

UNIVERSIDAD MICHOACANA DE SAN NICOLAS DE HIDALGO



FACULTAD DE CONTADURÍA Y CIENCIAS
ADMINISTRATIVAS



TESIS

SISTEMA DE ATAQUE DE DICCIONARIO PARA CLAVES
WPA

PARA OBTENER EL TITULO DE:
LICENCIADO EN INFORMATICA ADMINISTRATIVA

PRESENTA: JOSÉ CLAUDIO CHÁVEZ GONZÁLEZ

ASESOR: DOCTOR EN CIENCIAS PEDRO CHÁVEZ LUGO

MORELIA MICHOACÁN, MARZO DE 2015

Índice	
Resumen.....	1
Abstract.....	1
Capítulo 1 Planteamiento del Problema.....	2
1.1 Objetivos Generales.....	4
1.2 Preguntas de Investigación.....	4
1.3 Hipótesis.....	4
1.4 Justificación.....	5
Capítulo 2 Marco Teórico.....	6
2.1 Redes de Comunicaciones.....	7
2.1.1 Componentes de la red.....	7
2.1.2 Topologías de la red.....	14
2.1.3 Tecnologías de la red de área local.....	20
2.2. Tecnologías para red de área extensa.....	41
2.2.1 Redes inalámbricas.....	51
2.3 Protocolos para redes inalámbricas.....	54
2.3.1 WEP.....	54
2.3.2 WPA.....	56
Capítulo 3 Seguridad en Redes.....	58
3 Seguridad En Redes.....	59
3.1 Mecanismos de Seguridad.....	61
3.2 Tipos de Ataques.....	62
3.3 Hechos Relevantantes para la Seguridad en Sistemas.....	64
3.3.1 Los Años 60's.....	65
3.3.2 Los Años 70's.....	65
3.3.3 Los Años 80's.....	66
3.3.4 Los Años 90's.....	67
3.3.5 La Seguridad en la Actualidad.....	67
3.4 Autenticación.....	70
3.5 Mecanismo de Auditoría.....	72
3.6 Mecanismo de Control de Acceso.....	73
Capítulo 4 Herramientas de Seguridad.....	76
4.1 Scanners.....	77
4.2 Sniffers.....	82
4.3 Rompedores de contraseñas.....	83
4.4 Ataques a red.....	85
Capítulo 5 Desarrollo de la Herramienta.....	95

5.1 Introducción.....	96
5.2 Requerimientos de la Herramienta.....	97
5.3 Comandos Utilizados para Redes Inalámbricas.....	98
5.4 Diccionario de Claves o Contraseñas.....	99
5.5 Operación de la Herramienta.....	99
5.6 Diagrama de Flujo de Operaciones.....	100
5.7 Interface Gráfica de Usuario.....	101
Capítulo 6 Pruebas, Resultados y Recomendaciones.....	105
6.1 Introducción.....	106
6.2 Pruebas y Resultados.....	106
6.3 Recomendaciones.....	107
Capítulo 7 Conclusiones.....	109
Bibliografía.....	111
Apéndice A. Código Fuente de la Herramienta desarrollada en Java	114

Resumen

Este trabajo presenta el desarrollo de una herramienta de software basada en Java y en ambiente Linux para determinar si los usuarios emplean en los puntos de acceso a internet del hogar, trabajo, empresas, etc. contraseñas débiles que se pueden encontrar en un diccionario de contraseñas. Adicionalmente se pretende ofrecer una metodología de generación de contraseñas robustas, para que los usuarios puedan en un momento dado fortalecer sus contraseñas para el acceso a sus recursos de hardware y software.

Abstract

This document presents the development of a software tool based on a Java and on a Linux environment to determine if the users employs at access internet points from home, job, companies, etc. weak passwords which could be found in a passwords dictionaries. In addition, it pretends to offer a methodology of hard password creation in order for the users could get their passwords stronger, so they could access to their resources of hardware and software.

Capítulo 1

Planteamiento del Problema

1. Planteamiento del Problema

En los últimos años, alrededor de todo el mundo el tema de las redes inalámbricas ha salido a relucir, siendo esta un gran logro, ya que la información viaja a todos lados y llega de manera continua a las personas brindando una mayor comodidad y ayuda en las actividades y trabajos diarios de las personas.

La necesidad de comunicarse sin la ayuda de cables ha definido el nacimiento de la tecnología inalámbrica, este sistema utiliza la tecnología de radio frecuencia la cual utiliza como medio de transmisión el ambiente para transmitir y recibir datos. Todo esto permite la combinación de conectividad y movilidad para que los datos lleguen a varios metros o kilómetros de distancia.

Wifi es la tecnología inalámbrica más usada para las redes de área local ya sea de tipo personal o de trabajo. Una de sus principales ventajas es la movilidad, eliminando el tradicional cable Ethernet, permitiendo también conectividad en distancias considerables de metros a kilómetros.

La tecnología inalámbrica genera nuevos problemas de seguridad ya que las señales se encuentran en el ambiente y están disponibles tanto para los usuarios autorizados como para los no autorizados. Esto implica utilizar protocolos de autenticación que utilicen contraseñas robustas. Lamentablemente los estudios sobre las contraseñas utilizadas por las personas reflejan que un alto porcentaje de estos utilizan contraseñas débiles.

1.1 Objetivo General

Elaborar de una herramienta para determinar si una contraseña utilizada para proteger los puntos de acceso inalámbricos (wifi) es débil y se encuentra en un diccionario de contraseñas. Esta herramienta será desarrollada empleando el sistema operativo Linux con el lenguaje Java y el entorno de desarrollo Netbeans.

Objetivos Específicos

1. Probar la herramienta desarrollada en un entorno real.
2. Proponer a los usuarios una metodología para la generación de contraseñas robustas.

1.2 Pregunta De Investigación

Que tan común es el uso de contraseñas débiles en los puntos de acceso wifi con WPA (Wi-Fi Protected Access)

1.3 Hipótesis

La hipótesis de este trabajo se basa en los estudios [Leakedmail] [Realpass2006], que indican la alta probabilidad de que las personas utilicen contraseñas débiles para el acceso a sus recursos de hardware y software. Para esto se pretende probar estas hipótesis en el entorno local de la ciudad de Morelia, Michoacán México.

1.4 Justificación

El ser humano tiene una alta dependencia a los sistemas informáticos y a las redes de comunicación. Es vital importancia que el acceso a los recursos de hardware y software esté garantizados mediante una contraseña robusta. Ya que los usuarios en muchas ocasiones el acceso a sus recursos es sin contraseña y otras muchas veces más es mediante contraseñas débiles que hacen referencia a sus datos personales, palabras comunes de su mismo idioma u de algún idioma extranjero. Los datos sensibles que le pueden pertenecer a una persona u organización deben estar protegidos de la manera correcta y el empleo de contraseñas robustas es una instancia importante. De manera adicional es importante que los puntos de acceso de los hogares, trabajo, etc. tengan asociado una contraseña robusta para impedir que personas no autorizadas hagan un mal uso del servicio de internet. Lo cual puede implicar problemas legales, debido a posibles instalaciones de actividades no permitidas en equipos vulnerados.

En nuestro país existen innumerables tipos de redes inalámbricas las cuales permiten acelerar el proceso de información ofreciendo una mayor comodidad a los usuarios. La mayoría de los centros comerciales, edificios, restaurantes, aeropuertos, hoteles, cuentan con esta tecnología lo cual además de brindar a los usuarios acceso a Internet o alguna red de datos privada. Es por ello que se ha decidido realizar una herramienta para determinar si la contraseña es débil y se encuentra en un diccionario de contraseñas.

Capítulo 2 Marco Teórico

2.1 Redes de Comunicaciones

Entre todos los elementos esenciales para la existencia humana, la necesidad de interactuar está por debajo de la necesidad de sustentar la vida. La comunicación es casi tan importante para nosotros como el aire, el agua, los alimentos y un lugar para vivir.

Los métodos que utilizamos para compartir ideas e información están en constante cambio y evolución. Mientras la red humana estuvo limitada a conversaciones cara a cara, el avance de los medios ha ampliado el alcance de nuestras comunicaciones. Desde la prensa escrita hasta la televisión, cada nuevo desarrollo ha mejorado la comunicación.

Al igual que con cada avance en la tecnología de comunicación, la creación e interconexión de redes de datos sólidas tiene un profundo efecto.

Las primeras redes de datos estaban limitadas a intercambiar información basada en caracteres entre sistemas informáticos conectados. Las redes actuales evolucionaron para agregarle voz, flujos de video, texto y gráficos, a los diferentes tipos de dispositivos. Las formas de comunicación anteriormente individuales y diferentes se unieron en una plataforma común. Esta plataforma proporciona accesos a una amplia variedad de métodos de comunicación alternativos y nuevos que permiten a las personas interactuar directamente con otras en forma casi instantánea.

2.1.1 Componentes de la Red

Según [REDCOMTITTEL] la mayoría de las redes modernas se crean conectando diversos dispositivos físicos con el fin de establecer una ruta desde el dispositivo emisor hasta el receptor. El nivel del Modelo OSI en el que operan permiten generalmente clasificar estos dispositivos, vamos a concentrarnos en los más comunes: cables, concentradores, puentes, conmutadores y encaminadores.

Nivel del modelo OSI	Componente de Red
Nivel 3: Red	Encaminadores
Nivel 2: Enlace de datos	Puentes y conmutadores
Nivel 1: Físico	Concentradores y cables

Tabla 2.1 Asignación de los componentes de red al modelo OSI.

Cables

Según [REDCOMTITTEL], los cables transportan de un sitio a otro los datos. Podría ser de un extremo de una habitación a otro o de un extremo a otro de un país. La longitud máxima de un cable es un criterio de diseño importante y esta generalmente limitada por un factor denominado atenuación. La atenuación es una medida de la intensidad de la señal a medida que viaja a lo largo de un segmento de cable cada vez más largo; cuanto más largo sea el cable, mayor será la atenuación.

El apantallamiento es otro criterio de diseño importante. Algunos cables están apantallados para prevenir las interferencias externas, como por ejemplo las causadas por los motores o las luces fluorescentes, de modo que estas interferencias no puedan modificar la señal a medida que viaja por el cable.

El medio físico que forma el cable también es importante. La mayoría de los cables son de cobre, que transporta una señal eléctrica, o de fibra óptica, que permite transmitir un haz luminoso. Los cables basados en cobre suelen ser más resistentes, baratos y fáciles de usar, mientras que los cables de fibra óptica pueden alcanzar distancias mucho mayores y admiten frecuencias mucho más altas, lo que hace que tengan un ancho de banda mayor que los cables de cobre.

Los cables utilizados en la mayoría de las redes informáticas están normalizados por alguno de los siguientes organismos:

- American National Standards Institute (ANSI)

- Electronic Industry Association (EIA)
- Telecommunications Industry Association (TIA)

Nota: Los dos últimos organismos suelen considerarse conjuntamente, como EIA-TIA

Los cables también pueden clasificarse según la normativa aplicable al sector de la construcción. Generalmente, lo que determina estas clasificaciones son aspectos tales como la cantidad de humo o las llamas emitidas en caso de incendio. Estas clasificaciones incluyen:

- Cable restringido: debe de estar encerrado en un conducto.
- Cable de propósito general: para conexiones generales.
- Cable elevador: para conectar unas plantas con otras.
- Cable pleno: para su utilización es falsos techos y conductos de aire.

El tipo de cable de cobre más común es de par trenzado no apantallado (UTP, unshielded twisted-pair). Este cable está normalizado en las categorías 1 a 6, estas categorías suelen abreviarse como CAT1, CAT2, etc. Las categorías más comunes son:

- CAT2 Comúnmente utilizado para cableados telefónicos dentro de un edificio y especificado para una frecuencia máxima de 1Mhz. Este cable suele terminarse mediante un conector RJ-11. Este es el tipo de conector utilizado para los teléfonos domésticos.
- CAT3 Comúnmente utilizado para redes Ethernet 10Base-T y especificado para una frecuencia máxima de 16 MHz. El cable CAT3 suele terminarse mediante conectores RJ-45, que son similares a los conectores RJ-11, pero tienen ocho hilos en lugar de cuatro hilos.
- CAT5 La especificación mínima para Fast Ethernet 100Base-T. también admite 10Base-T, Token Ring y aparatos telefónicos. Debido a esto, es el tipo más común de cableado utilizado en las redes modernas. CAT5 utiliza también conectores RJ-45.

Otro tipo de cable de cobre es el cable coaxial, utilizado en los sistemas de televisión por cable. Está construido mediante anillos concéntricos de material conductor, separados por una capa aislante de algún tipo, en lugar de trenzar pares de hilos individuales. Normalmente, los cables coaxiales se terminan mediante un conector BNC.

El cable de fibra óptica también se suministra en diversas variantes. Los dos tipos más comunes son los cables de monomodo (SM, Single Mode) y multimodo (MM, Multimode). A diferencia de los cables de cobre, la especificación de los cables de fibra óptica incluye el diámetro del núcleo, que es a través de donde viaja la luz, y el diámetro de recubrimiento. Estos valores para las fibras multimodo suelen ser de 62,5 y 125 μm , respectivamente, en Estados Unidos, mientras que en Europa se utiliza a menudo un núcleo de 50 μm . La fibra monomodo tiene un núcleo mucho más pequeño, normalmente entre 5 y 10 μm . la especificación de los cables de fibra óptica suelen también indicar la longitud de onda permitida. Esta es, normalmente, de 850 nm(xxxxxxx) o 1350 nm. Observe que la diferencia principal entre los cables monomodo y multimodo es que para la fibra monomodo se utiliza un dispositivo laser, lo que hace que sea un sistema muy caro; por el contrario, para las fibras multimodo se utiliza un diodo electroluminiscente (LED, Light-Emitting Diode). Esto hace que las fibras monomodo admitan distancias mucho mayores (normalmente superiores a los 25 km.), mientras que la fibra multimodo se utiliza casi exclusivamente dentro de un mismo edificio o en redes universitarias de pequeño tamaño.

Otra diferencia entre la fibra óptica y el cable de cobre es que este suele tener vario pares de hilos; la fibra óptica solo tiene un único par (un hilo para transmitir y otro para recibir).

Los cables de fibra óptica suelen terminarse mediante conectores SC o ST. Estos conectores son similares, pero los conectores ST son redondos y los conectores SC, cuadrados. La mayoría de los sistemas de comunicaciones más antiguos, como las redes telefónicas y las redes Token Ring, utilizaban conectores ST, mientras que los sistemas más modernos como Ethernet, emplean conectores ST. Una desventaja de estos conectores es que cada fibra del par necesita su propio conector, por lo que para terminar un circuito (dos hilos) se requiere un espacio considerable, por comparación con los conectores típico para cables de cobre. Esto limita el número de conexiones que pueden terminarse con facilidad en un dispositivo.

Recientemente se ha desarrollado el conector MTRJ para solventar esta diferencia. Dicho conector permite terminar ambas fibras y es mucho más pequeño, aunque con la desventaja de ser algo frágil.

Otros aspectos interesantes son los siguientes:

- Todos los cables tienen un radio mínimo de curvatura, porque la señal se ve afectada al doblar el cable. Los cables de fibra óptica son mucho más sensibles que los de cobre a los efectos derivados de doblar el cable.
- Los cables de fibra óptica no son susceptibles a la interferencia electromagnética.

Concentradores

De acuerdo con [REDCOMTITTEL], para conectar varias computadoras dentro de un mismo edificio, se suelen tender cables desde el PC en la mesa de cada usuario hasta un armario de cableado. Allí, hace falta un dispositivo para conectar entre si los cables. Dicho dispositivo es normalmente un centrador. Los concentradores son dispositivos que proporcionan una ruta física para que una señal viaje de un cable a otro. Aunque se comportamiento está especificado por una determinada tecnología, como por ejemplo Ethernet o Token Ring, las cuales se consideran generalmente parte nivel 2 del modelo OSI, se considera que los concentradores operan en el nivel 1, que es el nivel físico. Esto se debe a que los concentradores actúan como repetidores multipuerto; en otras palabras, se limita a regenerar una señal eléctrica recibida en un puerto, retransmitiéndola a través de uno o más puertos diferentes, sin introducir ningún cambio.

Puesto que un concentrador simplemente repite la señal sin modificar la información, cada puerto del concentrador forma parte del mismo enlace de datos o segmento de red. Esto significa que, en una red Ethernet, todos los puertos de un concentrador forman parte del mismo dominio de colisiones. Esto quiere decir que, para el concentrador en su conjunto, solo puede haber una computadora enviando datos en cada momento. En una red Token Ring, todos los puertos de un concentrador forman parte del mismo anillo.

Puentes

Conforme a [REDCOMTITTEL], cuanto el número de usuarios crece y se empieza a ver cuestionados los límites de un único segmento de red, surge la necesidad de crear un nuevo segmento para enlazar entre si dos redes. Los dispositivos llamados puentes permiten conseguir precisamente esto. Originalmente, los puentes tenían dos puertos, uno para cada una de las dos redes que había que conectar. Sin embargo, a diferencia de los concentradores, los puentes sí que inspeccionan los datos que pasan a su través y toman decisiones sobre si deben enviarse a la otra red o no. Estas decisiones se basan en la dirección MAC en las redes Ethernet y en el número de anillo en las redes Token Ring. Debido a este comportamiento, se dice que los puentes son dispositivos de nivel 2.

Los puentes Ethernet analizan el tráfico enviado por las computadoras y otros dispositivos de red y registran la dirección MAC de la computadora, que está ubicada en el campo de la dirección de origen (Source Address) de la cabecera de trama Ethernet, y el puerto en el que la dirección fue detectada. Si el puente recibe entonces una trama procedente de la otra red que este destinada para la dirección MAC registrada para la primera red, enviara dicha trama a la primera red.

Los puentes Token Ring operan basándose en el número de anillo. A cada puente se le asigna un número de puente y un número de anillo. Las tramas Token Ring contienen un campo de información de encaminamiento (RIF, Routing Information Field), que es una lista de los números de anillo y números de puente que la trama debe atravesar para llegar a su destino. Cuando un puente Token Ring ve una trama en un anillo que está destinada para otro al cual el puente también está conectado, retransmitirá la trama hacia ese segundo anillo.

Conmutadores

Según [REDCOMTITTEL], a medida que las redes fueron creciendo todavía más y la cantidad de datos transmitidos por cada computadora se fue incrementando, se hizo toda vía más importante segmentar las redes. Los puentes de dos puertos ya no eran suficientes. Los

conmutadores iniciaron su andadura como puentes multipuerto y se consideran dispositivos nivel 2. La mayoría de los conmutadores tienen 12 o 24 puertos, pero muchos son moduladores y pueden tener varios centenares de puertos.

Otra distinción es que los conmutadores pueden gestionar varias conversaciones al mismo tiempo. Cada puerto 100Base-TX de un conmutador puede enviar y recibir tramas al mismo tiempo (lo que se le denomina comunicación dúplex). Lo que quiere decir que los conmutadores tienen una tarjeta de control bastante compleja que permitía que a cada puerto hable con cualquiera de los otros puertos. Aunque los detalles de funcionamiento de estos mecanismos suelen publicarse, no forman parte de ningún estándar tecnológico, sino que son propietarios de casa tipo de conmutador.

Los conmutadores emplean también búferes. Estos búferes son zonas de memoria que se utilizan para almacenar las tramas hasta que sean transmitidas. Esto resulta útil cuando hay muchos dispositivos conversando con un mismo dispositivo y todos ellos envían, simultáneamente, más datos de los que el enlace puede admitir en un momento dado. En este caso, las tramas esperan en el búfer hasta que esté disponible el suficiente ancho de banda en el enlace para transmitir la trama. Se utilizan comúnmente dos tipos de conmutadores:

- Conmutadores de almacenamiento y reenvío, que reciben la trama completa en el búfer antes de transmitirla. Esto permite al conmutador leer y comprobar las sumas de comprobación al final de las tramas, para garantizar que estas no se hayan corrompido.
- Conmutadores de anticipación, solo leen el campo de dirección destino de la cabecera de nivel 2 antes de comenzar a transmitir. Esta clase de conmutadores pueden retransmitir tramas erróneas y fragmentos de tramas, pero son más rápidos que los conmutadores de almacenamiento y reenvío.

Encaminadores

De acuerdo con [REDCOMTITTEL], aunque segregar los segmentos de red resulta útil, todos los dispositivos conectados a los concentradores, puentes y conmutadores continúan

estando en el mismo dominio de difusión, y hay una serie de límites prácticos al número de dispositivos que pueden coexistir en cualquier dominio de difusión. Por esto, para poder segregar los dominios de difusión, se crearon los encaminadores. Los encaminadores actúan como la frontera entre dominios de difusión. De forma similar a la manera en que los puentes actúan como la frontera entre dominios de difusión. Los encaminadores leen y toman decisiones basándose en las cabeceras de nivel 3, como por ejemplo las cabeceras TCP/IP o IPX. Por lo tanto, decimos que los encaminadores son dispositivos de nivel 3.

El trabajo de un encaminador consiste en inspeccionar cada paquete que se le envía y determinar si pertenece a la red IP o IPX local o a una red remota. Si el destino del paquete es una red remota y el encaminador conoce como llegar hasta esa red, el encaminador reenvía el paquete; en caso contrario el paquete se descarta.

Los encaminadores suelen utilizarse casi exclusivamente para conectar redes remotas a través de enlaces WAN, pero esto no guarda ninguna relación con la propia función de encaminamiento. Resulta posible utilizar puentes u otros dispositivos, como por ejemplo un PC, para conectar enlaces WAN.

2.1.2 Topologías de la Red

Como dice [REDCOMTITTEL] una determinada topología puede describir varios protocolos de red. Por ejemplo, FDDI como Token Ring son anillos. Los tipos de rutas se suelen dividir en varias categorías.

En el contexto físico, tiene su significado normal no hace referencia al nivel físico del modelo OSI. Los atributos físicos de una red describen cosas palpables, mientras que el nivel Físico describe el comportamiento de los electrones y ondas luminosas. Ambas acepciones se solapan en cierta manera en el área de las especificaciones de cables y conectores.

Por comparación, el término lógico en el contexto de las topologías de red describe el comportamiento especificado en el nivel 2 del modelo OSI.

Esta distinción es importante dado que una misma tecnología como Ethernet o Token Ring, pueden tener una topología física y otra topología lógica diferente.

Medio Compartido

Una topología es capaz de permitir que más de dos dispositivos compartan la red en cada momento [REDCOMTITTEL]. A diferencia de otros tipos de medios, los medios compartidos es posible que tengan lugar múltiples conversaciones simultáneamente. Esto significa que el medio debe disponer de algún método para controlar el acceso.

Ethernet es un medio compartido y su protocolo de nivel MAC es CSMA/CD. Por contraste, los dispositivos Token Ring solo pueden transmitir una trama cada vez. Después de eso, debe de entregar el testigo de transmisión al siguiente al siguiente dispositivo de la cadena. Incluso aunque los demás dispositivos no tengan tramas que transmitir, el testigo debe de dar la vuelta completa al anillo antes de que la primera estación pueda transmitir la segunda trama.

El método que Ethernet utiliza para controlar el acceso es la detección de colisiones (CD, collision detection). En pocas palabras CD significa que cada dispositivo está a la escucha para ver si otros dispositivos de la red están transmitiendo. Si no hay nadie más transmitiendo, el dispositivo puede transmitir el máximo de información posible.

