet filter” ou “screening router” um “firewall” em que as regras são estáticas e o seu comportamento é sempre igual independentemente das comunicações existentes.

As regras de filtragem são definidas com base em ACLs (“Access Control Lists”) que são percorridas sequencialmente para cada pacote analisado. Idealmente a política da ACL é bloquear, ou seja, chegando ao fim da lista sem encontrar nenhuma regra que permita explicitamente a passagem do pacote, ele deve ser bloqueado.

Normalmente o “packet filter” é implementado num “router” e analisa as propriedades do pacote na camada do protocolo de rede, tipicamente IP. Existem no entanto dispositivos de nível 2 (“switches”) capazes de exercer esta função.

As regras são definidas com base nas propriedades observadas no pacote, em especial: interfaces de entrada e saída, endereço de origem e destino, protocolo de transporte e números de porto de origem e destino (serviços de rede).

“Firewall” dinâmico – “Stateful Firewall”

A principal característica de um “firewall” dinâmico é a de detetar contextos específicos (tipicamente ataques maliciosos) e reagir criando regras temporárias nas ACLs.

Continuam a existir regras estáticas de filtragem pacote a pacote como as do “packet filter”, mas adicionalmente são previstas situações anómalas para as quais serão automaticamente inseridas regras adicionais. Esta capacidade é designada “stateful packet inspection” (SPI).

Embora os “firewall” dinâmicos possam ser implementados em vários níveis do modelo OSI, tipicamente analisam informação acima da camada de rede, incluindo até dados transportados pelos protocolos de aplicação como HTTP ou SMTP.

Enquanto um “packet filter” analisa cada pacote individualmente, fora de contexto, os “firewall” dinâmicos fazem o seguimento da todas as trocas de informação, mantendo informação de contexto sobre conexões e sessões entre clientes e servidores.

Os “firewalls” dinâmicos assumem o papel de ferramentas de deteção de intrusos.

“Firewalls” de aplicação e “Proxies”

Um “firewall” de aplicação tem princípios de funcionamento semelhantes ao de um “firewall” dinâmico, mas opera mais claramente na camada de aplicação analisando em detalhe as informações transferidas através dos protocolos desse nível.

Um “proxy” é um serviço de rede que funciona com intermediário no acesso a serviços de aplicação. Por exemplo no caso de um “proxy” HTTP o cliente (“browser”) não acede diretamente aos servidores HTTP, solicita ao “proxy” a página pretendida através do respetivo URL, o “proxy” obtém a página do servidor e fornece-a ao cliente.

A principal vantagem dos “proxy” é que podem manter cópias locais da informação e desta forma otimizar a utilização da ligação à “internet”. Sob o ponto de vista de segurança são o local ideal para implementar o “firewall” de aplicação.

Para cada protocolo de aplicação é necessário um “proxy” distinto, além disso muitos protocolos de aplicação não suportam o uso de “proxies”, por esta razão, em muitos casos, o “firewall” de aplicação não é associado a nenhum “proxy”.

DMZ – (“DeMilitarized Zone”)

O objetivo dos “firewalls” implementados em nós intermédios é a separação de zonas de infra-estrutura de rede com diferentes requisitos de segurança.

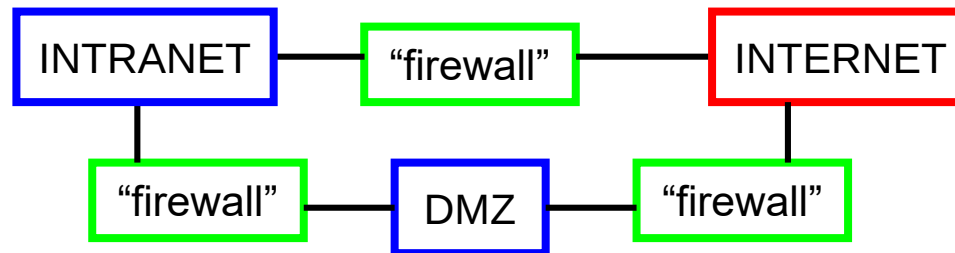
Normalmente existem três zonas claramente distintas que merecem esta separação:

- INTERNET – zona sobre a qual a organização não tem controlo e na qual se espera que tenham origem a maioria dos ataques.
- INTRANET – zona de trabalho dos utilizadores da organização. Também pode ser a origem de ataques, pelos próprios utilizadores ou por intrusos já que o controlo de acesso físico nesta zona pode ser precário. Por outro lado de forma não intencional os postos de trabalho podem ser infetados por “malware” que conduz os ataques.
- DMZ – zona do CPD onde se encontram os servidores, deve ter por isso o nível mais elevado de segurança, os “firewalls” que rodeiam a DMZ devem ser o mais restritivos possível, permitindo apenas a passagem do que é estritamente necessário.