Los medios compartidos también deben resolver el problema de direccionamiento. Cuando solo hay dos dispositivos en la red, resulta obvio en otras topologías que el otro dispositivo será siempre el receptor del tráfico que nosotros enviemos y viceversa. Una vez que se añade un tercer dispositivo, se necesita una manera de identificar cual es el dispositivo con el que queremos comunicarnos.

Aunque las implementaciones varían, casi todas las tecnologías de red modernas utilizan algún tipo de direccionamiento MAC. Hay muchos tipos de medios compartidos, pero los dos más comunes son la topología en bus y la topología en anillo.

Bus

Una topología de bus (ver Figura 2.1) es similar a la arquitectura de bus que conecta la memoria principal con el procesador y las unidades de disco de una computadora [REDCOMTITTEL]. Se trata de una ruta de datos simple en la que todos los dispositivos de la red están conectados a la misma ruta de comunicaciones, de manera que solo uno de los dispositivos puede usarla en cada momento. Dicha ruta puede ser física o lógica.

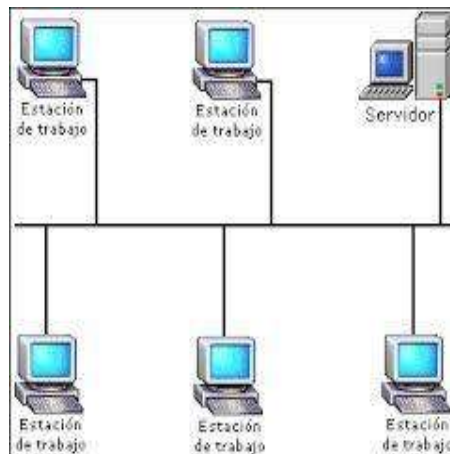


Figura 2.1 Topología en bus.

Desde el punto de vista lógico, Ethernet es siempre un bus. Sin embargo, desde el punto de vista físico, Ethernet puede ser un bus, en el que cada computadora está conectada a la siguiente computadora y se incluyen sendas terminaciones en los extremos del bus; pero Ethernet también puede ser físicamente una topología en estrella, en la que cada dispositivo tiene su propio cable que está conectado a un concentrador central.

Anillo

Desde el punto de vista lógico, una topología en anillo (ver Figura 2.2) es aquella en la que cada dispositivo transmite solo hacia su vecino situado aguas abajo y recibe solo desde su vecino situado aguas arriba [REDCOMTITTEL]. En otras palabras, si se quiere recibir una trama del vecino más próximo aguas abajo, la trama tendrá que recorrer todo el anillo, pasando por cada uno de los demás dispositivos, antes de recibirla.

Desde el punto de vista físico, una topología en anillo representa el mismo concepto. Cada dispositivo está conectado exactamente a otros dos (a menos que solo haya dos dispositivos en la red), de modo que todos los dispositivos forman un círculo. Otro ejemplo de la diferencia entre disposición física y disposición lógica es que Token Ring es un anillo desde el punto de vista lógico, mientras que físicamente es una estrella, al igual que una red Ethernet de par trenzado, en la que cada dispositivo está conectado a un concentrador central.

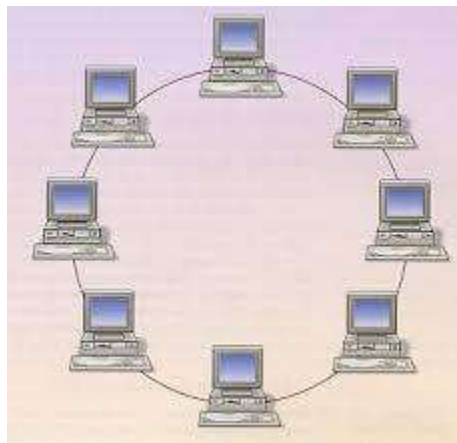


Figura 2.2 Red en anillo.

Redes Igualitarias

Una red igualitaria o peer to peer network (ver Figura 2.3), que a menudo se designa mediante la abreviatura P2P, es el tipo más simple de red cuando existen muy pocos nodos, pero si son muchos los dispositivos que necesitan comunicarse con más de un dispositivo, este tipo de redes deja muy pronto de ser manejable [REDCOMTITTEL]. Estas redes se

construyen conectando un único enlace entre dos dispositivos homólogos. Estos dispositivos homólogos pueden ser dos computadoras con un cable serie o dos encaminadores conectados mediante un circuito TI punto a punto. El aspecto más importante que hay que tener en cuenta es que solo hay dos dispositivos y una conexión. Aunque algunas veces los dispositivos utilizan direcciones MAC, resulta en realidad innecesario.

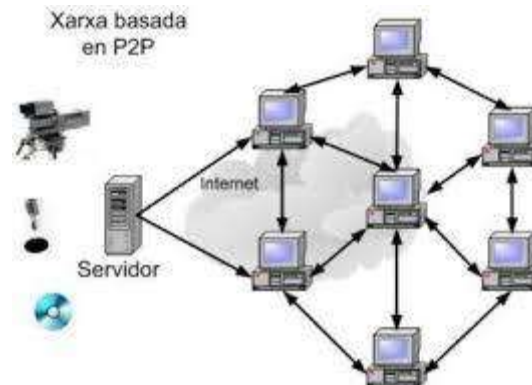


Figura 2.3 Red P2P

Otro punto importante es que el término de red igualitaria en los medios de comunicación hace usualmente referencia a la relación entre dos computadoras por encima del nivel de aplicación: cliente-servidor, servidor-servidor, o cliente-cliente, que es otra forma de hacer referencia a las redes formadas entre los dispositivos homólogos. Un ejemplo de redes P2P son los servicios de intercambio de archivos musicales, como por ejemplo Ares o Gnutella. Es importante mencionar que ese tipo de redes no tiene nada que ver con el concepto de red igualitaria en el contexto de las topologías de red.

Cuando hay muchos dispositivos involucrados y varios de ellos necesitan comunicarse en numerosos dispositivos involucrados, suelen combinar varias redes igualitarias para formar una de las siguientes configuraciones:

Estrella

En una configuración en estrella (ver figura 2.4), hay un dispositivo central que tiene conexiones con todos los demás dispositivos y se encarga de transmitir las comunicaciones

entre ellos [REDCOMTITTEL]. La configuración en estrella también se denomina a veces configuración de rueda, porque se asemeja a la rueda de un carromato.

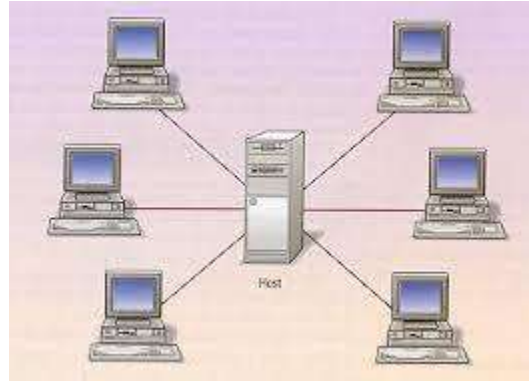


Figura 2.4 Red estrella

Malla

Para conectar todos los dispositivos entre sí, hacen falta $n(n-1)/2$ conexiones. Si solo hay 10 dispositivos en la red, serían necesarios 45 enlaces físicos y cada dispositivo necesitaría nueve interfaces para conectarse con los demás [REDCOMTITTEL]. La ventaja de este tipo de topología es que los datos nunca tienen que realizar más de un salto a través de la red; como consecuencia, esta red puede ser increíblemente rápida.

Una solución de compromiso bastante común es la que se denomina malla parcial (ver figura 2.5). En una malla parcial, lo que se hace es eliminar algunos de los enlaces de la malla. Es la realidad, los patrones de tráfico en cualquier red indican que cada dispositivo de la misma invierte la mayor parte del tiempo hablando únicamente con algunos otros dispositivos determinados, si eliminamos los enlaces raramente utilizados, se reduce enormemente el coste y la complejidad de la red.

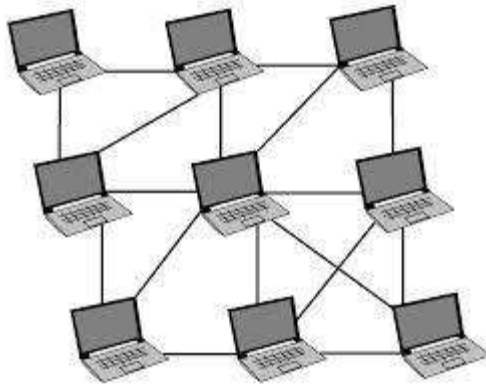


Figura 2.5 Red malla

Redes Híbridas

Por mucho empeño que pongamos en definir correctamente las reglas, siempre aparecerá alguien que pretenda romperlas [REDCOMTITTEL]. No resulta sorprendente, por tanto, que diversas compañías traten de tomar las mejores características de cada topología y combinarlas. En esta topología, cada bus puede hasta ocho dispositivos (ver figura 2.6), pudiéndose conectar varios buses entre sí en una configuración en estrella.

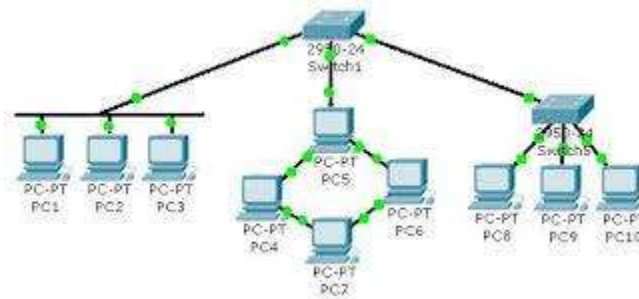


Figura 2.6 Buses conectados en una configuración en estrella

2.1.3 Tecnologías de Red de Área Local

De acuerdo con [REDCOMTITTEL], explica las tecnologías LAN en orden de importancia en las redes empresariales de hoy en día. La popularidad de algunas de estas tecnologías se

está incrementando, mientras que las de otras se desvanece, pero el aspecto más importante que hay que resaltar es que las especificaciones representan una determinada instantánea temporal. En la realidad, los fabricantes de equipos de redes suelen lanzar al mercado equipos funcionales antes de publicarse las especificaciones, para ser los primeros en vender. Y, por supuesto, estas tecnologías se están constantemente mejorando, especialmente el nivel Físico, que es donde se concentra la mayor parte de las investigaciones, por lo que las especificaciones solo pretenden ser lo suficientemente restrictivas como para garantizar la compatibilidad, sin impedir la continua mejora de las implementaciones hardware y software de dichas tecnologías.

Tecnologías Ethernet

Ethernet es, con mucho, la tecnología LAN más popular de hoy en día, y continuara siéndolo durante el futuro próximo. Su popularidad tiene más que ver con el coste de fabricación por cada puerto que con las ventajas prácticas. De hecho, hay otras tecnologías más rápidas y seguras que permiten comunicarse a distancias mucho más grandes, pero ninguna de ellas es más barata que Ethernet. En cualquier caso Ethernet tiene una historia bastante notable. Ethernet fue creada originalmente por Bob Metcalfe en 1976, en los laboratorios PARC (Palo Alto Research Center) de Xerox. Fue diseñada para conectar un PC a una impresora láser.

Tipos de Tramas de Ethernet

Debido a su larga historia, hay muchas versiones de las tramas y muchas especificaciones para las implementaciones del nivel Físico. Si se está utilizando Ethernet en un entorno empresarial actual, existen bastantes probabilidades de que se tope con dos o más de los cuatro tipos de trama más comunes. Estos tipos pueden llegar a ser algo confusos, porque las diferentes organizaciones los designan mediante nombres diferentes.

Los cuatro formatos de trama son bastante sencillos. Todos ellos proporcionan las siguientes informaciones:

- Un campo de dirección de destino.
- Un campo de dirección de origen.
- Un mecanismo para identificar el contenido de la carga útil.
- Un campo de carga útil, que trasporta los datos (por ejemplo, un paquete TCP/IP).
- Una suma de comprobación.

Tipo de trama	Cisco	Novell	Notas
Versión II IEEE 802.3 IEEE 802.3 SNAP Formato de Novell	ARPA LCC SNAP NOVELL	Ethernet II Ethernet 802.2 Ethernet SNAP Ethernet 802.3	A menudo denominada DIX Incluye la cabecera LLC 802.2 Utilizada para la compatibilidad Propietaria

Tabla 2.2 Tipos de tramas Ethernet

Ethernet Versión II

La primera versión de Ethernet ha sido completamente sustituida por la Versión II y ya no se emplea. La Versión II es la primera especificación que goza de una amplia aceptación. Comúnmente se le denomina DIX, un acrónimo formado a partir de las iniciales de las tres empresas que respaldaron el estándar Ethernet: DEC (Digital Equipment Corporation), Intel y Xerox. La versión II fue especificada por el consorcio DIX y utiliza el formato de trama en la siguiente Tabla.

6 bytes	6 bytes	2 bytes	Variable	4 bytes
Dirección de destino (unidifusión, difusión o multidifusión)	Dirección de Origen	Ethertype	Carga útil	FCS

Tabla 2.3 Formato de trama de Ethernet II.

Los campos de dirección de origen y destino contienen la dirección MAC de 6 bytes del transmisor y el receptor, respectivamente. Existen tres tipos de direcciones de destino:

- Unidifusión, que identifica un único nodo de la red.

- Difusión, que hace que la trama se envíe a todos los nodos de la red.
- Multidifusión, que hace que la trama se envíe a un grupo de nodos de la red.

Una dirección de destino de unidifusión comienza con el campo OUI (Organizationally Unique Identifier, identificador unívoco de organización), que ocupa los primeros tres bytes de la dirección MAC. Estos tres bytes son asignados por el IEEE para identificar de manera unívoca a un fabricante de hardware de red. El propio fabricante determina los últimos tres bytes de la dirección MAC. Este esquema permite garantizar que cada tarjeta de interfaz de red (NIC) del mundo tenga una dirección hardware unívoca.

Una dirección de destino de difusión está compuesta solo por unos binarios (11111111-11111111-11111111-11111111), lo que se muestra como FF-FF-FF-FF-FF-FF en hexadecimal. Un nodo de la red puede enviar un paquete de difusión cuando quiera compartir o solicitar información de todos los demás nodos de la red.

Una dirección de multidifusión comienza siempre con uno en el primer byte. Esto se debe a que el primer bit de la dirección MAC es el bit de Individuo/Grupo (I/G). Si el bit IG es un cero, la dirección representa a un adaptador individual. Si el bit IG es un 1 se trata de una dirección de grupo. La parte de OUI es siempre 0X01.00.5E. El resto de la dirección MAC indica la dirección IP de multidifusión. Desafortunadamente, solo hay 23 bits de espacio de dirección MAC, mientras que la dirección IP de multidifusión tiene 28 bits. Esto quiere decir que no se puede especificar la dirección multidifusión IP completa en una dirección MAC de multidifusión.

Una organización, IANA (Internet Assigned Numbers Authority, autoridad de asignación de números Internet), mantiene la lista de protocolos que se ejecutan por encima de Ethernet. Algunos valores comunes de ejemplo son 0X0800 para el tráfico IP, 0X0806 para el protocolo de resolución de direcciones (ARP) y 0X8137 para el protocolo IPX/SPX (Internetwork Packet Exchange/Sequenced Packet Exchange) de Novell. Cuando el software controlador Ethernet ve uno de estos valores en el campo Ethertype, sabe a qué proceso software o ubicación de memoria tiene que entregar el contenido de su carga útil. La carga útil es el paquete que está siendo transportado, que contiene la cabecera de nivel 3

y los datos. El adaptador Ethernet de origen calcula un polinomio complejo, denominado secuencia de control de trama (FCS, Frame Check Sequence), utilizado como entrada la cabecera Ethernet y la carga útil. Cuando se completa el proceso, se envía la trama a través del cable hasta su destino. El adaptador destino almacena todos los bits que recibe en su búfer de recepción. Cuando detecta que la tensión en la línea vuelve a cero o detecta una señal de inactividad, calcula el mismo polinomio complejo utilizando todos los bytes de búfer, salvo los últimos cuatro. Después compara el resultado con el contenido del búfer y, si los valores no son idénticos el adaptador asume que la trama esta corrompida y la descarta.

Ethernet IEEE 802.3 y Ethernet 802.2

El siguiente tipo de trama es IEEE 802.3

6 bytes	6 bytes	2 bytes	3 bytes	Variable	4 bytes
Dirección de destino	Dirección de origen	Longitud	Cabecera LLC	Carga útil	FCS

Tabla 2.4 Formato de trama Ethernet IEEE 802.3

La cabecera LLC (Logical Link Control, control de enlace lógico) que aparece después del campo Longitud está definida en la norma IEEE 802.2

Los campos de la cabecera IEEE 802.3 operan exactamente como los campos de la Versión II del mismo nombre, con dos excepciones. La primera es que el valor de Ethertype se ha sustituido por un campo de Longitud que contiene un valor igual a la longitud de la carga útil. En segundo lugar, se ha añadido una cabecera LLC adicional antes de la carga útil. El campo LLC opera igual que el campo Ethertype, salvo porque transporta información no solo sobre el protocolo de destino, sino también sobre el protocolo de origen. Esta función está definida en las normas sobre control del enlace lógico de IEEE 802.3.

1 byte	1 byte	1 byte
Punto de acceso de servicio de destino (DAP, Destination Service Access Point)	Punto de acceso de servicios de origen (SSAP, Source service Access Point)	Control

Tabla 2.5 Cabecera LLC de 3 bytes

Formato de Trama SNAP

El formato de trama SNAP es idéntico a la cabecera IEEE 802.3, salvo porque se incluye una cabecera SNAP

3 bytes	2 bytes
Código del fabricante	Código local

Tabla 2.6 formato de trama SNAP.

La cabecera SNAP contiene dos campos: un código de fabricante y un código local. Estos campos proporcionan compatibilidad descendente entre IEEE 802.3 y la Versión II. Para conseguir la compatibilidad, el código de fabricante suele hacerse igual a los primeros tres bytes de la dirección de origen, que es la parte OUI de la dirección MAC de la tarjeta adaptadora transmisora. El código local se hace igual al valor de Ethertype que habría sido utilizado si se trata de una trama Versión II.

El formato propietario de Novell para el tipo de trama SNAP es casi idéntico también al formato IEEE 802.3, salvo porque no incluye la cabecera LLC. Por esta razón, este formato se suele denominar <<en bruto>>. La otra diferencia es que comienza el campo de carga útil con una suma de control opcional, que esta desactivada de manera predetermina (es decir, se le asigna valor FF-FFh).

Versiones de Ethernet

No solo hay cuatro tipos de tramas, sino que también hay varias versiones diferentes de red Ethernet. Estas versiones se denominan normalmente, Ethernet 10 Mbps (la primera

versión de Ethernet), Fast Ethernet de 100 Mbps y Gigabit Ethernet de 1000 Mbps. Aunque todas ellas son <<Ethernet>>, difieren enormemente en el nivel Físico, porque utilizan diferentes esquemas de codificación.

Ethernet a 10 Mbps

Hay varias especificaciones de nivel Físico con Ethernet. Estas especificaciones comenzaron con el siguiente formato: en número representando la velocidad en Mbps, seguido de la palabra *Base*, para hacer referencia a <<banda base>>, o *Broad*, para hacer referencia a <<broadband (banda ancha)>> y luego un número, que originalmente representaba la máxima distancia en cientos de metros. Los cuatro tipos de Ethernet a 10 Mbps de los que se tratarán son: 10Base2, 10Base5, 10BaseT y 10Base-F.

- *10Base2*. Normalmente llamada Ethernet fina o thinnet, 10Base2 opera a 10 Mbps y tiene una distancia máxima de 200 metros. Esta especificación usa un cable coaxial fino y es, desde el punto de vista físico, una topología de bus, en el que cada nodo de la red Ethernet tiene una única conexión en algún punto del cable coaxial. La Ethernet fina especifica una resistencia de 50 ohmios para terminar cada extremo del cable coaxial, lo que evita que señales reboten (es decir, que se reflejen en el extremo del cable y vuelvan a introducirse en el segmento, pudiendo ser, como consecuencia, malinterpretadas como si fueran una señal diferente procedente de otro dispositivo). Los conectores utilizados para 10Base2 son denominados uniones en T tipo BNC.
- *10Base5*. El estándar 10Base5, llamado Ethernet gruesa o thicknet, también utiliza un cable coaxial similar, pero mucho más grueso. Este cable resulta menos adecuado para conectar computadoras en un entorno típico de oficina, porque el cable es mucho más rígido que el de la Ethernet fina. Puesto que esta especificación tiene una distancia máxima de 500 metros, normalmente se utilizaba este tipo de red como red troncal para interconectar armarios de cableado dentro de un edificio o, en ocasiones, para conectar unos edificios dentro de otros en un campus.

- *10BaseT*. Tanto 10Base2 como 10Base5 fueron abandonadas casi completamente a mediados de la década de 1990 en favor de la especificación 10BaseT. En este caso, la <<T>> quiere decir <<cable de par Trenzado>>. 10BaseT difiere sustancialmente de las otras dos redes Ethernet a 10 Mbps, en el sentido de que, desde el punto de vista físico, tiene una topología en estrella. En lugar de conectar unos nodos a otros en serie, cada nodo se conecta a un dispositivo central denominado <<concentrador>>. La especificación 10BaseT emplea normalmente cables de par trenzado no apantallado de tipo CAT3 o CAT5. Una de las principales ventajas de esta especificación es que puede enchufarse y desenchufarse un cable del concentrador sin interrumpir las comunicaciones de los otros dispositivos.
- *10Base-F*. El último estándar Ethernet que merece una mención es 10Base-F. En este caso, la <<F>> quiere decir <<Fibra Óptica>>, que es el medio físico utilizado en esta especificación. 10Base-F también utiliza una topología en estrella, pero en la práctica la mayoría de los concentradores contienen 8, 16 o 24 puertos 10Base-T y uno o dos puertos 10Base-F. esta combinación hace posible combinar varios segmentos de tender enlaces de muy larga distancia entre armarios de cableado o edificios.

Nota: aunque el término Banda Ancha se utiliza comúnmente de forma incorrecta para diferenciar las conexiones rápidas y lentas, el término no hace referencia en realidad a si están utilizando señales simultáneamente en un cable o, por el contrario, solo se utiliza una única señal. Los módems de cable son de banda ancha porque los datos y la televisión por cable pueden utilizar un mismo segmento de cable coaxial al mismo tiempo. De forma similar, las conexiones DSL (Digital Subscriber Loop, bucle digital de abonado) son de banda ancha porque se pueden enviar y recibir datos y utilizar el teléfono a través del mismo par de hilos de cobre simultáneamente. Aunque Ethernet es mucho más rápida que la conexiones por cable y de tipo DSL, es una tecnología de banda base, porque no puede compartir el medio físico con ninguna otra señal.

Fast Ethernet

De forma similar a la Ethernet a 10-Mbps, Fast Ethernet tiene varias especificaciones para diferentes medios físicos, pero además incluye unos cuantos modos opcionales de operación. El primero es 100BaseT, que es muy similar a 10BaseT. En este nivel Físico, sin embargo, la codificación es muy distinta. Esto se debe a que Fast Ethernet descompone el nivel Físico en tres subniveles diferentes:

- PCS (Physical Coding Sublayer, subnivel de codificación física).
- PMA (Physical Medium Attachment, conexión al medio físico).
- PMD (Physical Medium Dependent, dependiente del medio físico).