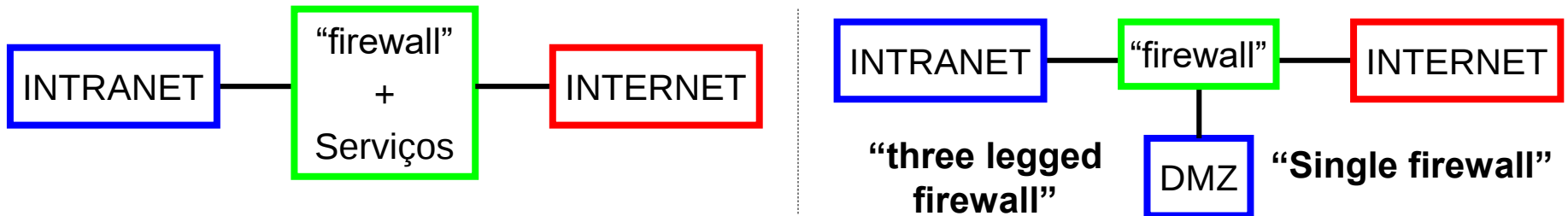
A DMZ e a INTRANET podem, ser subdivididas em mais zonas, por exemplo criando na INTRANET uma zona separada para a administração ou outros departamentos.

DMZ – Arquiteturas

Sob o ponto de vista lógico (regras aplicadas), as três zonas devem estar sempre separadas entre si por um “firewall”:



Sob o ponto de vista físico podem ser adotadas várias arquiteturas:



“Dual firewall”



Ataques – Inspeção Furtiva (“sniffing”)

É um ataque passivo em que o intruso se limita a observar os pacotes na rede, compromete a confidencialidade. Pode também ser usada como ponto de partida para levar a cabo outros tipos de ataque.

O controlo de acesso físico não resolve o problema pois teria de ser garantido em todo o percurso da informação, a única forma eficaz de contrariar o “sniffing” é recorrer à encriptação usando protocolos de comunicação seguros.

A monitorização de tráfego pelo administrador da rede é realizado através de um analisador de rede, o objetivo é recolher dados estatísticos sobre os tipos de pacotes que estão a ser usados e detetar ações maliciosas (deteção de intrusos). Trata-se de uma atividade legal integrada nos mecanismos de deteção de falhas.

A monitorização de tráfego deve preservar a privacidade, pode ser realizada internamente num dispositivo de interligação de nível 3 (“router”) ou através de um dispositivo externo (“packet analyzer”) ligado à porta de um comutador de nível 2. Para ligar um analisador de rede à porta de um comutador é necessário desactivar a filtragem nessa porta, colocando a porta no modo “Monitor” ou “Mirror”.

Ataques – Interposição (MITM – “Man-in-the-middle”)

O atacante interceta os dados de tal forma que além de os poder observar (“sniffing”) também pode alterar o seu conteúdo (compromete a integridade), o atacante terá de possuir uma posição privilegiada que lhe permite observar e opcionalmente alterar os dados que estão a ser transmitidos através da rede sem que os utilizadores legítimos se apercebam. Por exemplo no interior de um “router”.



Para conseguir a posição privilegiada de interceção das comunicações o atacante pode:

- tentar atacar um dispositivo já existente conseguindo acesso interno ao mesmo
- inserir na rede um novo dispositivo através de um ataque de usurpação de identidade, por exemplo fazendo-se passar por “router” legítimo da rede de um dos nós, ou fazendo-se passar por ponto de acesso de rede sem fios.

Estes ataques apenas são possíveis por existirem falhas de autenticação, a forma de os contrariar é por isso reforçar a autenticação. Tratando-se de um ataque ativo, a deteção de intrusos também pode ajudar.

Ataques – Negação de serviço (“Denial of Service” - DoS)

A atacante simplesmente inunda o serviço com um número muito elevado de pedidos “normais”. O objetivo é degradar o tempo de resposta do serviço até que eventualmente ele fique indisponível.

Os ataques DoS são difíceis de contrariar porque se assemelham a uma situação normal de sobrecarga do serviço.

Para contrariar um ataque DoS sem afetar a disponibilidade do serviço é necessário arranjar uma forma de distinguir os pedidos legítimos e apenas eliminar os pedidos do atacante. Esta distinção pode ser conseguida por deteção de:

- um número muito elevado de pedidos com o mesmo endereço de origem;
- um número muito elevado de pedidos todos iguais, sem qualquer seguimento;
- pedidos formalmente incorretos (por vezes com o objetivo de comprometer o serviço).

Em todos estes casos há necessidade de recorrer a um “firewall” dinâmico, que seja capaz de reconhecer o ataque e alterar as suas regras de forma a contrariar o ataque. Um “firewall” estático (com regras estáticas) não é suficiente.

Ataques – DDoS (“Distributed Denial of Service”)

O principal meio de combate aos ataques DoS é através de um “firewall” dinâmico que deteta um elevado número de pedidos todos com o mesmo endereço de origem, nesse caso é criada uma regra temporária para bloquear esse endereço.

O ataque DoS distribuído recorre a um número elevado de diferentes endereços de origem para contornar esta defesa.

Podem tratar-se de máquinas comprometidas através de “malware” que lançam um ataque sincronizado a determinada data/hora.

Os ataques DDoS mais simples podem ser desencadeados por um único nó, mas nesse caso os endereços serão forjados (“spoofing IP”) pelo que o atacante não poderá obter respostas nem estabelecer completamente as ligações TCP.

O ataque “SYN flood” é um destes casos em que o atacante emite apenas pedidos de estabelecimento de ligação TCP (“SYN”), sem nunca dar qualquer sequência aos mesmo pois nunca recebe os “ACK”.

Outros ataques DoS

Ataque DDoS refletido – o atacante envia pedidos para um conjunto vasto de servidores, esses pedidos têm o endereço de origem forjado (“spoofing IP”) de forma a corresponder ao da vítima, isso leva a que a vítima seja inundada com respostas de todos os servidores atacados.

As medidas defensivas DoS também podem ser aproveitadas pelos atacantes, se o objetivo do atacante for impedir o acesso de uma máquina A ao serviço, pode realizar um ataque DoS ao serviço com endereço de origem forjado, correspondente à máquina A. O “firewall” do serviço entende que se trata de um ataque DoS e bloqueia o endereço da máquina A.

De uma forma geral as medidas defensivas contra DoS podem causar problemas porque sobre tráfego elevado podem ocorrer falsas deteções. Grande parte dos postos de trabalho atuais encontra-se em redes privadas (NAT) isso leva a que sob o ponto de vista do “firewall” todos os respetivos pedidos aparentem ter o mesmo endereço de origem.

Ataques – Usurpação de identidade

A atacante faz-se passar por quem não é. Muitas vezes tento em vista de seguida um outro tipo de ataque como “sniffing”, MITM ou DoS.

“IP Spoofing” – consiste em emitir pacotes IP com endereço de origem falso, é uma técnica vulgarmente usada nos ataques DoS para escapar à deteção pelos “firewall”.

Os ataque “IP Spoofing” podem ser totalmente evitados nas redes de origem. É da responsabilidade dos administradores das redes garantir que todos os pacotes enviados para a “internet” têm como endereço de origem as redes locais. Isto é facilmente implementado com um simples “firewall” estático.

A ausência de autenticação em muitos serviços de rede torna-os vulneráveis a ataques de usurpação de identidade. A utilização de protocolos seguros, com autenticação, é a forma de resolver o problema, mas nem sempre estão disponíveis.