El principal desafío a la hora de crear redes de alta velocidad es que las señales de alta frecuencia no se propagan muy bien a través de muchos medios físicos. Todavía peor es el hecho de que Ethernet utiliza codificación Manchester para resolver los problemas de temporalización. La codificación Manchester funciona haciendo que la tensión cambie durante cada intervalo de bit.

Un cambio de una tensión superior a una inferior se considera un uno, mientras que un cambio de una tensión inferior a otra superior se interpreta como un cero binario. Estas transiciones regulares permiten al adaptador de recepción mantener sincronizado su reloj, incluso en el caso de que se envíen varios unos o ceros seguidos. El problema es que este tipo de mecanismo duplica en la práctica la frecuencia requerida. En otras palabras, una velocidad de datos de 10 Mbps requiere una forma de onda a 20 MHz.

Si 100Base-TX utilizara codificación Manchester, los 100 Mbps requerirían una frecuencia de casi 200 MHz. Para reducir los requisitos de frecuencia, la especificación 100Base-TX utiliza primero un mecanismo denominado MLT-3 (Multiple Level Transitions-3, transiciones multinivel con tres niveles) y un esquema de codificación NRZI (Non-Return-to-Zero, Invert-on-one; sin retorno a cero y con inversión a uno).

La codificación NRZI funciona manteniendo una de dos tensiones constantes (por ejemplo, una tensión positiva y una tensión negativa). Cuando se envía un bit de valor uno, el

transmisor hace cambiar la tensión. Cuando se envía un bit de valor cero, la tensión permanecer constante durante todo el intervalo del bit. MLT-3 simplemente añade otro nivel a este procedimiento, para poder utilizar una tensión positiva, una tensión cero y una tensión negativa. La ventaja de este sistema es que pueden representarse cuatro bits en un único ciclo de reloj. MLT-3 va recorriendo estos tres valores en orden. Por ejemplo, si se transmite un byte con valor 11111111, la tensión sería 1,0, -1, 0, 1, 0, -1, 0. Si se transmite un byte con un valor 11001100, este se representaría eléctricamente como 1,0, 0, 0, -1, 0, 0, 0.

Queda un problema por resolver con MLT-3 y NRZI: si se envía un gran número de ceros seguidos, no se producirá ningún cambio de tensión. Eso, a su vez, significa que el reloj de recepción podría desincronizarse. Para resolver esto, se utiliza la codificación 4B5B. Este esquema de codificación 4B5B simplemente crea una tabla de todos los posibles valores de un cuarteto (es decir, la mitad de un byte, cuatro bits) y luego asigna esos valores a un valor de código de cinco bits, cumpliéndose que todos los valores de código de cinco bits incluyen al menos dos binarios. Esto significa que ninguna combinación de los valores de datos permitirá que pasen más de unos cuantos intervalos de bit sin efectuar una transición de tensión.

Cuarteto de datos	Código de cinco bits utilizando para reemplazar en cuarteto
0000	11110
0001	01001
0010	10100
0011	10101
0100	01010
0101	01011
0110	01110
0111	01111
1000	10010
1001	10011
1010	10110
1011	10111
1100	11010
1101	11011
1110	11100
1111	11101

Tabla 2.7 Codificación 4B5B.

Hay tres valores de código importantes que tienen un significado especial:

- 11111 se utiliza cuando la línea esta inactiva.
- 00000 se utiliza cuando la línea está muerta.
- 00100 es una señal de detección.

El resultado final es que MLT-3, NRZI y 4B5B permiten al estándar 100Base-TX utilizar una frecuencia máxima de 31,25 MHz, que puede perfectamente emplearse con cables de tipo CAT5.

La especificación 100Base-T4 utiliza un método alternativo para resolver el problema de la alta frecuencia. En primer lugar, los datos se <<distribuyen>> sobre tres pares de hilos en lugar de utilizar un solo par. En segundo lugar, en vez de emplear 4B5B, se usa un esquema 8B6T que asigna a cada octeto un patrón de seis símbolos ternarios o tri-estados. El resultado es que solo, se requiere una frecuencia de 12,5 MHz, lo que significa que pueden emplearse los cables CAT3, más baratos, para transmitir 100 Mbps. Recuerde que los cables CAT3 admiten una frecuencia nominal máxima de 16 MHz.

Gigabit Ethernet

El estándar Gigabit Ethernet también tiene muchas especificaciones. Gigabit opera en dos modos primarios: dúplex y semidúplex. El modo semidúplex opera de forma muy parecida a Ethernet y Fast Ethernet, utilizando CSMA/CD. El modo dúplex utiliza un control de flujo basado en tramas definido en la norma IEEE 802.X.

Una de las especificaciones más comunes para Gigabit Ethernet es IEEE 802.3z. Esta especificación incluye las redes de tipo 1000Base-CX, 1000Base-SX y 1000Base-LX. La “C”, la “S” y la “L” quieren decir “Clúster”, “Short (corto)” y “Largo”, respectivamente. Las diferencias residen, principalmente, en el tipo de fibra utilizada (monomodo o multimodo) y el tipo de fuente luminosa (LED o Laser). Otra especificación popular es IEEE 802.3ab, que incluye 1000Base-T, donde se emplean cables de tipo CAT5. Todas las especificaciones forman parte del subnivel PHY del nivel Físico.

En el nivel PCS, Gigabit Ethernet emplea una codificación 8B10B, que es parecida al esquema 4B5B. La Ethernet a 10 Gigabits esta especificada en la norma IEEE 802.3ae. Este estándar opera con fibras multimodo más modernas, puede alcanzar hasta 300 m. existen varias implementaciones monomodo, que van desde 1300 nm con un rango máximo de 300 m, utilizando una tecnología denominada WWDM (Wide Wave Division Multiplexing, multiplexación por división de onda ancha), hasta 1550 nm con un rango máximo de 40 km.

La mayoría de estos formatos de trama y especificaciones físicas comparten un protocolo común en el subnivel MAC del nivel en Enlace de datos. Este protocolo es CSMA/CD (Carrier Sense, Multiple-Access/Collision Detect, acceso multiple y detección de colisiones mediante detección de portadora). El propósito de este protocolo consiste en gestionar el control de flujo y la contienda de acceso al medio. Puesto que Ethernet define los unos y ceros mediante transiciones en la tensión relativa del hilo, si dos dispositivos colocan una determinada tensión en el hilo al mismo tiempo, el bit resultara corrompido. Para evitar que dos dispositivos trasmitan al mismo tiempo, todas las especificaciones semidúplex empelan un mecanismo para detectar las colisiones (la palabra colisión es la que se utiliza en la jerga de Ethernet para describir el cambio de tensión que tiene lugar cuando dos señales se encuentran en un cable de transmisión, corrompiéndose ambas tramas). El circuito receptor esta ala escucha mientras el circuito trasmisor comienza a enviar y, si se detecta una colisión, ambos nodos dejan de transmitir datos y comienzan a transmitir una señal de interferencia de 32 bits. Un corto tiempo después, ambos nodos ejecutan el algoritmo de “reintento exponencial binario truncado” y vuelven a transmitir de nuevo. El algoritmo de reintento exponencial binario truncado significa que el nodo espera una cantidad aleatoria de tiempo a partir del momento en que se detectó la colisión y luego intenta volver a trasmitir adecuadamente sus tramas o hasta que se realicen 16 intentos, en cuyo caso el adaptador *informara* de que se ha producido un error.

El algoritmo utilizado para determinar el tiempo que ha de transcurrir hasta el reintento es el siguiente:

$$0 < r < 2^k$$

Dónde:

$$K = \min(n, 10)$$

En esta ecuación, r es el número de franjas temporales (intervalos de bit que hay que esperar antes de retransmitir y n es el número de intentos de retransmisión). A pesar de la existencia de numerosas opciones, tecnológicas y protocolos, Ethernet es una solución excelente en términos de facilidad de operación e instalación. En buena medida, implantar una red de este estilo es algo tan simple como enchufar un extremo de un cable en un concentrador y el otro extremo en una computadora.

Modos de Operación de Ethernet

Los dos modos principales de operación de Ethernet se denominan “dúplex” y “semidúplex”. La diferencia es, simplemente, que una conexión semidúplex permite el tráfico en ambas direcciones, pero solo en una dirección cada vez (una conexión simplex, como por ejemplo las emisiones de radio, solo permite que los datos se transfieran en una única dirección). Es decir, un nodo puede enviar o recibir, pero no hacer ambas cosas a la vez. El modo de operación dúplex, por el contrario, puede transmitir y recibir al mismo tiempo, lo que dobla en la práctica la tasa de transmisión. Es lo que a las especificaciones respecta, tanto 10BaseT como 100Base-TX. E incluso en este caso, la comunicación dúplex solo esta soportada entre un nodo y un conmutador. No puede utilizarse la operación dúplex al conectar un nodo a un concentrador. Observe también que en el modo de operación dúplex no se producen colisiones.

La segunda opción de operación es la que selecciona entre la autonegociación, la autodetección o la configuración manual. Puesto que se pueden utilizar múltiples formas de Ethernet sobre un cable de par trenzado, muchos fabricantes permiten que un solo adaptador soporte más de una especificación. Por ejemplo, la mayoría de las tarjetas de interfaz de red soportan tanto 10BaseT como 100Base-TX. Además, para 100Base-TX permiten tanto el modo de la operación dúplex como el semidúplex. Por tanto, cuando se conectan estas tarjetas a un concentrador, tanto este como la tarjeta NIC tienen que acordar que especificación deben utilizar para las comunicaciones.

El proceso de autonegociación está definido en la norma IEEE 802.3u, Sección 28, y permite que el nodo y el concentrador o conmutador determinen automáticamente el máximo común denominador. IEEE 802.3 sección 28B.3 enumera los siguientes modos soportados por la autonegociación:

- 100BaseT Dúplex
- 100BaseT4
- 100BaseT Semidúplex
- 10BaseT Dúplex
- 10BaseT Semidúplex

El proceso de autonegociación utiliza una señal de prueba especificada en 100BaseT. Esta señal se denomina pulso normal de enlace (NLP, Normal Link Pulse). Si un nodo recibe esta señal, presupondrá que el transmisor solo es capaz de operar según la norma 10BaseT. Si recibe un pulso de enlace rápido (FLP, Fast Link Pulse), que es un grupo especial de estas señales, dispuestos para formar una palabra de código, entonces el nodo sabrá que el transmisor es capaz de admitir cualquiera de las características que la palabra de código indique.

En la autonegociación, tanto el nodo como el concentrador o conmutador intercambian estas señales. En la autodetección, el dispositivo se limita a escuchar y a configurarse para ajustarse a su nodo homologo. La configuración manual, por supuesto, implica que una persona fije el modo de operación, obviando cualquier tipo de negociación.

Tecnologías Token Ring

En una arquitectura Token Ring, los datos fluyen en una única dirección. Si una estación quiere enviar una trama hacia el vecino situado aguas arriba (es decir, el nodo anterior dentro del anillo), lo que debe hacer en la práctica es enviar la señal hacia el nodo situado aguas abajo, que a su vez la reenviara hacia el siguiente nodo, etc., hasta que la señal viaje a lo largo del anillo y sea finalmente recibida por el nodo destino deseado. Pero la trama no habrá finalizado ahí su viaje, porque únicamente el transmisor puede eliminar la trama del

anillo. El nodo de destino regenera la trama al igual que los demás nodos del anillo y la reenvía, pero indicando, mediante un bit de la cabecera, que la trama ha sido recibida con éxito. Cuando el transmisor recibe la trama, comprueba el valor de ese bit indicador para saber si la trama fue recibida por el nodo de destino adecuadamente. En caso contrario, puede continuar retransmitiendo la trama. En la práctica, este mecanismo es una especie de sistema de confirmación integrado.

FDDI (Fiber Distributed Data Interface)

FDDI (Fiber Distributed Data Interface, interfaz de datos distribuidos por fibra) es un estándar ANSE que utiliza un método de acceso al anillo basado en el paso de testigo. La topología FDDI es, tanto desde el punto de vista lógico como físico, un anillo, aunque también poder ser físicamente una configuración en estrella. Sin embargo en lugar de utilizar un único anillo, como las redes Token Ring de IBM, FDDI utiliza dos anillos duales, dispuestos en sentido contrario. El anillo secundario solo se utiliza en case de fallo del anillo primario. Ambos anillos operan a una velocidad de 100 Mbps, pero se pueden conectar hasta 500 estaciones con toma dual, en una red de 100 km. Dado el gasto que implica la implementación de dos anillos, resulta posible conectar una estación a una red FDDI mediante un único cable. Este tipo de configuración se denomina estación de conexión simple.

Nota: segmentando el anillo, pueden combinarse dos anillos, lo que permite conectar hasta 1000 estaciones. Así que asegúrese de determinar si el anillo esta segmentado antes de responder a la pregunta de cuál es el número máximo de estaciones.

En un sistema de paso de testigo, solo el nodo que tiene el testigo puede transmitir datos hacia la red. Cuando una estación conectada a la red FDDI quiere transmitir, espera hasta que le llegue el testigo. De esta forma, podrá estar segura de que ninguno de otro dispositivo de la red intente transmitir. No existe la posibilidad de colisiones en una red de paso de testigo. FDDI gestiona la contienda por los recursos de red utilizando un protocolo de testigo temporizado. Cuando se inserta una estación en el anillo, la estación negocia la

cantidad de acceso de red de que podrá disfrutar. En otras palabras, se le garantiza un cierto tiempo de acceso. En el esquema de testigo temporizado, la estación puede comenzar a transmitir en cuanto capture el testigo y podrá continuar transmitiendo hasta que el intervalo de temporización caduque o hasta que no tenga más tramas para transmitir. Entonces, la estación enviara el testigo hacia el siguiente nodo situado aguas abajo.

El sistema FDDI también utiliza direcciones MAC y el subnivel LLC de IEEE 802.2, al igual que Ethernet y Token Ring. Una diferencia, sin embargo, es que en la especificación de las direcciones MAC se incluyen dos bits reservados. Aunque Ethernet implementa solo el bit I/G, FDDI implementa tanto el bit I/G como el bit U/L. Este último bit, U/L, indica si una dirección está administrada universal o localmente. Esto permite a los administradores de red cambiar las direcciones MAC por algún tipo de esquema interno; si lo hacen, serán los administradores los responsables, en lugar del IEEE y del fabricante, de garantizar que no haya ningún duplicado. En el campo de dirección de origen, mientras que Ethernet siempre comienza la dirección MAC con un cero, FDDI utiliza el primer bit como indicador de información de encaminamiento (RII, Routing Information Indicator). El campo RII indica si hay presente información de encaminamiento (RIF, Routing Information Field). FDDI no utiliza casi nunca el encaminamiento de origen.

Una de las características de FDDI (y Token Ring) es que en estas redes se transmiten los bytes en un orden diferente al que se utiliza en Ethernet. Específicamente, FDDI y Token Ring transmiten los bits, enviando en primer lugar el bit situado más a la derecha. Esto puede causar problemas serios cuando se trate de enviar tramas a través de un puente que conecte una red FDDI y una red Ethernet, porque las direcciones MAC estarán invertidas.

Otra distinción importante es que una trama FDDI puede contener hasta 4500 (4472) bytes de datos. Esta hace que FDDI pueda ser mucho más eficiente que Ethernet, aunque también puede hacer que aparezca un retardo significativo mientras que las otras estaciones esperan a que esa trama sea procesada. Y, por supuesto, si se envía una trama FDDI a través de un puente hacia una red Ethernet, cuyo tamaño máximo es de 1500 bytes de datos, la fragmentación de la trama puede constituir un problema.

Las redes basadas en paso de testigo, incluyendo FDDI y Token Ring, utilizan un monitor del anillo, para realizar funciones de gestión en el anillo. Sin embargo, FDDI difiere de Token Ring en que FDDI se distribuye el papel de monitor del anillo entre una serie de dispositivos, mientras que en las redes Token Ring se designa a un único nodo como monitor activo.

A diferencia de Ethernet, cada nodo FDDI emplea varios temporizadores, incluyendo:

- TTRT (target token rotational timer, temporizador objetivo de rotación del testigo), que se configura durante la inicialización del anillo y representan el retardo del anillo.
- TRT (token rotational timer, temporizador de rotación del testigo), que mide el tiempo que el testigo necesita para recorrer el anillo.
- THT (token holding timer, temporizador de retención del testigo), que establece la cantidad de tiempo que un nodo puede retener el testigo antes de pasárselo a la siguiente estación
- TVX (valid transmisión timer, temporizador de transmisión válida), que se calcula a partir del tiempo necesario para enviar la trama más larga (4500 bytes) alrededor del anillo más grande posible (200 km). Este temporizador se utiliza para permitir que los nodos se reinicialicen más rápidamente.

Estos temporizadores se utilizan para reiniciar el anillo en caso de que una trama tarde más de dos veces el tiempo TTRT en recorrer el anillo. FDDI también utiliza codificación 4B5B, al igual que Ethernet, aunque no emplea codificación Manchester. Como tecnología, FDDI fue bien recibida y se la implementó de forma general como red troncal o para conectar una serie de servidores de misión crítica y alto ancho de banda, debido a su velocidad, a las grandes distancias alcanzables y a su redundancia inherente. Sin embargo, FDDI ha sido prácticamente sustituida por Gigabit Ethernet y otras tecnologías.

Redes Token Ring de IBM e IEEE 802.5

Token Ring fue creada originalmente por IBM y posteriormente estandarizada en la norma IEEE 802.5 en 1985. Como el nombre sugiere, se trata de una tecnología de paso de testigo que, desde el punto de vista lógico, es un anillo, aunque puede ser algo confusa físicamente. Esto se debe a que el circuito eléctrico utilizado es físicamente un anillo, en el sentido de que debe pasar a través de cada dispositivo, formando un bucle de gran tamaño. Pero la red se asemeja a una estrella, porque todos los cables se suelen tender desde cada nodo hasta un concentrador. Por esta razón, resulta perfectamente aceptable hacer referencia a esta red como si tuviera una topología en anillo o en estrella.

Al igual que Ethernet y FDDI, Token Ring utiliza el sistema MAC de direccionamiento, como FDDI, los bytes se transmiten en el orden en el que se escriben. Originalmente, Token Ring permitía velocidades de anillo de 4 Mbps y 14 Mbps. Recientemente, se ha hecho esfuerzos de desarrollo de una red Token Ring de 100-Mbps, pero esa tecnología no ha sido ampliamente adoptada, ya que es casi dos órdenes de magnitud más lenta que ATM y que Ethernet.

Token Ring utiliza codificación Manchester diferencial, que es ligeramente distinta de la codificación Manchester utilizada en Ethernet. Es la codificación Manchester diferencial, resulta posible que la tensión efectúe una transición al principio del intervalo de bit, de la misma forma que en la codificación Manchester normal se hace en mitad del intervalo de bit. La presencia de una transición indica un uno, mientras que su ausencia indica un cero. La dirección de la transición al principio o en mitad de un intervalo de bit no tiene ningún significado.

El formato de una trama Token Ring es considerablemente más complejo que Ethernet, porque varios de los campos tienen un tamaño variable y el formato soporta la inclusión de información de encaminamiento de origen.

El tamaño máximo de una trama de Token Ring también resulta un tanto complicado de

calcular. El máximo es, usualmente, de 4096 bytes, 4500 bytes o 4472 bytes. Pero, en realidad, el tamaño máximo de trama está en función de la velocidad del anillo y del tiempo de retención del testigo (THT, Token Hold Time). Por tanto, en un anillo a 16 Mbps, con un tiempo máximo de retención resulta posible tener una trama de 17.800 bytes.

Aunque Token Ring utiliza el concepto del anillo, al igual que FDDI, la red Token Ring designara una única estación como monitor activo y otra estación como monitor de reserva, por si acaso el monitor activo falla. Esta estación actúa como fuente de información de monitorización y lleva a cabo una serie de funciones de mantenimiento del anillo.

1 byte	1 byte	1 byte
Delimitador de inicio	Control de acceso	Delimitador final

Tabla 2.8 los tres campos de un testigo.

Un ejemplo de este tipo de funciones de gestión seria la eliminación de una trama que estuviera circulando de manera continua y la generación de un nuevo testigo. Otra función, que es realizada por todas las estaciones, se denomina función de baliza. Este tipo de función se lleva a cabo cuando cualquiera de las estaciones detecta un problema en la red, como por ejemplo un cable roto. La función de baliza se utiliza para definir un dominio de fallo y hacer que el resto de la red entre en un proceso de autorreconfiguración.

Para que el proceso de autorreconfiguración funcione, cada estación tiene que conocer cuál es el vecino activo más próximo situado aguas arriba (NAUN, Nearest Active Upstream Neighbor). Para obtener esta información, se ejecuta un proceso denominado sondeo del anillo cada siete segundos. Al igual que Ethernet y FDDI, el estándar IEEE 802.5 emplea los protocolos LLC IEEE 802. Token Ring también tiene un proceso de control de prioridades que permiten a ciertas estaciones acceder al testigo más frecuentemente que otras. Parte de este proceso requiere que el propio testigo tenga una prioridad y, además, cada estación del anillo tiene también una prioridad asignada. Las estaciones pueden capturar el testigo solo si la prioridad actual de este es inferior a la prioridad de la estación. Una vez que una estación ha finalizado con el testigo, vuelve a ajustar la prioridad a su

valor original.

La función de encaminamiento de origen de Token Ring permite conectar varios anillos entre sí a través de puentes. A cada puente y a cada anillo se les asigna un número. Si un nodo desea comunicarse con otro nodo, envía una trama de exploración de rutas, (ARE, All Routes Explorer), que será detectada por el puente y coloca en todos los anillos adyacentes. Al hacer esto, el puente registra el número del anillo y su propio número de puente en la trama ARE. Este proceso se repite hasta que la trama llega eventualmente a su destino. El nodo destino marcará la trama y ejecutará el proceso inverso. Cuando la trama ARE vuelva a llegar al nodo origen, este dispondrá de un registro de todos los puentes y anillos que ha habido que cruzar para hacer el viaje de ida y vuelta. Si existen múltiples trayectos, el nodo de origen recibirá una trama ARE para cada una de las posibles rutas. Cuando el nodo de origen esté listo para enviar datos hacia el destino, evaluará todas las tramas ARE que haya recibido e incluirá la ruta más corta en el campo de información de encaminamiento (RIF, Routing Information Field) de la cabecera Token Ring. Cuando el puente reciba esta trama, leerá el campo RIF y reenviará la trama hacia el siguiente anillo indicado en dicho campo.

Aunque resulta algo compleja y no muy escalable, esta función de encaminamiento de origen permite que los nodos sean muy explícitos en lo que respecta a la manera en la que quieren que se gestionen sus tramas. Observe también que este proceso se basa en la utilización de puentes, no de encaminadores. El encaminamiento de origen se utiliza normalmente con protocolos no encaminables, como por ejemplo SNA y NetBIOS de IBM.

Otros Protocolos LAN

Algunos estándares, como IEEE 802.4 Token Bus, nunca han pasado de ser meros proyectos, mientras que otros como ARCNet, Bus & Tag y el protocolo banda base de IBM carecían de la capacidad de expansión suficiente y fueron, por tanto, abandonados. Muchas tecnologías más recientes, como USB (Universal Serial Bus, bus serie universal) y Firewire (IEEE 1394) están ganando una rápida aceptación. Por supuesto, hay muchas otras tecnologías más antiguas que se siguen utilizando, como los enlaces serie (el puerto COM

de un PC) y Fibre-Channel, que se utiliza principalmente en redes de área de almacenamiento. Y hay tecnologías como ESCON (Enterprise System Connection, conexión de sistemas empresariales), que se utilizan casi exclusivamente en entornos mainframe.