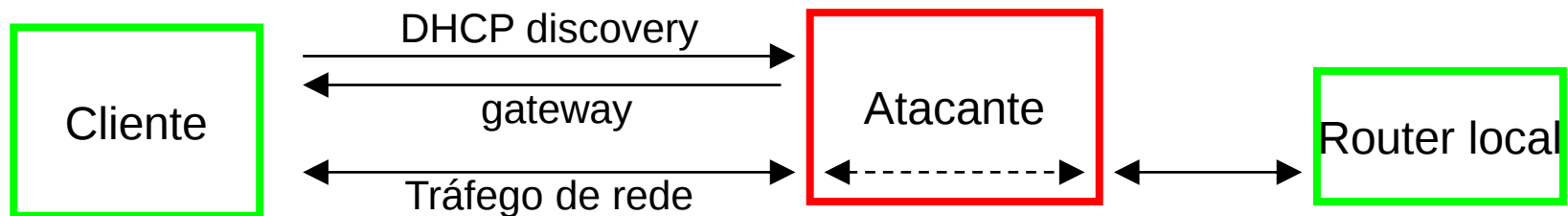
Os serviços de rede local que recorrem a “broadcast” estão particularmente expostos como por exemplo o DHCP e a resolução de endereços IP via ARP.

Ataque MITM ao serviço DHCP

O serviço DHCP não tem autenticação, o cliente aceita a primeira resposta ao pedido “DHCP discovery” que emitiu em “broadcast”.

O atacante introduz na rede um servidor DHCP falso, quando este servidor é adotado por um cliente, são fornecidos dados falsos para continuar o ataque MITM.

“router” falso – é fornecido ao cliente um endereço numa rede fictícia e respectivo “gateway” que corresponde a um “router” falso do atacante. O “router” do atacante encaminha o tráfego através do “router” legal pelo que o cliente não se apercebe de nada de anormal.



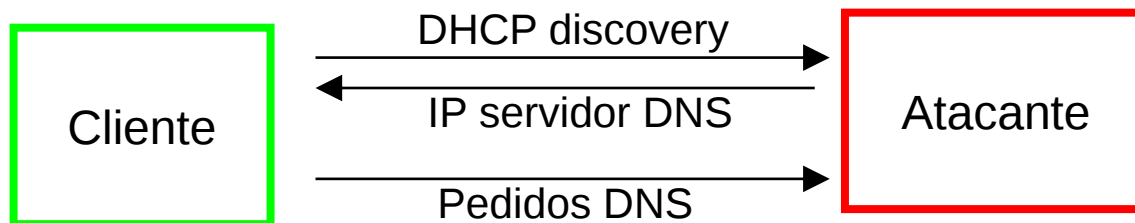
O atacante realiza a tradução dos endereços de origem (SNAT), isso torna mais difícil a sua deteção.

Ataque MITM ao serviço DHCP – DNS falso

É um ataque semelhante ao anterior, mas o objetivo é atacar o sistema de resolução de nomes.

Servidor DNS falso – são fornecidos ao cliente dados corretos para a rede em que se encontra, apenas o endereço do servidor de nomes DNS é falsificado.

Ao fornecer ao cliente um endereço de servidor de nomes correspondente a uma máquina controlada pelo atacante, este passa a controlar os endereços IP que são indicados ao cliente sempre que ele pretende resolver um nome.



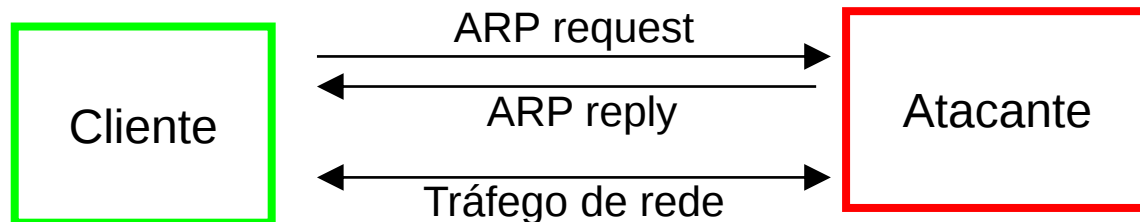
Por exemplo, o utilizador pode indicar um determinado nome no URL, mas esse nome não é transformado no endereço IP correcto, em seu lugar recebe o endereço que o atacante pretender, tipicamente o de um servidor falso. Desta forma o utilizador não se apercebe que não está a aceder ao servidor legítimo.

ARP “spoofing”

O ARP é usado para transformar endereços IP em endereços MAC (“ethernet”) dentro de uma mesma rede IP. Não dispõe de mecanismos de autenticação, por isso um intruso pode através dele anunciar-se como dono de um qualquer endereço IP da rede.

É particularmente tentador ao intruso anunciar-se como o detentor do endereço IP do “gateway” da rede, fica assim na posição apropriada para um ataque MITM.

O intruso fica a aguardar por um pedido ARP de resolução do IP correspondente ao “router” local, envia então várias respostas ARP, fornecendo o seu próprio endereço MAC e tentando anular a resposta do legítimo detentor do endereço.



Tendo sucesso a vítima começa a enviar os dados para o endereço MAC fornecido, pensando que se trata do endereço IP que solicitou.

ARP “poisoning”

O envenenamento ARP, tira proveito de algumas simplificações e optimizações que foram criadas em torno das implementações ARP. Muitas implementações observam a rede e “aproveitam” todas as respostas ARP que encontram na rede para actualizar a sua tabela ARP.

Para este tipo de implementações o atacante nem sequer necessita de esperar que surja o pedido ARP, pode desde logo começar a enviar respostas e introduzir na tabela ARP da vítima os endereços que pretende. Quando a vítima pretende usar o “IP”, verifica que ele se encontra na tabela e nem sequer emite nenhum pedido.

As implementações ARP permitem a alteração de alguns parâmetros de funcionamento que afectam a segurança.

Uma medida drástica para evitar os ataques ao ARP é usar tabelas ARP estáticas, a tabela é carregada de um ficheiro (normalmente designado “ethers”) e as entradas correspondentes são marcadas como permanentes não podendo ser sobrepostas por outras. Embora forneça algumas garantias a gestão é difícil. Mesmo que apenas sejam colocados na tabela estática os endereços IP mais importantes todos os ficheiros têm de ser mantidos atualizados.

“Spoofing IP” da INTRANET

O “firewall” INTERNET/DMZ é muitas vezes mais sólido do que o “firewall” INTRANET/DMZ é por isso tentador para os atacantes posicionarem-se na INTRANET.

O posicionamento na INTRANET pode ser apenas simulado através de “spoofing IP”, neste caso os pedidos terão como origem a INTERNET, mas transportam um endereço de origem falso, correspondendo à INTRANET. Isto pode ser facilmente evitado através de regras estáticas no “firewall” de ligação à INTERNET.

As regras para contrariar o “spoofing IP” são muito simples e devem existir em todos os “firewalls”, verificam se os endereços de origem são correctos:

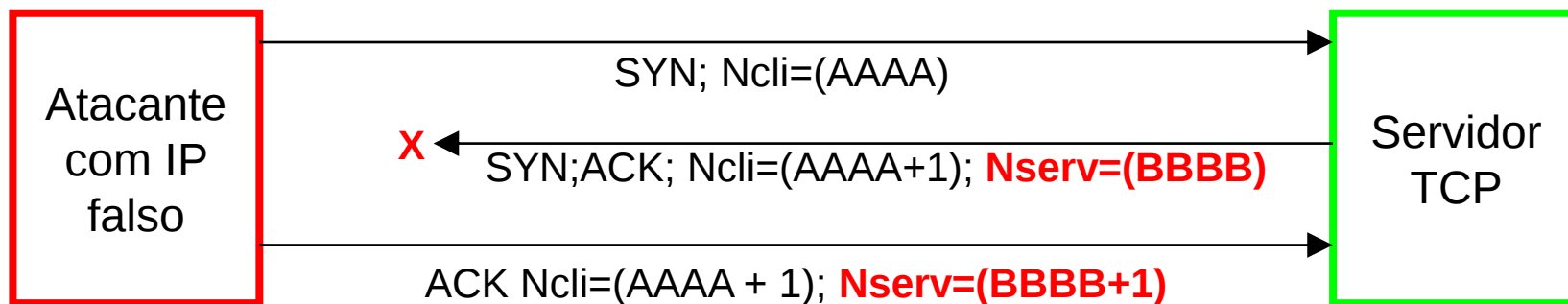
- no tráfego de entrada numa rede devem ser eliminados todos os pacotes com endereço de origem igual ao da rede.
- no tráfego de saída de uma rede devem ser eliminados todos os pacotes com endereço de origem diferente do da rede.