ARCNet fue un protocolo banda base que gozo de una popularidad efímera durante la década de 1980. Utilizaba un método de codificación basado en desplazamiento de frecuencia. Básicamente, durante un intervalo de bit, un ciclo completo representa un uno, mientras que dos ciclos completos representan un cero. ARCNet es también un protocolo basado en paso de testigo extremadamente simple. Desde el punto de vista lógico, se trata de un anillo y es determinista, al igual que la tecnología Token Ring de IBM. Los tamaños de paquete iban de 0 a 507 bytes.

IEEE 802.4 Token Bus es muy similar a ARCNet, en el sentido de que fue creado también en la década de 1980 (principalmente por General Motors), es un protocolo basado en paso de testigo y utiliza el mismo método de codificación por desplazamiento de frecuencia; la diferencia es que se trata de un protocolo de banda ancha, en lugar de banda base. Pero, a diferencia ARCNet, 802.4 es extremadamente complejo, constando la especificación de varios cientos de páginas. Además, se comporta como anillo lógico, pero se trata de un bus. Permitía velocidades de 5 y 10 Mbps. General Motors le dio el nombre de Manufacturing Automation Protocol, protocolo de automatización de fabricación.

ESCON (Enterprise System Connection) es una tecnología de conexión a canal creada por IBM en 1990 para reemplazar la tecnología Bus & Tag en cables de cobre. ESCOM utiliza cables de fibra óptica con un tamaño máximo de cada enlace de 3 km (sin repetidores). Esto proporciona dos tipos principales de canales, denominados multiplexor de bytes y multiplexor de bloques. El ancho de banda máximo de un único canal es de 18,6 Mbps (o 148,8 Mbps) y puede soportar una distancia máxima de 8 km.

2.2 Tecnologías Para Red de Área Extensa

La mayoría de las tecnologías [REDCOMTITTEL] para redes de área extensa (WAN, wide area network) difieren enormemente de sus equivalentes LAN en varios sentidos:

- Se les diseña normalmente teniendo presentes las necesidades de los operadores de telecomunicaciones (que a menudo necesitan dar conexión a decenas de miles de clientes) por lo que son extremadamente escalables.
- El nivel Físico tiene usualmente una distancia máxima de entre 3 y 60 km.
- Están especificadas con muchas velocidades diferentes, desde 56 Kbps o menos hasta 10 Gbps.
- A menudo utilizan técnicas de multiplexación para permitir la existencia de varios circuitos lógicos sobre un mismo circuito físico.

Específicamente, mientras que la mayoría de las tecnologías LAN se venden en forma de producto terminado (es decir, el cliente compra un concentrador Ethernet o Token Ring, algunos cables y tarjetas adaptadoras para sus computadoras y es libre de hacer lo que desee a continuación), las tecnologías WAN se han diseñado pensando en alquilar los circuitos, es decir, en que el cliente pague una tarifa mensual por la utilización del circuito, más una tarifa de uso basada en la cantidad de tráfico enviada a través de la red del operador de telecomunicaciones.

Frame Relay

Frame Relay hizo su aparición en 1988, cuando los desarrolladores de redes RDSI de dieron cuenta de que el protocolo LAPD (Link Access Protocol-D, protocolo de acceso a enlace-D), que se utilizaba para proporcionar señalización para el canal D de un circuito RDSI, podía emplearse para muchas otras cosas. Esto dio como resultado la recomendación 1.122 de la ITU-T, donde se especificaban una serie de servicios adicionales en modo

paquete. Este protocolo está compuesto por varios estándares ANSI e ITU-T y la mitad de ellos son compartidos con RDSI, por lo que no existe ninguna fuente centralizada donde poder ver el estándar completo, como sucede con los protocolos IEEE. Pero un buen lugar por donde empezar es el foro Frame Relay y los estándares ITU-T Q.922 y Q.933. El foro Frame Relay (www.frforum.com) es una organización sin ánimo de lucro compuesta por unas 300 empresas, que trata de promover Frame Relay y publica una serie de acuerdos de implementación. Estos acuerdos detallan la manera de resolver ciertos problemas, como los que atañen a los circuitos virtuales conmutados, a la fragmentación, al tráfico de voz y otros.

Frame Relay está basada en un concepto denominado circuito virtual (VC, virtual Circuit). Un circuito virtual es una ruta bidireccional a través de la red que se define por software. La ventaja principal del Frame Relay es que pueden implementarse muchos circuitos virtuales sobre una única conexión física. Por ejemplo, si la sede central de una empresa necesita comunicarse con tres oficinas remotas, en lugar de arrendar tres circuitos punto a punto para conectar cada una de estas tres oficinas, se puede utilizar un único circuito que conecte cada una de ella a la red Frame Relay, empleando tres circuitos virtuales. Desde el nivel 3 hacia arriba, lo que tenemos en realidad son tres circuitos físicos.

Frame Relay permite implementar los circuitos virtuales utilizando lo que se denomina identificadores de conexión de enlace de datos (DLCI, Data Link Connection Identifiers), que son números de 10 bits que tienen un significado local. Esto quiere decir que el DLCI es uno de los circuitos virtuales de la oficina central podría ser igual que el DLCI de dicho circuito virtual en la sucursal, o podría ser un número completamente distinto. No importa que estos números difieran. Lo que sí importa es que los tres circuitos virtuales existentes en el enlace de la oficina central deben ser diferentes.

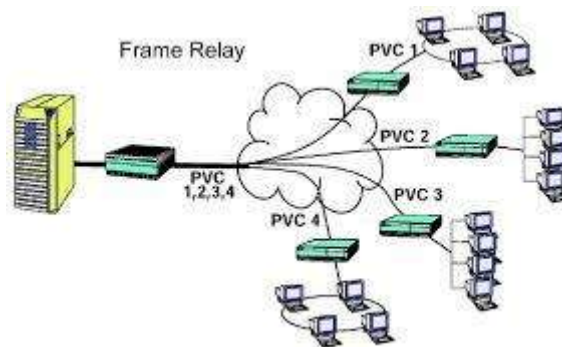


Figura 2.2 Red Frame Relay

6 bits	1 bit	1 byte	4 bits	1 bit	1 bit	1 bit	1 bit
DLCI	CR	Bit. Ext.	DLCI	FECN	BEDN	DE	EA

Tabla 2.9 Una Cabecera Frame Relay

Estos campos son:

- La primer parte del DLCI.
- Un bit Comando/Respuesta que indica si se trata de una trama de comando o de respuesta.
- Un bit de extensión que permite utilizar una cabecera de tres o cuatro bytes para ampliar el espacio de direcciones, en caso necesario.
- La segunda parte de DLCI.
- El bit FECN (Forward Explicit Congestion Notification, notificación explícita directa de congestión) indica al dispositivo receptor de esta trama que se ha detectado una situación de congestión.
- El bit DE (Discar Eligible, elegible para descarte) indica que esta trama debe descartarse en primer lugar en caso de congestión
- El bit BECN (Backward Explicit Congestion Notification, notificación explícita inversa de congestión) es activado por un conmutador de la red Frame Relay para permitir al dispositivo receptor conocer que las tramas que el receptor está transmitiendo están provocando congestión.
- El bit de Extensión permite emplear a una cabecera de tres o cuatro bytes.

Existen dos tipos de circuitos virtuales: circuitos virtuales permanentes (PVC, Permanent Virtual Circuit) y circuitos virtuales conmutados (SVC, Switched Virtual Circuit). Los circuitos virtuales permanentes se configuran de manera manual. En un circuito virtual conmutado, cuando un nodo tiene datos para enviar, la red Frame Relay establece un circuito virtual sobre la marcha (el SVC) y transmite luego los datos a su través. Una vez

terminada la transmisión, después de un periodo de inactividad, el circuito virtual conmutado se elimina. De esta forma, la red Frame Relay puede ser extremadamente flexible, porque se pueden mover los dispositivos de un sitio a otro sin tener que configurar de forma manual las rutas virtuales.

SMDS (Switched Multimegabit Data Service)

SMDS (Switched Multimegabit Data Service, servicio conmutado de datos multimegabit) es una red pública de conmutación de paquetes principalmente utilizada para enviar grandes cantidades de datos de forma esporádica. SMDS se parece a la RTGC (Red Telefónica General de Conmutación) en el sentido de que se asigna a las empresas direcciones que son números unívocos de 10 dígitos. A diferencia de Frame Relay, no hay circuitos virtuales. La ruta a través de la red es completamente dinámica, por lo que cualquier empresa puede enviar datos a cualquier otra sin necesidad de configurar ningún circuito. Puesto que los paquetes SMDS contienen hasta 7168 bytes, pueden enviarse la mayoría de las tramas LAN más comunes sin ninguna fragmentación.

RDSI (Red Digital de Servicios Integrados)

RDSI (Red Digital de Servicios Integrados), también llamadas ISDN (Integrated Services Digital Network), es un conjunto de estándares ITU-T para la transmisión de datos digitales, principalmente a través de líneas telefónicas estándar de cobre. Este servicio es bastante complejo pero muy flexible. RDSI tiene dos niveles de servicios: BRI (Basic Rate Interface, acceso básico) y PRI (Primary Rate Interface, acceso primario). Estos servicios utilizan dos tipos de canales: canales “B” y canales “D”. Los canales B proporcionan 64 Kpbs y se utilizan para transportar voz o datos. En el acceso básico, hay dos canales B y un canal D, el cual tiene 16 Kpbs. El canal D se utiliza para transportar información de control y de señalización. En el acceso primario (PRI), hay 23 canales B de 64 Kpbs cada uno y un canal D de 64 Kpbs. De hecho hay varios tipos de canales, incluyendo A, B, C, D, E y H, pero los más comúnmente utilizados son los canales B y D.

Aunque muchas organizaciones de tamaño medio y grande tienen al menos un acceso de tipo PRI, resulta muy común utilizar el término “conexión RDSI” para referirse solo a un acceso de tipo BRI.

En un entorno típico, se utiliza una conexión BRI para llevar un circuito de voz y otro de datos a un domicilio o a una oficina de pequeño tamaño, mientras que un acceso PRI se utiliza para transportar 23 canales de voz o algún tipo de combinaciones de 23 canales de voz y datos. Los accesos BRI se suelen utilizar en configuraciones de <<acceso telefónico>>, en las que un modem RDSI solo se conecta al proveedor de servicios (y por tanto solo genera facturación) cuando hay datos para enviar. Los accesos PRI son, usualmente, conexiones permanentes.

RDSI se implementa mediante una compleja serie de dispositivos. Los dispositivos se conectan mediante interfaces denominadas “puntos de referencia”. Comenzando por la red del operador de telecomunicaciones y yendo hacia el usuario final, lo que tenemos es un conmutador RDSI conectado a un dispositivo de terminación de red tipo 1 (NT1, Network Terminating device 1) a través de la interfaz U. El NT1 se conecta a un NT2 a través de la interfaz T. El N2 se conecta a un terminal de datos (TE1) o aun teléfono RDSI (TE1) utilizando el punto de referencia S. En la práctica, a menudo se suelen combinar puntos de referencia S y U en un único elemento hardware, denominado S/T. las interfaces son:

Terminal de datos-S-NT2-T-NT1-U-Conmutador RSDI.

Existen en la práctica otro punto de referencia adicional: el punto de referencia R, que se utiliza para conectar un dispositivo no RDSI a un adaptador de terminal RDSI (TA, Terminal Adapter). El adaptador de terminal tiene un punto de referencia S que permite que un dispositivo no RDSI de comuniqué a través de una red RDSI, de la forma siguiente:

Terminal no RDSI-R-TA-A-NT2-T-NT1-U-Conmutador RDSI.

SONET (Synchronous Optical NETwork)

SONET (Synchronous Optical NETwork, red óptica síncrona) es un conjunto de estándares que definen las velocidades y formatos para redes ópticas especificado en ANSI T1.105, ANSI Y1.106 y ANSI T1.117. SDH (Synchronous Digital Hierarchy, jerarquía digital síncrona) es un estándar similar definido por ITU-T y utilizado principalmente en Europa. El formato de trama utilizado por SONET es SMT (Synchronous Transport Module, módulo de transporte síncrono). STM-1 es la señal de nivel base, que es de 155 Mbps y esta transportada mediante una señal OC-3. Se dice que este sistema es jerárquico porque se pueden multiplexar entre si varios niveles de señal pequeños para formar otros más grandes.

Una trama STS-1 tiene nueve filas de 90 Bytes. Los primeros tres bytes de cada fila son bytes de control que contienen bits de tramado. La información de control de línea se transmite dentro de la carga útil, en una posición variable que está determinada por el puntero contenido en la información de control de sección. Una trama STS-1 se transmite en 125 μ sec, lo que equivale a 8000 tramas por segundo (a 810 bytes por trama X 8000 tramas por segundo = 51,84 Mbps).

Existen varios niveles de orden superior, como STS-3, que está compuesto por ocho filas de 270 bytes, con nueve bytes de información de control por cada fila. SONET se utiliza principalmente en redes de área metropolitana (MAN, Metropolitan Area Networks), donde el operador telefónico tiene cables de fibra óptica para formar una serie de bucles alrededor de una ciudad. Esta tecnología de red resulta ideal, porque el anillo puede operar a una velocidad de nivel OC-12 (622 Mbps) y cuatro empresas distintas pueden compartir el alquiler de un circuito, con lo que este será para ellas un circuito de tipo OC-3 (155 Mbps). La acción de reservar ancho de banda de un circuito para crear otro circuito se denomina “provisión de ancho de banda”.

SONET tiene una ventaja adicional, que es su redundancia inherente, muy similar a la de FDDI. Utilizando una arquitectura dual contrapuesta para los dos anillos, el anillo puede compensar de forma inmediata la ruptura de un enlace de fibra o el fallo de un único

equipo, esta propiedad se denomina en ocasiones “auto- curación”.

PPP (Protocolo Punto a Punto)

PPP (protocolo punto a punto) es un elemento muy importante del puzzle de las conexiones en red. Originalmente diseñado para encapsular IP en enlaces serie punto a punto, PPP soporta ahora muchos otros protocolos, como por ejemplo IPX de Novell y DECnet de DEC. También tiene una multitud de opciones y funcionalidades, incluyendo gestión de direcciones IP, autenticación, multiplexación y otras funciones de gestión, tales como configuración, pruebas, detección de redes, etc. Se utiliza comúnmente en las computadoras dotadas de modem, para acceder telefónicamente a internet o a una red corporativa. También se utiliza comúnmente en las redes de área extensa empresariales para enlaces a velocidades comprendidas entre 56 K y T1 (1,544 Mbps).

PPP está compuesto de un protocolo de control de enlace de datos de alto nivel (HDLC, High-Level Data Link Control), un protocolo de control de enlace de red (NCP, Network Control Protocol). HDLC se utiliza para encapsular datagramas a través de enlaces serie. LCP establece, configura y prueba la conexión de nivel de enlace de datos. El protocolo NCP se emplea para establecer y configurar uno o más protocolos de nivel de red. PPP opera entre un equipo terminal de datos (DCE, Data Communications Equipment). El enlace entre estos dispositivos debe ser dúplex, y puede operar en modo síncrono o asíncrono.

1 byte	1 byte	1 byte	2 byte	Variable	2 0 4 bytes
Bandera	Dirección	Control	Protocolo	Datos	FCS

Tabla 2.10 una trama PPP

La bandera simplemente marca el inicio de una trama. Siempre tiene el valor 01111110 en binario. El campo de direcciones es siempre 11111111, que es una dirección de difusión, porque PPP no define direcciones de estación. El campo de control es siempre 00000011, lo

que indica un servicio de enlace sin conexión, similar a LLCI. El campo de datos, por supuesto, contiene el datagrama, que tiene teóricamente un máximo de 1500 bytes, pero puede cambiarse en algunas circunstancias. La secuencia de control de trama (FCS, Frame Check Sequence) es un valor de 16 o 32 bits que se utilizan para detectar errores en la trama. Funciona igual que los campos FCS de casi todos los otros protocolos de enlace de datos que se han descrito anteriormente.

PPP utiliza los procedimientos HDLC de ISO (ISO 3309-1979), modificados por la norma ISO 3309:1984/PDAD1 <<Anexo 1: transmisión inicio/parada>>. ISO 3309-1979 se utiliza en entornos sincronizados, mientras que ISO 3309: 1984/PDAD1 se emplea en entornos asíncronos. La norma ISO 4335-1979/Anexo 1- 1979 especifica los procedimientos de control y el esquema de codificación.

HDLC (High-Level Data Link Control)

HDLC (High-level Data Link Control, control de enlace de datos de alto nivel) es un derivado del protocolo SDLC (Synchronous Data Link Control, control de enlace de datos síncrono) desarrollado por IBM a mediados de la década de 1970 para SNA, SDLC fue el primer protocolo síncrono de nivel de enlace de datos orientado a bit. IBM envió SDLC para su estandarización y, como suele suceder, ISO lo modifico y lo denomino HDLC, mientras que ITU-T también lo modifico y lo denomino LAPB (Link Access Procedure Balanced, procedimiento equilibrado de acceso de enlace). Después IEEE modifico HDLC para entornos LAN, creando la norma IEEE 802.2.

Resulta importante comprender HDLC, porque son muchos protocolos, como PPP que lo utilizan. Y muchos otros emplean el protocolo LLC IEEE 802.2, incluyendo Ethernet, Token Ring, FDDI y otros. Fundamentalmente, el formato de trama de HDLC es idéntico al de SDLC.

1 byte	1 0 2 bytes	1 0 2 bytes	2 bytes	Variable	2 0 4 bytes	1 byte
Bandera	Dirección	Control	Protocolo	Datos	FCS	Bandera

Tabla 2.11 formato de trama HDLC.

Estos campos ya nos resultan familiares, a partir de las explicaciones sobre PPP. Pero, a diferencia de PPP, en SDLC y HDLC estos campos sí que tienen un significado. El campo de dirección contiene la dirección de la estación secundaria, que puede ser una dirección física, una dirección de grupo o una dirección de difusión. El campo de control utiliza tres diferentes formatos: I, S y U. las tramas I, o de información, transportan información de los niveles superiores e información de control.

7 bits	1 bit	7 bits	1 bit
Numero de secuencia de recepción	Bit E/F	Numero de secuencia de envío	0

Tabla 2.12 Formato del campo de control en una trama I.

La operación de los números de secuencia en estas tramas es similar a la de TCP. La estación primaria utiliza el bit P/F (Poll/Final, sondeo/final) para sondear a la estación secundaria, lo que indica a esta que se espera de ella una respuesta inmediata. La estación secundaria utiliza este mismo método bit en las tramas que viajan en la dirección opuesta, para informar a la estación primaria de si la trama actual es, o no, la última de su respuesta. Las tramas S son tramas de supervisión que proporcionan más información de control, con la que se puede, por ejemplo, suspender la transmisión o confirmar la recepción de tramas I.

7 bits	1 bit	7 bits	1 bit	1 bit
Numero de secuencia de recepción	Bit P/F	Código de función	I	I

Tabla 2.13 Formato del campo de control de las Tramas S.

Las tramas U son tramas no numeradas (Unnumbered). Permiten realizar funciones de control, como por ejemplo la inicialización de dispositivos. Estas tramas no tienen ningún número de secuencia.

7 bits	1 bit	7 bits	1 bit	1 bit
Código de función	Bit P/F	Código de función	I	I

Tabla 2.14 Formato del campo de control en las tramas U.

Mientras que SDLC solo permite un modo de transporte, HDLC soporta tres: el modo NRM (Normal Response Mode, modo normal de respuesta) es el utilizado por SLDC, y en él la estación secundaria no puede comunicarse con la primaria hasta obtener el permiso correspondiente. El modo ARM (Asynchronous Response Mode, modo asíncrono de respuesta) permite a las estaciones secundarias iniciar una comunicación sin previo permiso. En el modo ABM (Asynchronous Balanced Mode, modo asíncrono equilibrado), cada nodo concreto puede ser tanto nodo primario como secundario al mismo tiempo. Este modo se denomina “modo combinado”.

LLC (Logical Link Control)

Aunque la especificación LLC (Logical Link Control, control de enlace lógico) de la norma IEEE 802.2 se utiliza principalmente en redes LAN. LCC proporciona tres tipos de servicio denominados LLC1, LLC2 y LLC3. Estas clases de servicio ofrecen una amplia variedad de enlaces, desde enlaces rápidos con poca información de control y escasos mecanismo de control de flujo, hasta enlaces más lentos y fiables, orientados a conexión. Las clases de servicios son las siguientes:

- LLC de tipo 1 es un servicio sin conexión y sin confirmaciones. Depende de la utilización de protocolos tales como TCP para realizar las tareas de corrección de errores. Debido a esto, es el protocolo más comúnmente utilizado. Esencialmente, solo proporciona los campos DSAP (Destination Service Access Point, punto de acceso a servicio de origen), que tienen una función similar a la del campo Ethertype en las tramas de Ethernet versión II).
- LLC de tipo 2 es un servicio orientado a conexión y con confirmaciones. Esto quiere decir que se establece una conexión fiable entre dos nodos antes de que pueda transmitirse ningún dato y que, una vez transmitidos los datos, se confirman cada transmisión. Si se produce un error, LLC tipo 2 volverá a enviar la trama. Este tipo de servicio raramente se utiliza.
- LLC de tipo 3 es un compromiso entre los tipos 1 y 2. Se trata de un enlace con confirmaciones, pero que no está orientado a conexión.

2.2.1 Redes Inalámbricas

Aunque las redes inalámbricas en la forma de radio AM, FM y de onda corta son ya conocidas desde hace tiempo, las redes inalámbricas dúplex con soporte de voz y datos son relativamente recientes y están ganando en popularidad rápidamente [REDCOMTITTEL]. Al igual que en las redes cableadas, el número de oscilaciones por segundo de onda se denomina “frecuencia” y se mide en hercios (Hz). La distancia entre dos picos o dos valles adyacentes en la onda se denomina “longitud de onda”, normalmente representada con la letra griega lambda (λ).

Frecuencias de Radios

Las ondas de radios son muy comunes de en mundo de las redes, porque pueden viajar a grandes distancias, atravesar las paredes y son relativamente baratas de generar. El comportamiento de las ondas de radio también depende de la frecuencia, de modo que a altas frecuencias las ondas tienden a viajar en línea recta y a reflejarse en los obstáculos. Sin embargo, a bajas frecuencias, las ondas tienden a pasar a través de las paredes, aunque tienen unas limitaciones de distancia mucho más estrictas, de modo que el medio de transmisión utilizado es, como se puede ver, bastante flexible.

Un problema con las redes de radio es que dos dispositivos que utilicen la misma frecuencia se interfieran mutuamente. En un intento de evitar, los gobiernos de la mayoría de los países regulan la utilización de las frecuencias. Los gobiernos conceden licencias sobre ciertas bandas de frecuencia a determinadas empresas y se reservan otras bandas para su utilización general por parte del público con bajos niveles de potencia.

Nombres de la banda	Extremo inferior	Extremo superior
Muy baja frecuencia (VLF)	3 KHz	30 kHz
Baja frecuencia (LF)	300 kHz	3 MHz
Alta frecuencia (HF)	3 MHz	30 MHz
Muy alta frecuencia (VHF)	30 MHz	300 MHz
Ultra alta frecuencia (UHF)	300 MHz	3 GHz

Tabla 2.15 Banda de radiofrecuencia.

Existen numerosas bandas con frecuencias más bajas, pero normalmente se las utiliza, porque, al ser de frecuencia más baja, también lo es el ancho de banda.

Las bandas de radio continúan por encima de la banda de UHF, aunque esas bandas raramente se utilizan:

- Súper alta frecuencia (SHF, Super High Frequency).
- Extremadamente alta frecuencia (EHF, Extremely High Frequency).
- Tremendamente alta frecuencia (THF, Tremendously High Frequency).

Nombre de la banda	Extremo inferior	Extremo superior
Subaudible (por debajo del rango de audición humano)	1 Hz	30 Hz
Subaudible (audible, pero por debajo del rango típico de la voz)	30 Hz	300 Hz
Audible (rango típico de la voz humana)	300 Hz	3000 Hz
Ultrasonidos (por encima del rango de audición humano, aunque los perros pueden continuar oyendo estos sonidos)	3 kHz	30 KHz

Tabla 2.16 Bandas de radio de baja frecuencia.

Frecuencias de Microondas

Los microondas son un subconjunto de las frecuencias de radios que generalmente se considera que comienzan a 1 GHz y terminan a unos 18 GHz. Las bandas por encima de 18 GHz se suelen denominar bandas de ondas “milimétricas”. La mayoría de los técnicos del mundo profesional y militar asocian una serie de designaciones de letras con las bandas del rango de las microondas.

Nombre de la banda	Extremo inferior	Extremo superior
Banda L	1 GHz	2 GHz
Banda S	2 GHz	4 GHz
Banda C	4 GHz	8 GHz
Banda X	8 GHz	12 GHz
Banda Ku	12 GHz	18 GHz

Tabla 2.17 Bandas de frecuencia de microondas.

Las bandas milimétricas también tienen designaciones de letras, incluyendo K, Ka, W y otras. Hay un tipo importante de banda que es preciso resaltar: las bandas ISM (Industrial/Scientific/ Medical, industrial/científica/médica). Estas bandas caen dentro del rango de frecuencias comprendido entre 2,400 y 2,484 GHz. En algunos países, también se utilizan las bandas comprendidas entre 902 y 928 MHz y entre 5,725 y 5,850 GHz. Si se examinan un teléfono inalámbrico, probablemente vea una indicación en el que se especifican 900 MHz o 2,4 GHz.

Una de las tecnologías inalámbricas más populares para las redes LAN es la especificada en la norma IEEE 802.11b, que opera en el rango de 2,4 GHz. Este estándar es muy similar a Ethernet en el nivel 2, pero utiliza tecnologías de expansión de espectro (DSSS, Direct Sequence Spread Spectrum) en el nivel físico.

Ondas Infrarrojas

Las ondas infrarrojas están comprendidas entre el espectro visible (es decir, el espectro correspondiente al arco iris) y las microondas. La tecnología infrarroja se utiliza normalmente en configuraciones de carácter muy local, como por ejemplo, en controles remotos de televisiones y magnetoscopios y en los puertos infrarrojos de las mayorías de las computadoras portátiles. Esto se debe a que las transmisiones infrarrojas son direccionales (es preciso apuntar aproximadamente hacia el dispositivo con el que nos queremos comunicar) y a que las ondas infrarrojas no pueden atravesar objetos sólidos, como por ejemplo las paredes.

Como ventaja, estas características hacen que la tecnología infrarroja sea más segura que las tecnologías de radio y de microondas, y los gobiernos no regulan la utilización de este tipo de comunicaciones. El estándar IEEE 802.11b también incluye una especificación para comunicaciones infrarrojas, aunque estas comunicaciones están limitadas a una distancia de 10 metros y no pueden atravesar las paredes.

2.3 Protocolos para Redes Inalámbricas

Las redes inalámbricas son aquellas que posibilitan la interconexión de dos o más equipos entre sí sin que intervengan cables, constituyendo así un eficaz medio para la transmisión de cualquier tipo de dato. Las redes inalámbricas se basan en un enlace que utiliza ondas electromagnéticas en lugar de cableado estándar.

2.3.1 WEP

Las redes Wi-Fi incorporan la posibilidad de encriptar la comunicación [Compredes]. Es una práctica recomendable ya que al ser un medio inalámbrico, de no hacerlo sería muy simple capturar el tráfico que por ella circula y por tanto la captura, por personas no deseadas, de datos sensibles.

A lo largo del desarrollo de las redes Wi-Fi han ido surgiendo diferentes métodos de encriptación de las comunicaciones, evolución necesaria pues los distintos métodos han resultado ser vulnerables y ha sido necesario implementar algoritmos más seguros que solventaran los problemas de los anteriores. Estos, a su vez, van demandando más recursos de los equipos que los implementan por lo que la solución adoptada será siempre un compromiso entre rendimiento y seguridad. Los métodos estándar disponibles se detallan a continuación.

WEP: Al inicio de las redes Wi-Fi ya se vio que las redes inalámbricas tenían problemas de seguridad intrínsecos a su naturaleza. Por esta razón, dichas redes nacieron con la posibilidad de activar encriptación y accesos mediante claves, siendo WEP el primer método que se implementó. Las siglas WEP provienen del inglés Wired Equivalent Privacy (Privacidad equivalente al cable). Ya en el mismo nombre se observa cual era el objetivo de esta encriptación, dar a las redes inalámbricas la misma seguridad que existía en las redes cableadas. Sin embargo la implementación de este protocolo adolece de problemas de diseño, que hace que si un equipo se encuentra dentro del alcance de la red, pueda capturar los paquetes de esta, y con la suficiente cantidad de paquetes capturados se pueda averiguar

la clave de la red, y por tanto tener acceso a ella. El proceso de captación de la clave de la red se puede hacer con herramientas públicas gratuitas y tan solo tarda unos pocos minutos. WEP permite claves de diversas longitudes de bits, lo cual teóricamente aumenta su seguridad, pero en la práctica, y debidos a los problemas existentes en la implementación de este protocolo, la única repercusión de utilizar una clave más larga es que aumenta el tiempo necesario para averiguar la clave de la red, pero esta sigue siendo vulnerable.

Dentro de WEP se reconocen dos métodos de autenticación de usuarios: Open System y Shared Key.

El método denominado Open System no implementa realmente autenticación. El punto de acceso permitirá que se una cualquier cliente, aunque posteriormente se obligará a que toda comunicación de datos sea codificada según el algoritmo dictado por WEP.

Por el contrario Shared Key dicta que los clientes tendrán que utilizar su clave WEP para autenticarse con el punto de acceso y solo aquellos que tengan las credenciales correctas serán admitidos por el punto de acceso como clientes.

En la práctica es recomendable utilizar autenticación Shared Key, pues Open System no proporciona realmente una autenticación de los clientes, solo encriptación de las comunicaciones, y aunque sería suficiente para preservar los datos, expone al punto de acceso a ataques de denegación de servicio (DoS).

Este protocolo no implementa ninguna gestión de claves. La clave utilizada es compartida por el punto de acceso y todos los clientes y debe ser distribuida a estos manualmente. Una consecuencia de ello es que con tener acceso a un solo equipo, se tiene la clave que compromete a todos los de la red.

Actualmente, por los problemas descritos se aconseja utilizar algún otro método de los disponibles, recurriendo solo a este sistema si no existiera ninguna alternativa viable y procurando acompañarlo de algún otro protocolo de encriptación general como puede ser IPsec o SSL.

2.3.2 WPA

WPA: El protocolo de seguridad WPA (Wi-Fi Protected Access) [Compredes] fue desarrollado por la Wi-Fi Alliance como respuesta a los fallos de seguridad detectados en WEP. Sin embargo, la seguridad proporcionada por este nuevo protocolo, se demostró que podía ser rota si se capturaban los paquetes que intercambian el punto de acceso y el cliente durante el proceso de autenticación. Con esa información, si la clave es corta y sencilla, lo cual, aunque no debiera, suele ser lo más normal, se puede averiguar la clave y por tanto acceder a los datos de la red. También se detectaron puntos de inseguridad en el protocolo que, aunque a día de hoy no han sido explotados por herramientas públicas, no se descarta que aparezca el software necesario para aprovechar dicha vulnerabilidad.

WPA incorpora varios sistemas de autenticación y encriptación que aportan seguridad extra, entre los que cabe destacar:

- **TKIP:** Siglas de Temporal Key Integrity Protocol, se basa en un sistema de verificación de integridad del paquete, es decir, que este no ha sido alterado durante la transmisión, y el uso de una clave que varía durante la comunicación, con lo que se solucionan problemas de WAP, pues la clave variará en menor tiempo y número de paquetes de los que se necesitan para averiguarla, por lo que no se dispondrá de información suficiente para hacerlo, y aunque se llegara a obtener esta, ya no sería válida para la comunicación en curso, pues la clave habría cambiado.
- **AES:** Algoritmo de encriptación más seguro que TKIP, cuya implementación no es obligatoria en sistemas WPAv1. Como contrapartida a esta mayor seguridad, demanda una mayor capacidad de proceso por parte de los puntos de acceso y los clientes. No obstante debería ser elegido, si es posible, ante TKIP.
- **EAP:** Es un protocolo de autenticación y encriptación que va asociado al protocolo 802.1x y que, por tanto, trabaja en conjunción con servidores de autenticación tipo RADIUS. Hace años se encontraron problemas de seguridad en el protocolo EAP, lo que desencadenó en nuevas variantes que,

mediante el uso de protocolos de seguridad asociados, pretendían solventar los problemas descubiertos. De este proceso surgió lo que se llamó Extended EAP, con diferentes variantes. Es de destacar que alguna de estas variantes como EAP-LEAP tienen fallos conocidos y su seguridad puede ser evitada con herramientas públicas y gratuitas.

WPA2: Ante la detección de la existencia de una brecha en la seguridad del protocolo utilizado por WPA, la Wi-Fi Alliance desarrolló una segunda versión que corrige dicho problema. Esta segunda versión obliga a la implementación del protocolo de encriptación AES, siendo este de uso por defecto en la norma WPA2.

Los protocolos WPA permiten la autenticación mediante una clave compartida entre cliente y punto de acceso, o haciendo uso de mecanismos más elaborados mediante el uso de un servidor de credenciales. Originalmente ambos tipos de arquitectura no tenían un nombre normalizado, y solían recibir el nombre de WPA el que hacía uso de servidor centralizado y WPA-PSK el que hacía uso de clave compartida (que es el significado de PSK, Pre-shared Shared Key). Actualmente se ha normalizado el uso de los términos “personal” para el uso de clave compartida, y “Enterprise” a aquella que provee autenticación contra un servidor RADIUS mediante protocolo 802.1x.

Capítulo 3 Seguridad en Redes

3 Seguridad en Redes

En los sistemas se tiene un modelo estándar de seguridad [Red03a], **CID** (Confidencialidad, Integridad, Disponibilidad) o por sus siglas en inglés (**CIA**). Este modelo de seguridad promueve que los recursos del sistema (datos, programas, parámetros para control de acceso, llaves criptográficas, memoria, tiempo de CPU, etc.) a proteger mantengan los siguientes atributos: Confidencialidad, Integridad y Disponibilidad.

La confidencialidad implica que los recursos no sean leídos o adquiridos por personas no autorizadas. La integridad implica que los recursos no sean alterados por personas no autorizadas. La disponibilidad implica que los recursos deben estar disponibles a las personas autorizadas en el momento que se requiera.

La palabra *intrusión* [Webster94], se refiere a la acción de introducirse sin derecho a un sitio e intruso [Webster94] es la persona que realiza la intrusión. En los sistemas ambos conceptos son aplicados en el contexto de la seguridad. La intrusión puede ser definida [Mukkamala02] como un conjunto de acciones que intentan comprometer la *confidencialidad, integridad y la disponibilidad* de los recursos del sistema.

Los intrusos se valen de ciertos tipos ataques (ver Tabla 3.1) para ingresar a un sistema específico. La Figura 3.1, muestra un ataque común de red [CERT03] utilizado para conseguir una intrusión.

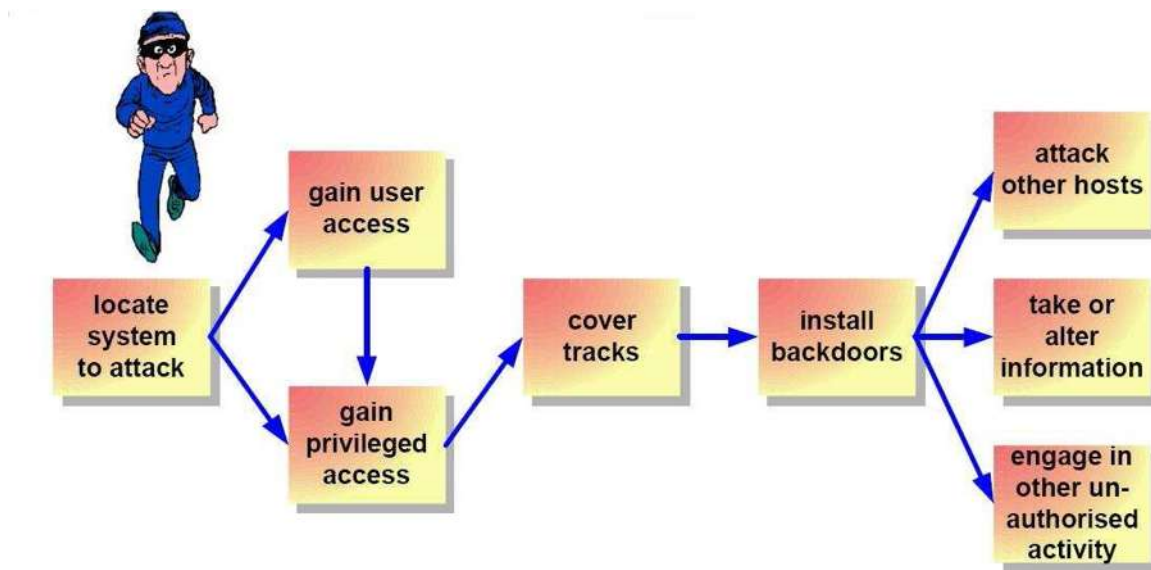


Figura 3.1 Ataque común de red

En la Figura 3.1, el atacante fija un sistema como meta para una intrusión. El ataque de tipo sondeo (ver sección 3.2) permite identificar vulnerabilidades o huecos que tienen asociados los sistemas. Una vez identificado el sistema y las vulnerabilidades o huecos, el atacante puede hacer uso de explotaciones (del Inglés exploits) de los ataques del tipo Usuario a Superusuario o Remoto a Local, para ingresar (intrusión) de una u otra forma al sistema. Una vez conseguida la intrusión y comprometida la cuenta del administrador, el intruso procede a eliminar toda evidencia originada de los ataques e instala puertas traseras (por ejemplo rootkit), las cuales le permitirán ocultar la intrusión y el libre acceso al sistema. Una Intrusión puede tener varios objetivos: atacar a otros sistemas, robar o alterar información, realizar actividades no autorizadas.

3.1 Mecanismos de Seguridad

A un alto nivel de abstracción, las personas emplean los sistemas para desarrollar sus tareas o actividades diarias mediante una entidad activa del sistema operativo denominada proceso. Los sistemas de cómputo tienen asociados diferentes recursos (archivos, directorios, tuberías, sistemas de archivos, sockets, etc.), los cuales son empleados por los procesos. En términos del sistema operativo el proceso y recurso se llaman *sujeto* y *objeto*, respectivamente.

Los sistemas deben contar con mecanismos de seguridad que permitan identificar usuarios, limitar el acceso a los recursos, y registrar acciones desarrolladas por los usuarios. Tales mecanismos deben aplicarse en el momento requerido, y deben contar con una protección que garantice su uso. Algunos ejemplos de tales mecanismos son:

- Autenticación.
- Auditoría.
- Control de acceso.

La pregunta obligada en la prueba de estos mecanismos, es el cómo determinar si un sistema que los emplea es seguro. Para responder esta pregunta algunas organizaciones forman equipos de especialistas que tienen la finalidad de obtener acceso no autorizado a los recursos, ya que la principal ventaja de un atacante es conocer y aplicar técnicas que le permiten descubrir y aprovechar vulnerabilidades que no son conocidas por quienes lo diseñaron, o administran un sistema.

Bob Toxen [Toxen00] identifica varias prácticas de inseguridad que son comúnmente realizadas por los administradores de sistemas. A tales prácticas Toxen les denomina los “siete pecados mortales”:

- Contraseñas débiles.
- Puertos de red abiertos.
- Versiones no actualizadas de software.
- Seguridad física pobre.
- CGIs (del Inglés Common Gateway Interface) inseguros.
- Cuentas viejas e innecesarias.

- Demora en actualizaciones por parte del administrador.

Un sistema debe ser administrado correctamente para minimizar el riesgo de las posibles amenazas que puedan afectar la confidencialidad, integridad y disponibilidad de los recursos. Las estructuras y los mecanismos del sistema operativo deben estar basados en un modelo formal de seguridad y para su implementación es conveniente seguir los criterios de diseño, implementación certificación y acreditación, establecidos por instituciones como el DoD con el “Libro Naranja” [dod85a], o como los criterios comunes de la ISO [iso08].

3.2 Tipos de Ataques

En el año de 1998 la DARPA (del Inglés Defense Advanced Research Projects Agency) [ofDefense05], lanzó una convocatoria para realizar evaluaciones en la detección de intrusión [Mukkamala02]. En esta convocatoria participaron varias empresas desarrolladoras de software e instituciones de educación. Para esta evaluación fue necesario utilizar ataques que permitiesen realizar una intrusión. Los ataques empleados se clasificaron en las siguientes categorías:

- DOS: (del Inglés *Denial of Service*) Negación de Servicio.

Estos ataques intentan mantener ocupado algún o algunos recursos del sistema, para impedir el manejo de solicitudes válidas.

- U2R: (del Inglés *User to Root*) Usuario a Superusuario.

Estos ataques intentan obtener el acceso a la cuenta del superusuario (administrador) desde una cuenta de un usuario normal.

- R2L: (del Inglés *Remote to Local*) Remoto a Local.

Estos ataques tienen la finalidad de explotar una vulnerabilidad que permita obtener el acceso a una cuenta de usuario de una máquina local desde una máquina remota.

- Sondeo: Búsqueda de Vulnerabilidades o Huecos

Los ataques de sondeo consisten en revisar una red de computadoras, en busca de vulnerabilidades y servicios disponibles.

La Tabla 3.1 muestra los tipos de ataques y sus correspondientes explotaciones conocidas.

Clase de ataque	Explotación	Clase de ataque	Explotación
Negación de Servicio	Apache 2 Arpoison Back Crashiis Dosnuke Land Mailbomb SYN Flood (Neptune) Ping of Dead (POD) Process Table selfping Smurf sshproccetable Syslogd tcpreset Teardrop Udpstorm	Remoto a Local	Dictionary ftpwrite Guest Httpunnel Imap Named ncftp netbus netcat phf ppmacro sendmail sshotrojan Xlock Xsnoop
Usuario a Superusuario	Anypw casesen Eject Ftbconfig Fdformat Loadmodule ntfsdos Perl Ps sechole Xterm Yaga	Sondeo	insidesniffer Ipsweep Is_domain Mscan NTinfoscan Nmap queso resetscan Saint Satan

Tabla 3.1 Lista de ataques difundidos por DARPA

Sin duda el arma más efectiva para cualquier atacante o intruso, es la habilidad de encontrar defectos de un sistema, que puede que no sean evidentes para quienes lo diseñaron o para quienes lo utilizan a diario. Se conocen varias formas para lograr introducirse sin permiso a un sistema.

La Figura 3.2 muestra los estudios de CERT (del Inglés *Computer Emergency*

Response Team) que evidencian el crecimiento de las vulnerabilidades. Esto trae como consecuencia la necesidad de contar con las actualizaciones necesarias de software.

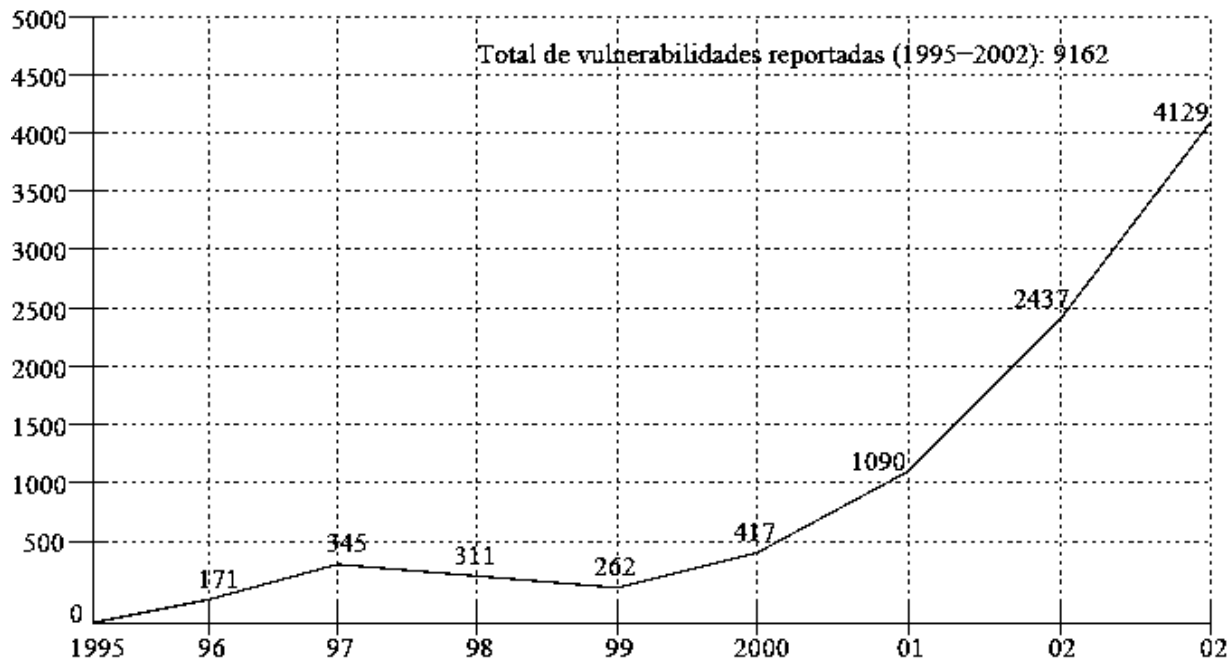


Figura 3.2 Vulnerabilidades reportadas por CERT desde 1995-2002

Diariamente se conocen nuevas vulnerabilidades. Para una referencia actualizada, se recomienda consultar la base de datos de ataques y explotaciones ofrecida por DARPA ([Lincoln Laboratory99], CERT [CERT05]).

3.3 Hechos Relevantes para la Seguridad en Sistemas

Algunos eventos han contribuido al nacimiento y crecimiento de la seguridad de sistemas. A continuación se listan algunos de los más importantes desde los años sesentas hasta los noventas, tomados de [Red03a].

3.3.1 *Los Años 60's*

Estudiantes del MIT (del Inglés Massachusetts Institute of Technology) miembros TMRC (del Inglés Tech Model Railroad Club), inventaron el término “hacker” bajo el contexto que hoy es conocido e iniciaron con la exploración y programación de un sistema de PDP-16.

A mediados de la década de 1960, el DoD (del Inglés Department of Defense) de Estados Unidos quería una red de comando y control que pudiera sobrevivir a una guerra nuclear. Mediante ARPA (del Inglés Advanced Research Agency) lanzó una convocatoria y de esta surgió la ARPANet. Esto dio pausa para la creación de la red conocida en la actualidad como Internet.