Uma forma de um atacante se posicionar realmente na INTRANET é através da introdução de “malware”, com a “ajuda” de algum utilizador menos cuidadoso.

“Spoofing IP” sobre conexões TCP

O estabelecimento de ligações TCP com “spoofing” do endereço IP é dificultado pelos números de sequência do protocolo TCP.

Ao usar um endereço IP de origem falso, o atacante fica impossibilitado de receber as respostas do servidor, para conseguir estabelecer a ligação tem de “adivinhar” quais foram as respostas, particularmente o número de sequência do servidor.



Os números de sequência iniciais (AAAA e BBBB) são gerados pelas duas entidades e são diferentes para cada ligação, como o atacante está a usar um endereço de origem falso, nunca obtém BBBB. Sem o valor correto de BBBB o atacante não consegue completar o protocolo de estabelecimento da ligação TCP.

Um ataque DoS realizado desta forma, sem responder ao SYN/ACK, designa-se “SYN flood”, mas tem pouco impacto.

Ataque aos números de sequência TCP

Infelizmente a maioria das implementações TCP não usa números de sequência iniciais completamente aleatórios.

Trata-se de números de 32 bits, a RFC793 (“Transmission Control Protocol”) indica um algoritmo de geração destes números com incrementos de uma unidade em cada 4 microsegundos.

A maioria das implementações não segue este algoritmo, normalmente os números de sequência iniciais que cada nó utiliza são incrementados em 128000 unidades em cada segundo e adicionalmente são incrementados 64000 unidades para cada nova ligação.

Esta simplificação facilita a tarefa de “adivinhação” do atacante. Por observação (“sniffing”) ou por estabelecimento de uma ligação com endereço IP de origem legítimo, o atacante consegue saber qual o último número usado e assim prever o valor que lhe será enviado seguidamente pelo servidor quando usar o endereço IP de origem falsificado.

Outros ataques

Existe uma grande variedade de ataque, muitos dos quais tiram partido de implementações dos protocolos pouco sólidas:

- Ataques de “loopback” em que são enviados pedidos com endereços de origem iguais aos de destino, por exemplo o “TCP Loopback DoS Attack” (land.c).
- Envio de informação incoerente para provocar o “crash” do servidor, como por exemplo o “ping of death” (POD) ou o “INVITE of Death” que afecta algumas implementações VoIP.
- Envio de um volume de informação superior ao que é suportado por uma implementação não protegida contra o “overflow”.

“ICMP Redirect” – estas mensagens ICMP são enviadas pelos “routers” aos nós da rede quando pretendem que estes utilizem um outro “router” para encaminhamento da informação, normalmente porque esse “router” alternativo corresponde a um caminho mais direto.

Um atacante pode enviar estas mensagens para fazer com que um nó utilize um “router” do atacante, ficando assim em situação de levar a cabo um ataque MITM. Para evitar esta situação os nós ignoram este tipo de mensagem.

Deteção de intrusos e autenticação

Os “firewall” são fundamentais para evitar muitos tipos de ataque como DoS e “Spoofing IP”, mas são apenas uma parte da solução.

A deteção de intrusos, através da “monitorização da rede” permite assinalar situações de ataque, nomeadamente usurpação de identidade e MITM.

Como parte da deteção de intrusos pode ser colocado na rede um nó destinado a atrair os potenciais atacantes, normalmente designado “honeypot”. Trata-se de uma armadilha para atacantes, não está especialmente protegido, mas é intensivamente monitorizado.

Além do controlo de acesso (“firewalls” estáticos) e da deteção de intrusos (“firewalls” dinâmicos e outro “software” especializado), a terceira vertente fundamental da segurança de rede é a utilização de protocolos de comunicação seguros com autenticação sólida.

Mesmo que a rede esteja a ser vítima de um ataque grave como MITM, se forem usados protocolos seguros com autenticação, o ataque acaba por ter pouco efeito.

Referências bibliográficas

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

(...)