El Sistema Operativo Unix fue diseñado a finales de los años sesenta y principios de los años setenta por Ken Thompson de AT&T. El núcleo de Unix fue escrito originalmente en lenguaje ensamblador y reescrito en 1973 en el lenguaje de programación C.

3.3.2 *Los Años 70's*

Bolt, Beranek y Newman, investigadores y desarrolladores contratados por el gobierno y la industria de Estados Unidos, desarrollaron el programa telnet para una extensión pública de ARPANet. Esto abrió las puertas al uso público de las redes de datos que estaba restringida al gobierno e investigadores académicos. Telnet, es el programa más inseguro para redes públicas, por no contar con la opción de cifrar información.

Steve Jobs y Steve Wozniak fundaron la compañía Apple Computer e iniciaron la venta de Computadoras Personales (PC).

Jim Ellis y Tom Truscott crearon USENET, un boletín informativo electrónico para la comunicación entre usuarios distantes. USENET pasó a ser uno de los foros más populares para el intercambio de ideas en computación, redes y por supuesto cracking.

En 1978 nació el lenguaje C, con la publicación de The C Programming Language por Brian Kernighan y Dennis Ritchie. Desde su nacimiento C, ha mostrado eficacia y potencia, ya que este lenguaje no está prácticamente asociado a ningún sistema operativo, ni a ninguna

máquina. Es por esta razón fundamental, que C es conocido como el lenguaje de programación de sistemas por excelencia.

3.3.3 *Los Años 80's*

La IBM desarrolló y vendió PCs basadas en el microprocesador 8086 de Intel, la cual era una arquitectura barata para el cómputo de oficina y casero.

El Protocolo de Control en la Transmisión (TCP), fue desarrollado por Vint Cerf, el cual fue dividido en dos partes. Una parte fue nombrada como el protocolo de Internet. De la combinación de ambas partes se obtuvo el protocolo TCP/IP, el cual es en la actualidad un estándar para toda la comunidad en Internet.

La Legión de Doom y Chaos Computer Club son dos grupos pioneros de hackers que iniciaron con la explotación de vulnerabilidades en computadoras y redes.

En 1986 el acto de fraude y abuso de cómputo fue anexado a las leyes de Estados Unidos por su congreso. Estas leyes se basaron en la explotación de Ian Murphy, también conocido como el capitán Zap, el cual ingreso a computadoras militares, hizo varios pedidos a algunas compañías y realizó llamadas desde el sistema telefónico del gobierno.

El CERT fue creado para alertar a los usuarios sobre incidentes⁷ de seguridad.

Probablemente la mayor violación de seguridad de todos los tiempos [Tanenbaum97] sucedió la noche del 2 de Noviembre de 1988, cuando un estudiante de Cornell, llamado Robert Tappan Morris, disemino el programa de un gusano⁸ (del Inglés Worm) que inhabilito a miles de computadoras en todo el mundo. Para reproducirse el gusano explotaba una vulnerabilidad que le permitía obtener las claves de los usuarios y para su propagación utiliza a las tablas de ruteo. Morris no tuvo mucho que investigar, pues obtuvo información confidencial de su padre, un experto de seguridad de la NSA de Estados Unidos. Morris fue sentenciado a pagar una multa por 10,000 dólares, 3 años de libertad condicional y 400 horas de servicio a la comunidad. Se estima que los gastos del juicio excedieron los 150,000 dólares.

3.3.4 *Los Años 90's*

La ARPANet es decomisada y su tráfico fue transferido a Internet. Linus Torvalds desarrolló el kernel de Linux para ser usado como sistema operativo de GNU. Linux es más popular entre los hackers, quienes lo encuentran útil para construir sistemas seguros.

Vladimir Levin y sus cómplices transfirieron 10 millones de dólares en varias cuentas de la base de datos central de CityBank. Levin fue arrestado y se dice que todo el dinero fue recuperado.

Un estudiante Israelita de 19 años fue arrestado y condenado por ingresar a los sistemas del gobierno de Estados Unidos, durante la guerra del Golfo Pérsico. Los oficiales militares llamaron a este suceso como el más organizado y sistemático ataque ocurrido en la historia de los sistemas del gobierno.

El control de los satélites de Inglaterra fue tomado por varios delincuentes, finalmente el gobierno retomó el control.

3.3.5 *La Seguridad en la Actualidad.*

En febrero del 2000 el ataque DDOS (del Inglés Distributed Denial of Service) fue desencadenado en varios de los sitios más frecuentados en Internet. El ataque afectó a yahoo.com, cnn.com, amazon.com, fbi.gov y otros sitios dejándolos completamente inalcanzables para los usuarios normales.

En la actualidad se estima que 400 millones de personas usan o hacen uso de Internet al mismo tiempo. En el primer cuarto del 2002, el número de incidentes reportados por CERT llegó a 73,359 de 52,658 reportados en el 2001.

El incidente más recientemente ocurrido es el robo de datos de 40 millones de tarjetas de crédito [mas05]. Dicho incidente fue reportado al FBI el 22 de mayo del 2005 por Visa y MasterCard. La empresa MasterCard culpa a CardSystems Solutions, ya que la falla se atribuye al procesador. MasterCard asegura que los datos sensibles (número de seguro social, fecha de nacimiento, etc.) de los clientes afectados no están incluidos en el robo.

Estudios realizados por CERT [Allen00], indican que el crecimiento de Internet ha generado un incremento de incidentes de seguridad (ver Figura 3.3).

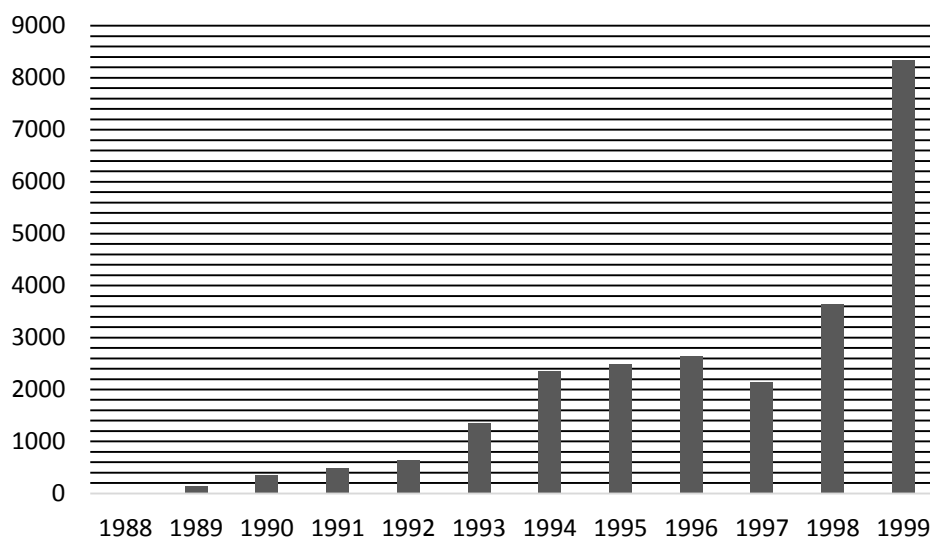


Figura 3.3 Crecimiento de incidentes en Internet según CERT 1988-1999

Los últimos estudios realizados por CERT [CERT03], revelan que del año 1998-2002 se registraron 182463 incidentes y con tendencia al aumento (ver Figura 3.4).

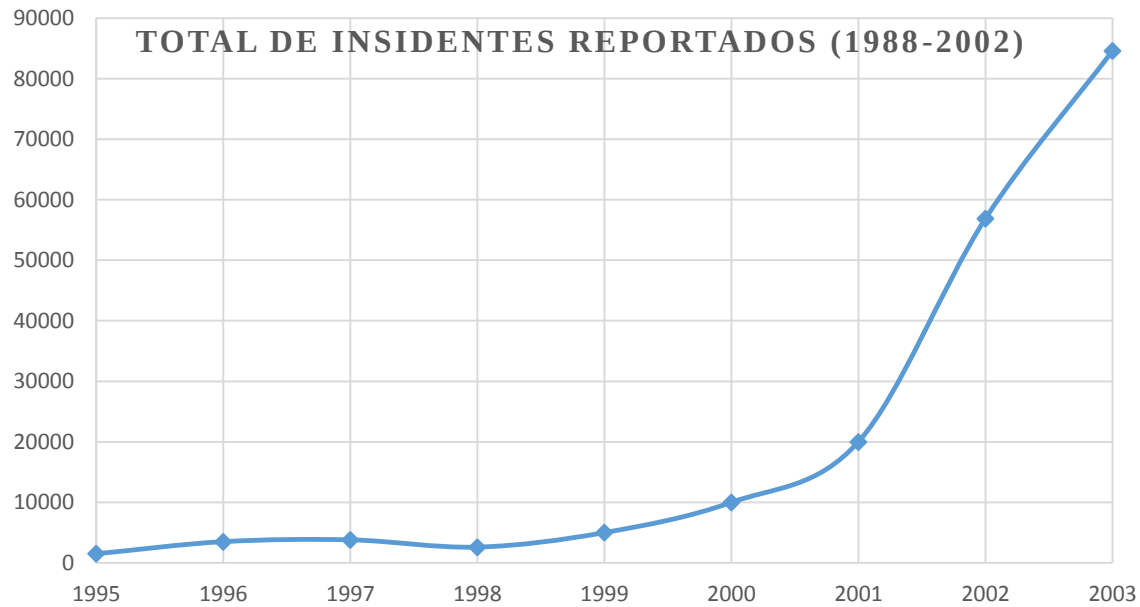


Figura 3.4: Crecimiento de incidentes en Internet según CERT 1995-2002

Los estudios de CERT [CERT03], indican que el grado de conocimientos técnicos para los atacantes es muy bajo, ya que en la actualidad existe un gran número de herramientas desarrolladas (ver Figura 3.5) para este fin. Lo cual permite obtener una alta complejidad. El contar con los medios de protección adecuados pudiera garantizar que los intrusos, atacantes o gente maliciosa no ingresen al sistema, evitando la observación, alteración o eliminación de información. Una librería disponible para ataque es libnet [Schiffman00], la cual ofrece la facilidad de crear e inyectar paquetes a la red. Los parámetros de los paquetes pueden ser manejados a conveniencia o gusto, lo cual permite crear explotaciones. Algunas herramientas que hacen uso de esta librería son: Nemesis [nem05] y nmap [nma05].

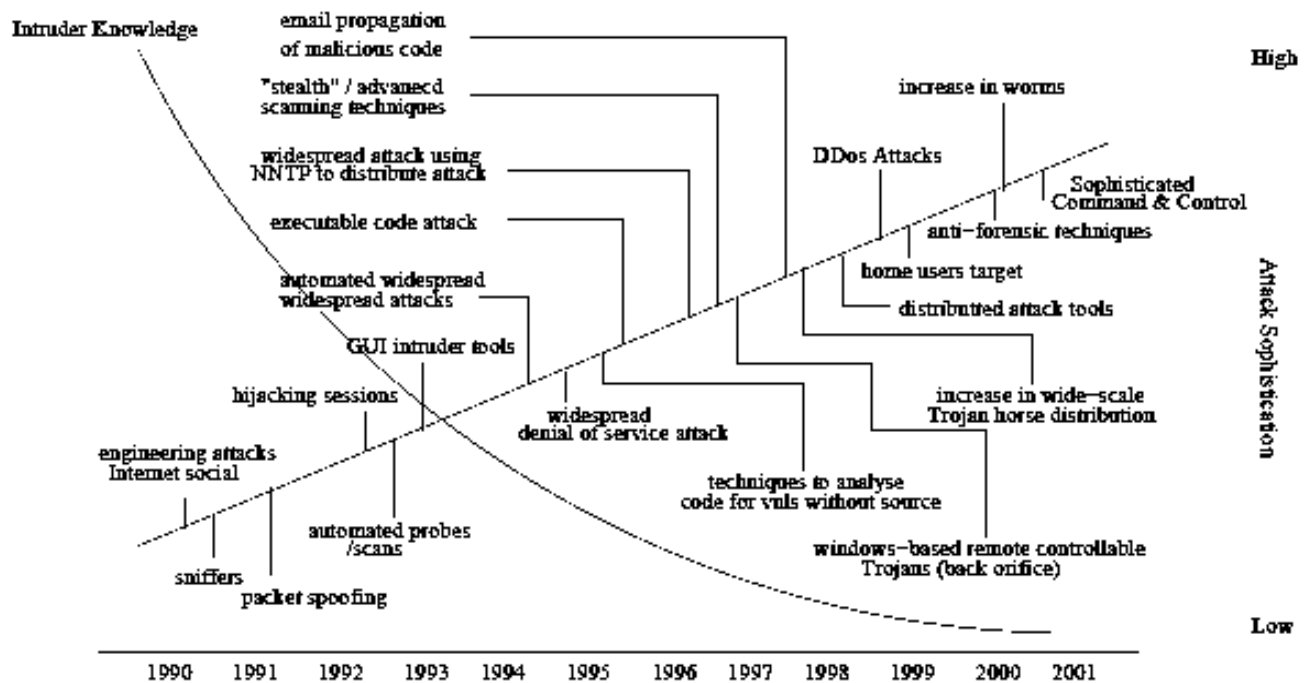


Figura 3.5 Nivel de conocimientos de atacantes y complejidad de los ataques según CERT

Con el crecimiento de la Internet y el aumento de vulnerabilidades e incidentes, la seguridad aparece en el horizonte como un problema potencial de grandes proporciones. Es por esta razón que se requiere de nuevas alternativas para salvaguardar los recursos de los sistemas. Son varios proyectos que actualmente se están desarrollando y que darán notables ventajas sobre los ya existentes; tal es el caso de SELinux10 [sel05], LIDS11 [lid05], entre otros. Dichos proyectos intentan ofrecer una nueva arquitectura de sistema, que elimine los problemas de un gran número de arquitecturas actuales.

3.4 Autenticación

Todo sistema debe contar con la capacidad de identificación de usuarios y crear para cada usuario válido una sesión de operación que le permita interactuar con los recursos del sistema. El mecanismo de autenticación es la primera barrera técnica de seguridad y tiene como finalidad verificar la identidad de los usuarios. La autenticación puede estar basada en uno o más de los siguientes factores:

- Algo que se conoce.
- Algo que se tiene.
- Algo que se es o Biométrico.

Factor *algo que se conoce*

El factor “algo que se conoce” se convierte en el secreto compartido entre usuario y sistema; tal secreto puede consistir de un número, palabra o frase. El sistema debe poseer al secreto compartido, aunque la experiencia ha demostrado que se corre un riesgo al tenerlo registrado en un medio de almacenamiento. Para esto se emplean los algoritmos de resumen (md5 [Group92], sha [ofStandards93], etc.), los cuales generan una cadena de bits correspondiente a la huella o resumen de la frase secreta. Este factor es tradicionalmente empleado debido a su fácil implementación y bajo costo. Al emplear este factor se deben eliminar los ataques de adivinanza, fijando un límite de intentos de autenticación, además de no informar la invalidez de usuarios.

Factor *algo que se tiene*

El factor “algo que se tiene” posee un bajo nivel de fiabilidad debido al empleo de partes físicas susceptibles a extravío o robo, por tales razones debe ser combinado con los factores “algo que se conoce” o “biométrico”. Las tarjetas magnéticas, generadores de números aleatorios son algunos ejemplos para el factor algo que se tiene.

Factor *algo que se es Biométrico*

El factor biométrico considera que los usuarios pueden reproducir sus características físicas de una manera repetitiva y precisa. Para este factor existen en el mercado una amplia gama de lectores biométricos de diversos tipos y marcas. Se pueden encontrar sensores de huella

dactilar, retina, iris, contorno de mano, etc., los cuales pueden realizar búsquedas de uno a muchos o uno a uno.

El número de factores empleado por el mecanismo de autenticación se basa normalmente en la sensibilidad de la información contenida en el sistema. Un mecanismo de autenticación robusto combina dos o más factores, lo cual implica un costo alto de implementación y mantenimiento. En general proteger información sensible implica una mayor inversión de recursos (tiempo y dinero) en la implementación de los mecanismos de seguridad, aunado a un posible aumento de complejidad para la operación y administración. Algunos trabajos sobre mecanismos biométricos, administradores de contraseña, y tarjetas magnéticas pueden encontrarse en [Lucas Ballard06, Sonia Chiasson06, Drimer07, Valentin Sgarciu06].

3.5 Mecanismo de Auditoría

El mecanismo de auditoría tiene la meta de registrar diversos eventos realizados por los usuarios y el propio sistema. Los diseñadores, programadores, y administradores necesitan analizar los registros de auditoría para intentar resolver problemas de seguridad.

La identidad, acción, y el tiempo son los aspectos mínimos a registrar. Los registros de eventos deben contener la información necesaria para la:

- Reconstrucción cronológica de eventos.
- Detección de eventos no autorizados.
- Identificación de problemas de configuración.

Un sistema puede registrar un gran número de actividades y para evitar la carencia de espacio en medios de almacenamiento los registros de auditoría deben limitarse, además de ser necesario una constante depuración. Más información detallada sobre los registros de auditoría se puede encontrar en [Swanson96].

Aspecto legal

Los sistemas auditores pueden tener la capacidad de registrar eventos locales y/o remotos.

Las organizaciones deben contar con el respaldo legal que les permita registrar, analizar y sancionar a los miembros que incurran en actividades no permitidas, intentado con esto no afectar la estabilidad de la organización. En nuestro país existen algunos intentos de abordar a la legislación informática, de tal manera que el código penal federal en su artículo 211 de los incisos 1 al 7 [cpf09] abordan el acceso ilícito a sistemas y equipos de informática.

Algunos estados por su parte han realizado esfuerzos para cubrir a la legislación informática, como ejemplo citamos al código penal del Distrito Federal [cpd07], código penal del estado de Sinaloa [cps06], y la Ley de protección de datos del estado de Colima [lpd03]. En nuestro país es necesario un mayor esfuerzo para regular los delitos informáticos y de manera global también es necesario integrar la reglamentación informática que involucre a todos los países del mundo, ya que las comunicaciones electrónicas han venido a modificar a la interacción entre las naciones y sus pobladores.

3.6 Mecanismo de Control de Acceso

Un sistema cuenta con un conjunto de sujetos o usuarios y un conjunto de objetos o recursos. Esto implica que se debe regular la interacción entre los elementos de tales conjuntos, y para ello el mecanismo de control de acceso tiene la finalidad de limitar la interacción entre sujetos y objetos. La autorización es parte del control de acceso y su función es otorgar o negar el acceso a un objeto por parte de la acción de un sujeto. La Figura 3.6, muestra un sujeto s_1 que pretende tener acceso de lectura sobre el objeto o_1 y dependiendo de las reglas de la autorización el acceso puede ser otorgado o negado.

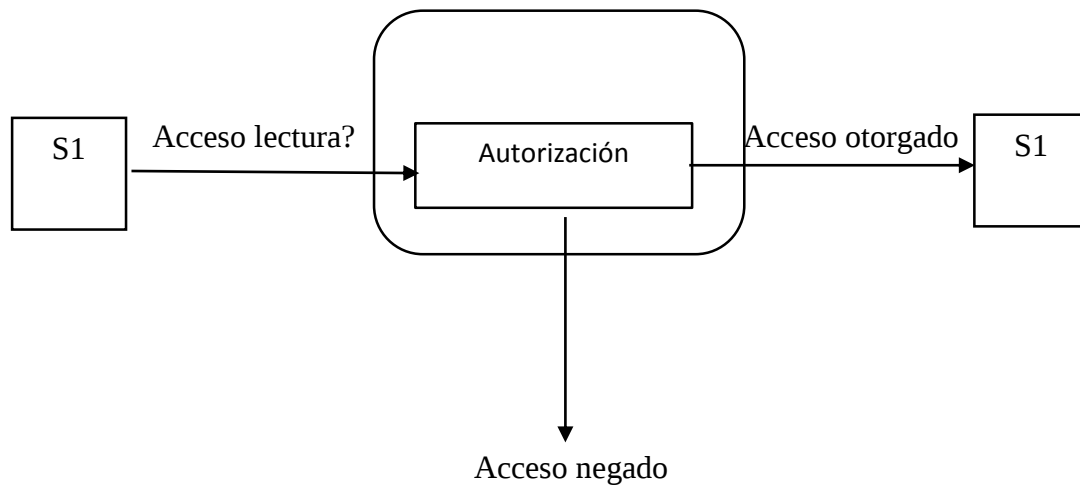


Figura 3.6 Mecanismo de control de acceso

Morrie Gasser [Gasser88] cita tres tareas fundamentales para el mecanismo de control de acceso:

- *Autorización*. Determina si un sujeto tiene el privilegio de tener acceso sobre un objeto.
- *Determinar los derechos de acceso*. Los derechos de acceso se obtienen de la combinación de los modos de acceso tales como lectura, escritura, ejecución, etc.
- *Cumplimiento de los derechos de acceso*.

El control de acceso establece que los sujetos y objetos deben tener asociado un conjunto de atributos de seguridad que garantice la validez de acceso. La determinación de acceso a un objeto se fundamenta en tales atributos y en una política de seguridad que consiste de un conjunto de reglas basadas en uno o más modelos de control de acceso.

Capítulo 4 Herramientas de Seguridad

Herramientas de Seguridad

La seguridad informática es el área que se enfoca en la protección de la infraestructura de los sistemas y todo lo relacionado con este, y especialmente, la información contenida o circulante. Para ello existen una serie de estándares, protocolos, métodos, reglas, herramientas y leyes concebidas para minimizar los posibles riesgos a la infraestructura o a la información.

Utilizar técnicas de desarrollo que cumplan con los criterios de seguridad al uso para todo el software que se implante en los sistemas, partiendo de estándares y de personal suficientemente formado y concienciado con la seguridad.

4.1 Scanners

El escaneo de puertos [Smith1ed], es el proceso de consulta de puertos para descubrir si los servicios están escuchando sobre ellos; esto es útil para las cookies, ya que rápidamente les permite ver los posibles puntos de entrada en el sistema. Los escáneres de puertos son cada vez más inteligentes, y en lugar de limitarse a informar si un puerto está abierto o cerrado, el nombre y la versión del servicio que se ejecuta en el puerto, y el sistema operativo que se ejecuta el servicio, lo cual ofrece mayor información a un atacante potencial.

Análisis de vulnerabilidad es una extensión de escaneo de puertos, e implica el intento de descubrir si alguno de los servicios que se encuentran a la escucha en un puerto TCP / UDP son explotables. Los escáneres de vulnerabilidades suelen utilizar una mezcla de técnicas pasivas, y técnicas de ataque. Una vez más, esto hace la vida más fácil para el atacante, que puede encontrar una lista de debilidades en el sistema que está escaneando.

Como se puede ver, por lo tanto, es importante utilizar este tipo de herramientas en contra de nuestra propia red, antes de que alguien con intenciones menos honorables lo haga.

Al escanear, es importante recordar que los resultados serán diferentes dependiendo de en qué parte de la red se realiza la exploración. El mayor peligro proviene de un ataque

externo, por lo que para escanear desde esta perspectiva, es necesario una máquina externa de preferencia con acceso root. En su defecto, un shell de la cuenta de una empresa de alojamiento remota por lo menos le dará la oportunidad de realizar la exploración basada en consola como usuario sin privilegios; pero hay que tener en cuenta que muchos proveedores de Internet bloquean ciertos puertos (por ejemplo, NetBIOS), por lo que estos resultados pueden no ser del todo exacto.

Dependiendo de la naturaleza de su organización, un ataque interno también puede ser una amenaza, y el escaneo debe realizarse desde la perspectiva de un usuario de estación de trabajo típico en el firewall y DMZ. Por el contrario, es importante analizar el firewall y LAN interna desde el punto de vista de la zona de distensión: prevemos la posibilidad de una entrada atacante obtener de la zona de distensión, y una de sus primeras acciones será la de realizar un escaneo de puertos del resto de la red, con la esperanza de eludir la seguridad de las fronteras estrictas.

Escaneando Puertos con Nmap

Nmap ha logrado gran éxito en los últimos años, llegando a ser fácilmente una de las herramientas de auditoría de seguridad más populares en todo; de hecho, la palabra en sí se ha convertido en sinónimo de escaneo de puertos, con gente regularmente se habla de "nmapping" o de haber sido "nmapped". ¿Qué es tan especial acerca de Nmap?. En el pasado, el escaneo de puertos ha sido algo de un arte pasado por alto? Aunque una variedad de técnicas de exploración se han conocido por un largo tiempo, la mayoría de los escaneos de puertos sólo implementan estándar de conexión TCP de exploración, quizás considerando el área demasiado trivial como para justificar el desarrollo grave. Nmap cambió todo eso, la combinación de una serie de técnicas nuevas y ya establecidas, dando finalmente el fascinante tema del escaneo de puertos, la atención que merecía.

Vulnerabilidad de Auditoría con Nessus

Nessus es algo más que una base de datos, sin embargo, es un escáner inteligente que puede reconocer los servicios que se ejecutan en puertos no estándar, y no cree ciegamente la versión del demonio informa en su bandera. Nessus también cuenta con una arquitectura de plug-in para las pruebas de vulnerabilidades, lo que significa que las nuevas pruebas se pueden agregar fácilmente sin tener que actualizar el escáner.

Nessus utiliza una arquitectura cliente / servidor, permitiendo a los usuarios a través de la red para iniciar sesión en el demonio y llevar a cabo auditorías de seguridad; Sin embargo, el escenario más común es que el cliente y el servidor, tanto se ejecutan en el mismo sistema. El demonio de Nessus se lanza al emitir el comando `nessusd` (como root). El demonio es compatible con las siguientes opciones:

-A <dirección>: Sólo escuchar conexiones en la dirección IP dada. Esto es útil si la máquina que ejecuta el demonio de Nessus tiene varias interfaces, o si desea desactivar el acceso remoto todos juntos especificando 127.0.0.1 como la dirección.

-P <número>: El puerto TCP en el que el demonio debe escuchar. De forma predeterminada es 1241.

-D: Enviar el demonio en el fondo. Alternativamente, un símbolo de unión puede ser utilizado para lograr esto: `y nessusd`.

Características del cliente

Una vez que el demonio está en ejecución, puede iniciar el cliente Nessus con el comando `nessus`, y se puede iniciar sesión con el nombre de usuario y la contraseña creada con el `nessus-adduser` utilidad. Vamos a suponer que usted está utilizando el cliente gráfico.

Plug-Ins

Los Plug-ins le proporciona un control preciso sobre el que se incluyen pruebas de vulnerabilidad en la exploración. A menos que necesite limitar el escaneo de una mayor velocidad, se podría pensar que por lo general es mejor simplemente elegir Habilitar todo. Tenga en cuenta, sin embargo, que esto incluye algunos controles que pueden causar que el

sistema se bloquee sondeado; así habilitar todo. Pero Plug-ins (peligrosas) suele ser la opción más segura para los servidores de producción. Si es posible, sin embargo, hacer una cita cuando se puede realizar un análisis completo; aparte de las vulnerabilidades adicionales que puede descubrir, también es importante saber si (las cookies) causarán sus máquinas se bloquee si se ejecutan exploraciones con Nessus en ellos.

Plug-Ins Preferencias

Las Preferencias Plug-ins da el control sobre casi todos los aspectos de la exploración. Nessus utiliza Nmap para escaneo de puertos, por lo que la primera opción es el rango de modo de escaneo y el puerto a utilizar.

Una gran cantidad de espacio está dedicado a las preferencias del módulo a la configuración de las sondas HTTP.

Ataques de fuerza bruta de acceso, aunque útil para eliminar las contraseñas débiles, que pueden aumentar la longitud de la exploración considerablemente. Servicios tales como Telnet, FTP, POP3, e IMAP (Internet Message Access Protocol) suelen autenticar usuarios a través de /etc/passwd, por lo que suponiendo que tiene acceso de root en el sistema que se está digitalizado, agrietamiento este archivo con una herramienta como John The Ripper es una opción mucho más rápido que el de fuerza bruta usando el demonio Nessus. Del mismo modo, crackers de contraseñas se pueden encontrar generalmente en aplicaciones como ICQ y htpasswd, archivos que pueden ser quebrados por John lo que podría tener otras mejores opciones para la comprobación de contraseñas de Nessus.

Una de las mejores características de Nessus es que proporciona información no sólo sobre la vulnerabilidad, sino también sobre la forma de solucionarlo. El BID valor se refiere a la vulnerabilidad de Bugtraq ID, una base de datos en línea de los errores de acceso a <http://www.securityfocus.com/bid/bugtraqid/>. La base de datos ofrece información detallada sobre las plataformas de afectados, el código de explotación (en su caso), y las soluciones al problema.

Por último, Nessus ofrece una variedad de formatos en los que desea guardar el informe, que van desde texto ASCII a HTML completo con gráficos de sectores que muestra estadísticas sobre los anfitriones perfectos para crear presentaciones escaneada.

Auditoría con Nikto

Los servidores web son uno de los servicios más comúnmente explotadas, en parte debido a su potencial complejidad de la configuración, sino también por el número de secuencias de comandos de terceros utilizados en ellos. Considerando que el servidor Apache tiene un historial bastante bueno de seguridad, muchos CGI (Common Gateway Interface) scripts no lo hacen. Escáneres especializados también están disponibles para prácticamente todos los servicios populares (tales como MySQL, Sendmail y FTP).

Nikto es un impresionante escáner servidor Web, las pruebas de más de 2.600 vulnerabilidades en 625 servidores, y es una herramienta importante para cualquier administración de un sitio Web. Basado en el escáner popular Whiskers, Nikto sigue utilizando el módulo LibWhiskers Perl; pero aunque Whiskers ya no existe, Nikto continúa prosperando. Las características incluyen:

- Plug-in de apoyo, lo que permite a los usuarios añadir análisis personalizados
- Usuario adivinando
- IDS (Intrusion Detection System) la evasión
- SSL (Secure Sockets Layer) de apoyo
- El soporte de proxy
- Salida en texto plano, HTML o CSV (valores separados por comas) servidor Web para localizar los puertos no estándar
- Un gran número de comprobaciones de vulnerabilidad
- Nikto requiere dos módulos Perl, Net: SSLeay (que a su vez requiere OpenSSL) y LibWhiskers.

Una de las características más interesantes de Nikto es sus IDS (Intrusion Detection System) técnicas de evasión, los cuales pueden ser utilizados para realizar escaneos sigilosos que podrían no ser captadas por muchos IDS. Esta es una característica común de este tipo de herramientas (incluyendo Nmap y Nessus), y es uno que debe tener en cuenta, ya que puede permitir a un atacante para escanear su red sin ser detectado.

4.2 Sniffers

Un sniffer es un programa para monitorizar y analizar el tráfico en una red de computadoras, detectando los cuellos de botella y problemas que existan. También puede ser utilizado para "captar", lícitamente o no, los datos que son transmitidos en la red.

Tcpdump

Tcpdump muestra las cabeceras de los paquetes que captura de un interfaz de red dado y que cumplen la expresión pasada al ejecutar [Wikisnif], es decir, te permite monitorizar tráfico de red en tiempo real.

Los filtros que se pueden crear para mostrar tan sólo la información que nos interesa, hacen de tcpdump una herramienta muy potente para el análisis de tráfico en redes de comunicaciones. Tcpdump es una aplicación “peligrosa”, por lo que en los sistemas UNÍX sólo se permite su utilización al usuario root. Luego, tcpdump permite examinar todas las conversaciones, incluyendo mensajes de difusión SMB y NMB. Mientras que sus capacidades en detección de errores están principalmente a nivel de capa de red OSI, todavía puedes usar su salida para obtener una idea general de qué están intentando hacer el servidor y el cliente. En Windows, en lugar del tcpdump se usa el Windump.

Wireshark

Es un potente analizador libre de protocolos de redes, funciona bajo Unix, Mac OS X y Windows [Wikisnif]. Nos permite capturar los datos directamente de una red u obtener la información a partir de una captura en disco (puede leer más de 20 tipos de formato

distintos). Destaca también por su impresionante soporte de más de 300 protocolos, gracias sin duda a la licencia GPL y sus más de 200 colaboradores de todo el mundo.

Kismet

Es un sniffer específico a Linux para redes inalámbricas [Wikisnif]. Específicamente, es un detector de la red 802.11, un succionador, y un sistema sin hilos para la detección de la intrusión. Funciona correctamente con los dos principales tipos de tarjetas inalámbricas, es decir, trabajará con cualquier tarjeta de red sin hilos que apoye modo de supervisión crudo (rfmon) y puede "oler" 802.11b, 802.11a y el tráfico 802.11g.

Kismet identifica redes de modo pasivo, recogiendo paquetes y detecta redes nombradas estándares, redes ocultas e infiere la presencia de redes sin balanceamiento vía tráfico de los datos.

4.3 Rompedores de Contraseñas

Aircrack

Es un programa crakeador de claves WEP y WPA que es capaz de recuperar las contraseñas una vez que haya conseguido suficiente paquetes de datos [Aircrack]. Implementa el ataque estándar FMS junto con algunas optimizaciones como los ataques Korek, así como todos los nuevos ataques PTW y como consecuencia obtiene un resultado de ataque mucho más rápido comparado con otras herramientas de crakeo, de hecho Aircrack es un conjunto de herramientas para auditar redes inalámbricas.

De manera general, las herramientas que se incluyen en Aircrack se dividen en varias categorías, recolección de información, ataques sobre dicha información, aceleración de la obtención de información. Estas son las tres herramientas más destacadas y más conocidas

de la suite aircrack-ng, sin embargo, no son las únicas y además se incluyen en dicha suite otra gran cantidad de herramientas útiles en estas y otras tareas que permitirán auditar tu red de manera mucho más eficiente.

John the Ripper

Es un programa de criptografía que aplica fuerza abierta para descifrar contraseñas. Es capaz de romper varios algoritmos de cifrado o hash, como DES, SHA-1 y otros [Wikijonh]. Es una herramienta de seguridad muy popular, ya que permite a los administradores de sistemas comprobar que las contraseñas de los usuarios son suficientemente buenas.

John the Ripper es capaz de autodetectar el tipo de cifrado de entre muchos disponibles, y se puede personalizar su algoritmo de prueba de contraseñas. Eso ha hecho que sea uno de los más usados en este campo.

John the Ripper usa un ataque por diccionario: tiene un diccionario con palabras, que pueden ser contraseñas típicas, y las va probando todas. Para cada palabra, la cifra y la compara con el hash a descifrar. Si coinciden, es que la palabra era la correcta.

Esto funciona bien porque la mayor parte de las contraseñas que usa la gente son palabras de diccionario. Pero John the Ripper también prueba con variaciones de estas palabras: les añade números, signos, mayúsculas y minúsculas, cambia letras, combina palabras, etc.

Además ofrece el típico sistema de fuerza bruta en el que se prueban todas las combinaciones posibles, sean palabras o no. Éste es el sistema más lento, y usado sólo en casos concretos, dado que los sistemas anteriores (el ataque por diccionario) ya permiten descubrir muy rápidamente las contraseñas débiles.

4.4 Ataques a Red

Según [Smith1ed], muchos de los servicios en el pasado, tales como el *Berkley R * Suite* (`rlogin`, `rsh`, y así sucesivamente), se basaban exclusivamente en el nombre de usuario y la dirección IP del cliente como un medio de autenticación; si su dirección aparece en la otra máquina. `Rhosts` archivos, puede acceder a la máquina sin necesidad de suministrar un nombre de usuario y contraseña. Desde el punto de vista de seguridad, esto no es bueno. Hoy SSH (Secure Shell), SCP (copia segura) y SFTP (Secure File Transfer Protocol) proporcionan alternativas más seguras, cifradas a R (Suite de Berkley), pero algunos protocolos de autenticación basados en host son todavía de uso común. Es posible argumentar que esta es una moneda de dos caras, después de todo, si no se transmite ninguna contraseña, no hay ninguna amenaza de un atacante olfateando la conexión y, posteriormente, con la contraseña de sí mismo. A pesar de ello, el alejamiento de la autenticación basada en host a los servicios cifrados protegidos por contraseña, es considerado por casi todo el mundo para ser una buena cosa.

Denegación de Servicio (DOS)

En un sentido general, el ataque de *la denegación de servicio (DoS)* se puede considerar como cualquier ataque que intenta privar a los usuarios legítimos de un servicio ofrecido por el sistema o red por la sobrecarga de un recurso limitado como el ancho de banda, memoria, espacio en disco o CPU. El DoS más popular ataca alrededor de ancho de banda y, sobre todo cuando se distribuyen, han sido un gran problema en los últimos años, con muchos ataques de alto perfil contra compañías como Yahoo! y eBay.

La forma más simple de los ataques de denegación de servicio de ancho de banda de limitación es simplemente enviar más datos a una máquina que tiene los recursos para hacer frente a .Si todo el ancho de banda disponible o los recursos para el objetivo se puede utilizar encima, el tráfico legítimo no se puede procesar. Una forma primitiva de este ataque es la *inundación de ping*, donde el objetivo es bombardeado con ICMP (Internet Control Message Protocol) eco peticiones, obstruyendo el ancho de banda en ambas direcciones, y

poner una tensión en la pila TCP / IP del sistema, ya que los intentos de destino para responder a los pings. Muchos administradores están bajo la idea errónea de que el bloqueo de las solicitudes de ping entrantes en el servidor de seguridad va a resolver este problema. Aunque esto hace detener el flujo de los paquetes ICMP que salen de la red, el ancho de banda aguas abajo entre el ISP y el perímetro cortafuegos / router está siendo afectada, por lo que esta solución es sólo parcialmente efectiva. Por cierto, si aguas abajo (desde la perspectiva de la víctima) de ancho de banda está completamente saturado, este se detiene efectivamente todo el tráfico TCP aguas arriba también. Esto podría ir en contra de lo que el sentido común le dice-después de todo, es común el uso de una analogía de la carretera para describir la red de tráfico, pero la respuesta está en el hecho de que TCP es un protocolo fiable. Parte de la fiabilidad de TCP proviene del hecho de que todos los datos enviados deben ser reconocidos por el receptor a la llegada (de lo contrario, se asume perdido y retransmitido); en cualquier conexión TCP, los datos deben ser capaces de fluir en ambos sentidos. UDP e ICMP no ofrezca tales garantías, pero aun así son generalmente bidireccional en la naturaleza; Por lo tanto, estos protocolos también se verán afectados.

Inundaciones Ping es una batalla de ancho de banda, debido a que el atacante debe saturar la línea de la víctima en la medida que el tráfico legítimo no puede fluir de una manera eficiente; aun así, a menos que la pila TCP / IP de la víctima puede ser abrumado al responder a estos paquetes ICMP, paquetes legítimos todavía pasa. Ahora la red corporativa típica utiliza una línea arrendada (generalmente de al menos 4 MB) para la conexión a Internet, mientras que el usuario doméstico medio sólo tiene acceso a acceso telefónico o DSL / cable-a menudo con el tráfico de subida tope significativamente menor que aguas abajo. Esto presenta un problema para el atacante, porque las probabilidades de ancho de banda se apilan firmemente contra él. Una "solución" a lo que ha sido la de poner en marcha la inundación ping desde una máquina comprometida en una red rápida; el otro ha sido la introducción del ataque DoS distribuido (DDoS). Ataques DDoS-el siguiente paso lógico en la mesa de inundaciones-ping utilizar el ancho de banda disponible de muchas redes para operar. Si la idea de la recepción de 100 kbps de tráfico ICMP de una máquina comprometida te preocupa, imaginar lo que sucede cuando 10 máquinas más comienzan unirse en el ataque.

Aparte de las inundaciones ICMP (y no sólo tiene que ser ICMP de eco peticiones que se

utilizan), inundaciones UDP también es común. Al igual que con las inundaciones de ping, es ventajoso si el equipo de destino puede ser persuadido para responder a los datagramas UDP, por lo que los puertos que ejecutan los servicios basados en UDP como `chargen`, `eco` y `cita` comúnmente se dirigen. El tráfico ICMP y UDP se pueden suplantar fácilmente (es decir, la dirección de origen ha cambiado), que puede conducir a que el administrador nave suponiendo que miles de máquinas están participando en el ataque, y señalando con el dedo a partes inocentes.

Ataque de Ping-Pong

Además de ofrecer al atacante la capacidad de ocultar su origen, fuente de la falsificación de direcciones también abre las puertas para los llamados ataques de ping-pong, llamado así porque los paquetes rebotan como pelotas de ping-pong (figura 4.1).

Así es como funciona, paso a paso:

El atacante identifica dos máquinas a la vez que ejecutan un servicio UDP como `chargen` o `eco`.

El atacante envía datagramas UDP al puerto 7 (el puerto de la `eco`, en donde el daemon se ejecuta) de la máquina A, con los paquetes falsificados para demostrar la máquina B como dirección de origen, y el puerto UDP 7 como el puerto de origen.

El `eco` daemon en el equipo A recibe el datagrama, y el `eco` es de nuevo a lo que piensa que es el remitente (equipo B).

Máquina B recibe el datagrama en su `eco` demonio, y el `eco` es de nuevo a la máquina A.

Este proceso continúa hasta el infinito, hasta que una máquina se cuelga, o comienza a caer los datagramas.

El ataque ping pong tiene el potencial de paralizar la red entre las dos máquinas durante mucho tiempo, porque una vez que se inicia el ataque, que no requiere mayor intervención por parte del atacante. Desactivar los servicios innecesarios, como `eco` y `chargen` puede eliminar el potencial de este tipo de ataque.

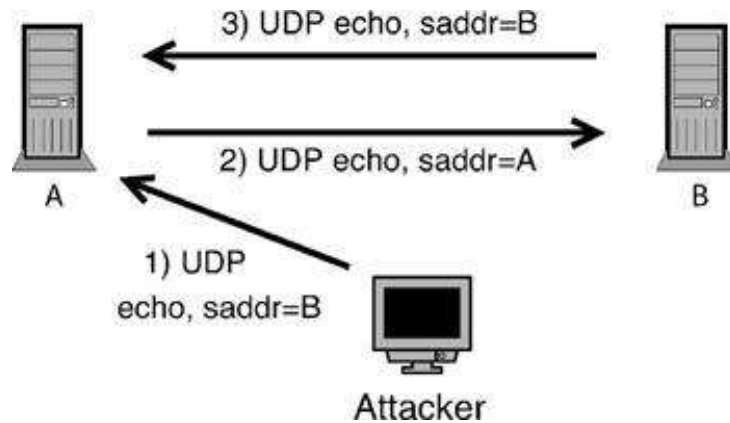


Figura 4.1 El ataque de ping pong en la acción.

Inundación de Redes Distribuidas

Con sólo un puñado de máquinas comprometidas que participan en una inundación de ping distribuida, el atacante podría fácilmente dentro y manualmente iniciar el ataque a cada uno, pero cuando el número de zombies (como se les llama comúnmente) se eleva a unos pocos cientos, esto se convierte en poco práctico. En 1999 vio un aumento dramático en los ataques DDoS como el resultado de la liberación de dos herramientas desarrolladas para coordinar esas redes comprometidas. Aunque estas dos herramientas (nombrados Trinoo y TFN) ven primitiva para los estándares de hoy en día y se han convertido en gran medida obsoleta, que vale la pena revisar porque constituir la base para muchos agentes DDoS más recientes.

El [Dittrich99] Trinoo red consta de dos partes ‘masters’ y ‘demonios’ todos los que se ejecutan en las máquinas que han sido comprometidos previamente por otros medios. Cada maestro controla muchos demonios, con el atacante controla cada maestro (figura 4.2).

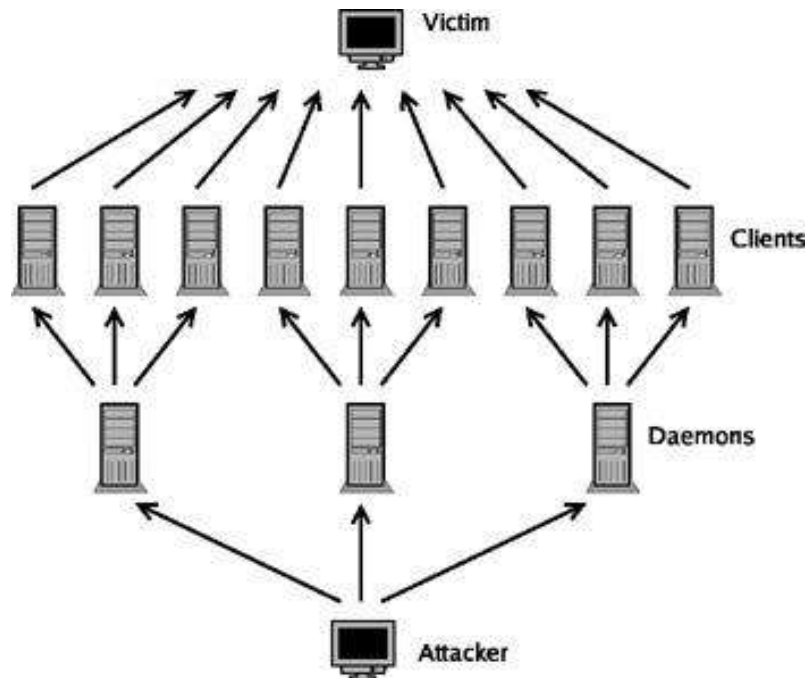


Figura 4.2 Arquitectura de la red de inundaciones Trinoo.

Para lanzar un ataque, el cracker abre una conexión TCP (es decir, en TELNETs) a cada maestro en el puerto 27665, y entra su contraseña. Una vez conectado, puede ejecutar comandos para realizar tareas como la gestión de los demonios activos y el lanzamiento de ataques de denegación de servicio. El maestro entonces retransmite estos comandos para cada demonio.

El otro jugador grande, TFN (Flood Network Tribal), utiliza una arquitectura maestro-daemon similar, pero es más avanzado, que ofrece la capacidad de realizar ICMP, SYN, UDP e inundaciones, así como proporcionar una puerta trasera rootshell on-demand. A los TFN se comunican con los demonios a través de ICMP de eco paquetes de respuesta, el comando a ejecutar está llevando en el segmento de datos del paquete. Esto tiene fuertes ventajas sobre TCP o UDP de comunicación debido a que muchos administradores de red no pensarían para comprobar la carga útil de un paquete ICMP.

La liberación de TFN y Trinoo fue seguido de cerca por la aparición de otros agentes DDoS, que utiliza una arquitectura maestro-daemon similar. Los ejemplos notables incluyen TFN2K (una versión mejorada del TFN, que utiliza el cifrado y altera su nombre en la lista de procesos para evitar la detección), Shaft, Mstream y Stacheldraht ("alambre de

púas"), todos los cuales datan de alrededor 1999 a 2001. En tiempos más recientes, el enfoque parece haber sido en la creación de agentes DDoS para Windows, con el desarrollo de agentes de Linux relativamente tranquilos. Aunque es improbable que desaparezca el problema de la floodnets (inundación de redes) coordinados en Linux, sí parece estar disminuyendo en el momento.

Febrero de 2000 se produjo una serie de ataques DDoS contra los gigantes de Internet como *Amazon.com*, *eBay*, *ZDNet*, *CNN.com*, *buy.com*, *Datek*, *E * Trade* y *Yahoo!*, golpeándolos desconectado por completo durante varias horas a la vez. Bien orquestada, y carente de cualquier motivo obvio, los ataques fueron de una intensidad sin precedentes. Esto no parecía ser el trabajo de script kiddies, y los rumores comenzaron a volar sobre quién estaba detrás de los ataques. Teóricos de la conspiración culparon al gobierno de los EE.UU.: el gobierno de Clinton estaba presionando para aumentos de poderes de vigilancia electrónica y, los teóricos argumentaban, ¿qué mejor manera de demostrar la necesidad de que estos poderes que por asustar al público? Y otros afirmaron que el dinero estaba detrás de los ataques, al citar los cambios en los precios de las acciones de las empresas afectadas antes y después de los ataques.

El Ataque Smurf

Tal vez el más peligroso de los ataques de denegación de servicio de ancho de banda que consume es el Smurf, que primero ganó el reconocimiento en 1997 con el lanzamiento del código de prueba de concepto por Tfreak. Desde entonces, el Smurf-y es descendiente de la gran popularidad Fraggles han logrado con los script kiddies en todas partes. Para entender cómo funciona el Smurf, en primer lugar hay que tomar un pequeño desvío en las redes IP. IP introduce el concepto de una dirección de difusión, que se calcula mediante la aplicación de la máscara de subred para una dirección en la red; los datos destinados a esta dirección se envían a todos los hosts de la red. Para las clases A, redes B y C (tabla 4.1), es simplemente un caso de sustituir la sección de host con 255, como se muestra a continuación:

Clase	Ejemplo de dirección	Máscara de subred	Dirección de difusión
A	10.2.8.34	255.0.0.0	10.255.255.255
B	172.16.10.1	255.255.0.0	172.16.255.255
C	192.168.53.19	255.255.255.0	192.168.53.255

Tabla 4.1 Cálculo de direcciones de difusión

Las direcciones de difusión son útiles para fines de diagnóstico ping a una dirección de difusión muestra de un vistazo qué hosts están vivos en la red.

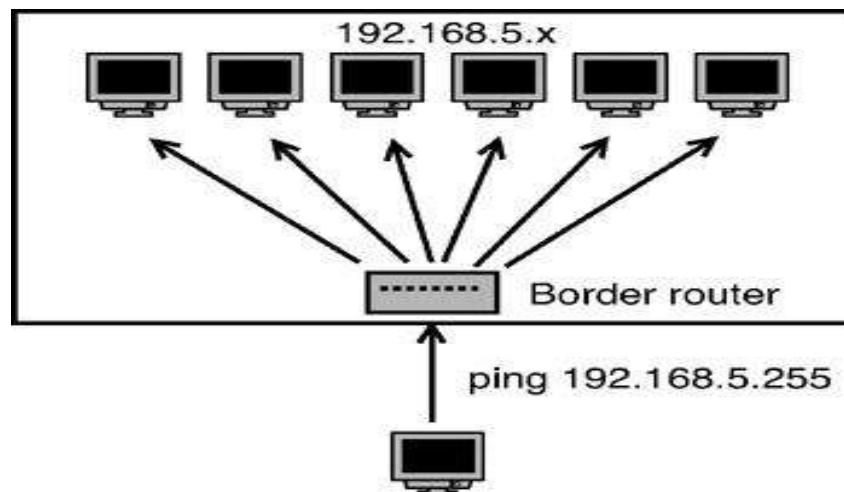


Figura 4.3 Las peticiones ICMP enviados a la dirección de difusión son vistos por todos los nodos de la red.

No todas las redes están configuradas para responder a transmitir el tráfico de esta manera, y muchos de los que no permiten que pase a través de los routers de frontera (ver figura 4.3); pero algunas redes, y estos se conocen como *amplificadores de difusión*.

No es inusual para un amplificador de emisión para contener varios cientos de hosts que

responden en su red; de vez en cuando un amplificador que contiene varios miles de los ejércitos hará emerger. Las matemáticas muestran que, incluso desde una conexión de acceso telefónico, un atacante todavía podría generar suficiente tráfico como para saturar un T1. Igualmente preocupante sobre el ataque Smurf es la facilidad con que se puede realizar. A diferencia de los ataques DDoS se miraron antes, el atacante no necesita pasar tiempo para comprometer las máquinas y la instalación de DDoS, herramientas y listas de amplificadores de difusión fácilmente se puede encontrar en Internet; y de curso al igual que muchas de estas herramientas que son muy fáciles de usar.

Las Inundaciones SYN

Uno de los ataques más populares DoS es el de inundaciones SYN, en los que la víctima es bombardeada con solicitudes de conexión, lo que en última instancia, las conexiones legítimas tienden a ser rechazadas, mientras que el consumo de recursos del sistema se agota.

Vamos a empezar por revisar cómo se crea una conexión TCP entre dos hosts.

En la primera el cliente envía un paquete TCP con el SYN (sincronización) establecido. Tras la recepción de este paquete, el servidor responde con un paquete TCP, esta vez con los SYN y ACK banderas (confirmación) se han establecido. Por último, el cliente responde a la SYN-ACK con su propia ACK. La conexión se ha establecido, y los datos pueden fluir.

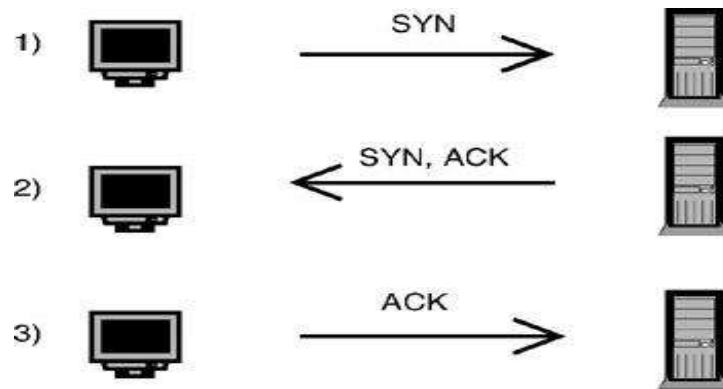


Figura 4.4 El apretón de manos de tres vías.

¿Qué sucede si el cliente no responde al SYN-ACK en el paso 3? El servidor se encuentra en espera de un corto período de tiempo (generalmente 180 segundos) y, a continuación, se da por vencido; pero durante este período de espera, la memoria asignada para la conexión está atado. La idea detrás de inundaciones SYN es bombardear el servidor con SYN paquetes, pero no seguir ACK al final. Esto deja a cientos de conexiones medio abiertas, toda la memoria que consumen. Finalmente, el servidor se quede sin memoria, o el núcleo decidirá que hay demasiadas conexiones pendientes. De cualquier manera, el resultado será que los nuevos intentos de conexión de hosts legítimos se les niega (figura 4.4).

La mayoría de SYN inundaciones utilizan campos de origen IP forjado en parte para ocultar el origen del ataque, y en parte porque algunos demonios definen un número máximo de conexiones por cliente y se reducirá de inmediato los intentos de conexión adicionales. El inconveniente, sin embargo, es que si una máquina inocente recibe un SYN-ACK (debido a que su dirección IP está siendo utilizado como una dirección de origen en los paquetes falsificados), se asume que ha habido algún tipo de error, y envía un RST (reset) de ella misma. El servidor recibe la RST, y libera la memoria reservada por la inicial SYN. Esto no es lo que quiere el atacante. La solución consiste en suplantar las direcciones IP de las máquinas inexistentes, un rápido `nmap` revelaría algunos bloques de red no utilizados y adecuados.

Los Ataques de Denegación de Servicio-Oriented Nonbandwidth

La mayoría de los demonios de red registran sus actividades, el nivel de detalle de este tipo de registros en general, siendo configurable por el administrador. Los archivos de registro pueden llegar a ser muy grande, lo que podría llenar el sistema de archivo. Esto puede hacer que los propios demonios se bloqueen, por no mencionar que causa decenas de potenciales problemas con el resto del sistema.

Si un atacante puede hacer archivos de registro para escribir una gran cantidad de datos, que por lo tanto, puede realizar un ataque de denegación de servicio en lugar crudo. La mayoría de los demonios de registro de los intentos de conexión, por lo que un explotador de este problema podría consistir en la creación de varias ocasiones y derribando las conexiones al demonio. Aunque esto podría tomar muchas horas para generar el registro significativo, en última instancia, sería eficaz.

El tipo más popular de ataque a esta categoría, sin embargo, es el envío de miles de mensajes de correo electrónico a la meta, un proceso conocido como mail bombing. Esto no sólo utilizar el espacio en disco significativo, sino que también causa congestión de la red y aumenta la memoria / CPU que utiliza el MTA (Mail Transfer Agent), lo que molesta enormemente al administrador que tiene que tratar de separar a estos e-mails de legítimos e-mails. En los servidores de correo de ocupados (como las pertenecientes a los ISP), no solicitado de correo no deseado (spam) puede tener un efecto similar, debido al enorme volumen.

En una línea similar, cualquier demonio que procesa peticiones de los usuarios (y que, después de todo, es el propósito principal de un demonio) es susceptible a ataques / RAM consume CPU. Esto podría ser en forma de reiteradas solicitudes a Apache, BIND, y así sucesivamente. La buena noticia es que Linux es bastante resistente a este tipo de ataques; que pueden frenar la caída de servidor, pero no es probable hacer que se caiga.

Capítulo 5 Desarrollo de la Herramienta

5.1 Introducción

Hoy en día existen multitud de herramientas útiles para realizar auditorías, test de penetración, ataques o simplemente comprobaciones y pruebas de seguridad sobre Redes Wi-Fi.

La gran mayoría de las herramientas se centran en una parte determinada del proceso de análisis y ataque, bien sea en la detección de redes, captura de tráfico, descifrado de paquetes, cracking de contraseñas, etc.

Adicionalmente, las aplicaciones más utilizadas suelen ejecutarse desde línea de comandos y, debido a la gran cantidad de parámetros que pueden ajustarse en determinados ataques, puede resultar complicado e incómodo escribir el comando para su ejecución, pues en algunos casos estos pueden llegar a ocupar varias líneas. Otro inconveniente de esto, es que difícilmente se va a conocer la sintaxis exacta del programa, por lo que se deberá consultar la ayuda constantemente.

Para solucionar el primer problema, desarrolladores independientes han desarrollado soluciones que liberan al usuario de tener que realizar la instalación de cada una de las herramientas para cada función, creando los llamados live-CD, sistemas operativos que se cargan desde un CD o DVD o USB sin necesidad de una instalación previa, que tienen ya instalados algunas de las aplicaciones más útiles para el análisis de seguridad. Algunas de las más populares son BackTrack, WifiSlax, etc, estas distribuciones contienen no sólo herramientas para auditoría Wi-Fi, también disponen de todo tipo de aplicaciones relacionadas con el mundo de la Seguridad Informática.

Si no desea utilizar este tipo de distribuciones live-CD existen programas que contienen las distintas herramientas que se suelen utilizar, por ejemplo, para la rotura de los protocolos de cifrado, como WEP y WPA, el programa más popular es aircrack, que contiene una aplicación para la captura de tráfico, otra para su inyección y otra para el proceso de cracking de la contraseña, además de alguna otra utilidad como una pequeña aplicación para unir los archivos de captura.

No obstante, se elija la opción que se elija la ejecución de las aplicaciones sigue teniendo que realizarse independientemente de forma manual. Para tratar de solucionar esto se han desarrollado interfaces gráficas para alguna de estas herramientas, sin embargo, la mayoría

de usuarios continúan utilizando las herramientas originales, pues estas interfaces no permiten mantener el control. Otro tipo de solución ha sido crear pequeñas aplicaciones donde simplemente se incluyen los parámetros de cada herramienta en una interfaz y muestra el comando que se debe ejecutar. Algo tan sencillo como esto, para el uso de determinadas aplicaciones, puede resultar algo muy práctico y una gran liberación de trabajo.

Por los motivos anteriores que se han descrito, se hace necesaria la existencia de una aplicación que nos brinde la comodidad de reducir la escritura de comandos con diversas opciones, la facilidad de uso de interfaces gráficas, ya que ofreciendo la información de una forma clara, el uso de la aplicación se podrá realizar prácticamente de una forma intuitiva, sin necesidad de tener conocimiento profundo de dicha aplicación.

5.2 Requerimientos de la Herramienta

Hay que destacar que para el funcionamiento de la herramienta desarrollada es necesario contar con lo siguiente:

- Sistema operativo Linux. Cualquier distribución con kernel Linux servirá, para la ejecución de la aplicación, a pesar de que la aplicación se puede ejecutar en otro tipo de sistemas operativos, las funcionalidades no podrán ser utilizadas.
- Máquina virtual de Java. La GUI está desarrollada en Java, por lo que es necesario disponer de este interprete para poder ejecutar el código.
- Tarjeta de red inalámbrica. Es el dispositivo físico más importante, dado que dependiendo de las características y de la compatibilidad de esta y el Sistema Operativo hace la diferencia entre que funcione o no la aplicación. Antes de ejecutar la aplicación, el usuario debe de verificar que su tarjeta de red inalámbrica sea compatible con las diferentes versiones de Linux.
- Punto de Acceso Inalámbrico Wi-Fi. Será necesario disponer de un AP para realizar el análisis de seguridad.
- Diccionario. Es necesario la elaboración o la obtención de unos buenos diccionarios de datos para realizar el ataque de fuerza bruta, y así poder encontrar la contraseña.

5.3 Comandos Utilizados para Redes Inalámbricas

Network-manager

Network-manager [Archinetman], es un programa que proporciona a los sistemas la detección y configuración automática para conectarse a la red. Las funcionalidades de este programa son útiles tanto para redes inalámbricas como por cable. Adicionalmente este programa permite que las interfaces puedan cambiar de modalidad en línea y fuera de línea, además de dar preferencia a las conexiones por cable antes que a las inalámbricas, tiene soporte para conexiones por módem y para ciertos tipos de Red Privada Virtual (VPN).

Ifconfig

El comando ifconfig [Linuxcom], se utiliza para mostrar información sobre las interfaces de red conectadas y disponibles del sistema y también provee la funcionalidad de configuración de la interfaz de red.

Iwlist

El comando iwlist [Linuxiwlist], se utiliza para conocer los datos de las redes que la interface de red es capaz de detectar, proporciona el nombre de la red, el canal en el que transmite, el modo en el que funciona el Access Point, si tiene o no el manejo de criptografía, etc.

wpa_supplicant

El comando wpa_supplicant [Archsupp], ofrece soporte para Wi-Fi Protected Access (WPA) y Wi-Fi Protected Access 2 (WPA2), está adaptado para el hardware tanto portátiles/escritorio como para sistemas integrados. Es el componente IEEE 802.1x/WPA utilizado por las estaciones cliente. Implementa las negociaciones entre la clave y un WPA Autenticador, y controla el marcado y la asociación/autenticación con IEEE 802.11.

wpa_passphrase

El comando wpa_passphrase [Linuxpass], genera un WPA PSK de una frase de contraseña ASCII para un cierto identificado de red SSID.

Dhclient

El programa Dhclient [Lindhc], se encarga de obtener y renovar con un servidor de DHCP las variables de configuración de red. Estas variables o parámetros son necesarios para la operación de la interface en un ámbito de red.

5.4 Diccionario de Claves o Contraseñas

Un diccionario es una colección de diversos números y/o palabras o combinación de estos, que regularmente son comunes que se emplean en un cierto idioma. Conforme a [WikiAtakdic], un ataque de diccionario es un método de ataque que consiste en intentar averiguar una contraseña probando toda la gama del diccionario.

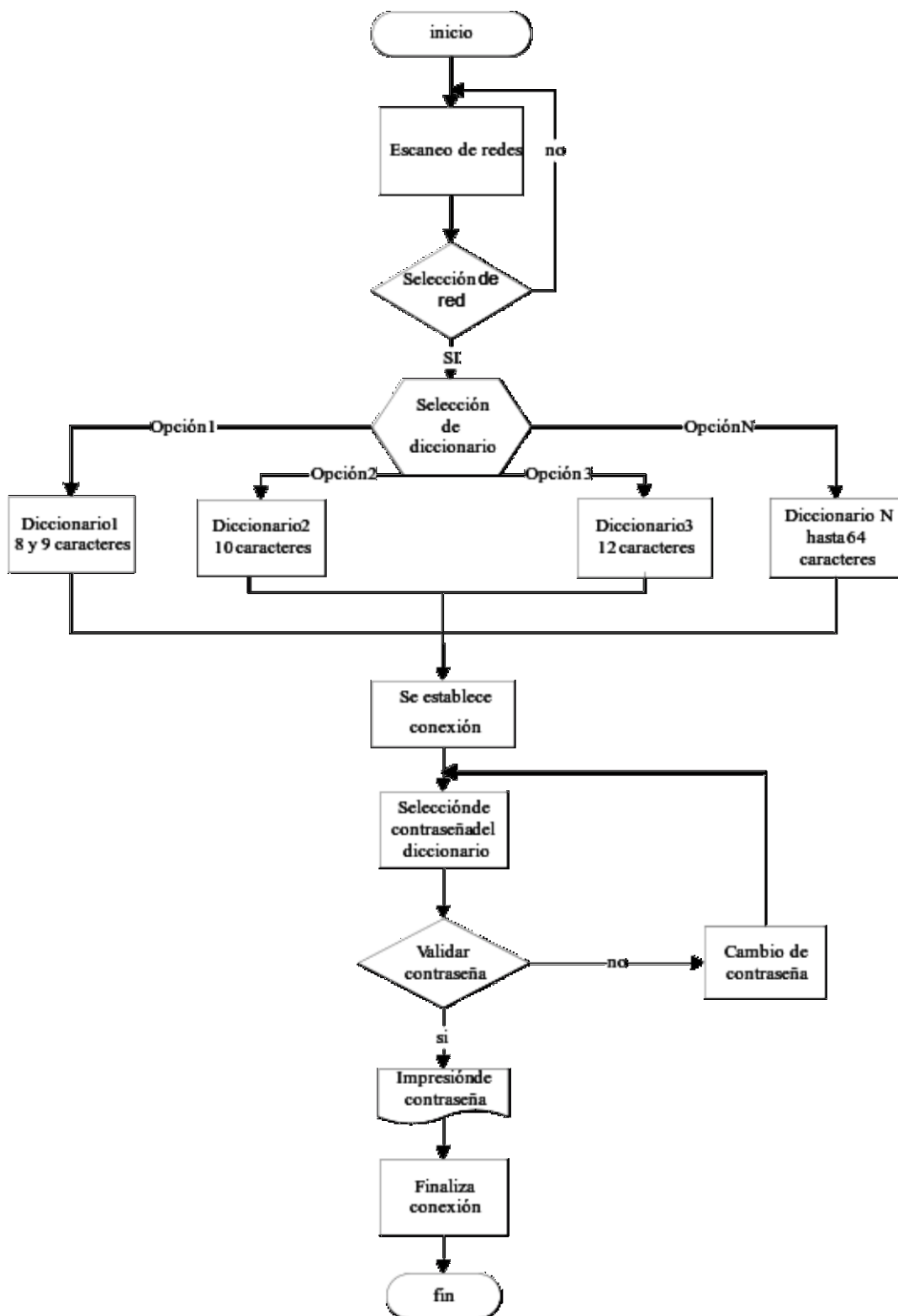
Los ataques de diccionario tienen pocas probabilidades de éxito con sistemas que emplean contraseñas fuertes con letras en mayúsculas y minúsculas mezcladas con números (alfanuméricos) y con cualquier otro tipo de símbolo. Sin embargo, para la mayoría de los usuarios recordar contraseñas tan complejas resulta complicado.

Existen variantes de ataque de diccionario que comprueban también algunas de las típicas sustituciones (determinadas letras por números, intercambio de dos letras, abreviaciones), así como distintas combinaciones de mayúsculas y minúsculas.

5.5 Operación de la Herramienta

La operación de herramienta consiste en escanear el ambiente inalámbrico, posteriormente se requiere seleccionar una red disponible y la elección de un diccionario de datos. Posteriormente a estos pasos iniciales, la herramienta comienza a probar de manera secuencial todas las opciones del diccionario seleccionado. Esta operación se presenta en el diagrama 5.6, que corresponde al diagrama de Flujo de Operaciones.

5.6 Diagrama de Flujo de Operaciones



5.7 Interface Gráfica de Usuario

Nos explica [Deitel7ma] una interfaz Gráfica de usuario (GUI) debe presentar un mecanismo amigable al usuario para interactuar con una aplicación. Una GUI proporciona a una aplicación una apariencia visual única. Al proporcionar distintas aplicaciones en las que los componentes de la interfaz de usuario sean consistentes e intuitivos, los usuarios pueden familiarizarse en cierto modo con una aplicación, de manera que pueden aprender a utilizarla en menor tiempo y con mayor productividad. En la figura 6.1 se presenta la GUI de la herramienta desarrollada con el lenguaje Java y el entorno de desarrollo Netbeans.

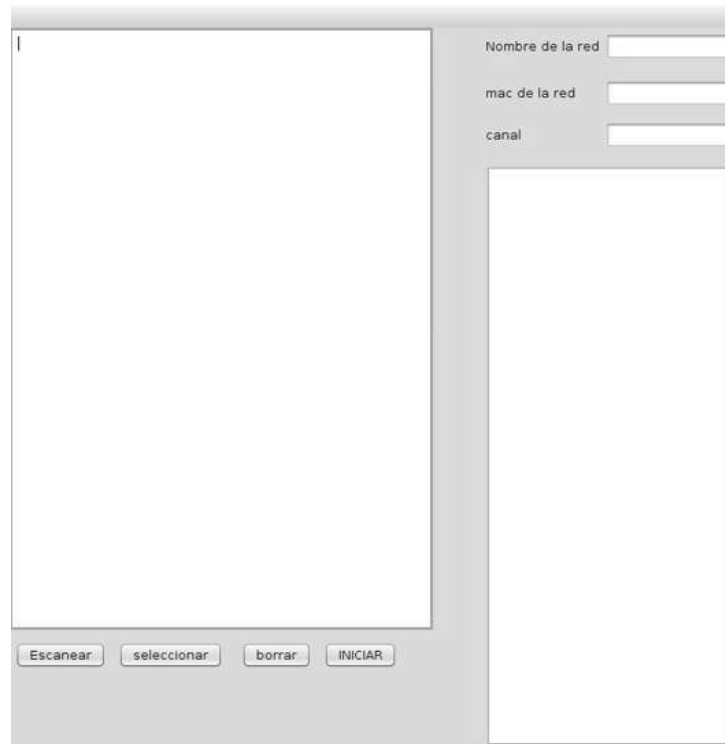


Figura 6.1 Interfaz Grafica

En este apartado, vamos a utilizar el botón de “Escanear” (figura 6.2), el cual nos dará un listado de todas las redes inalámbricas que se encuentran dentro del alcance de nuestra tarjeta de red.

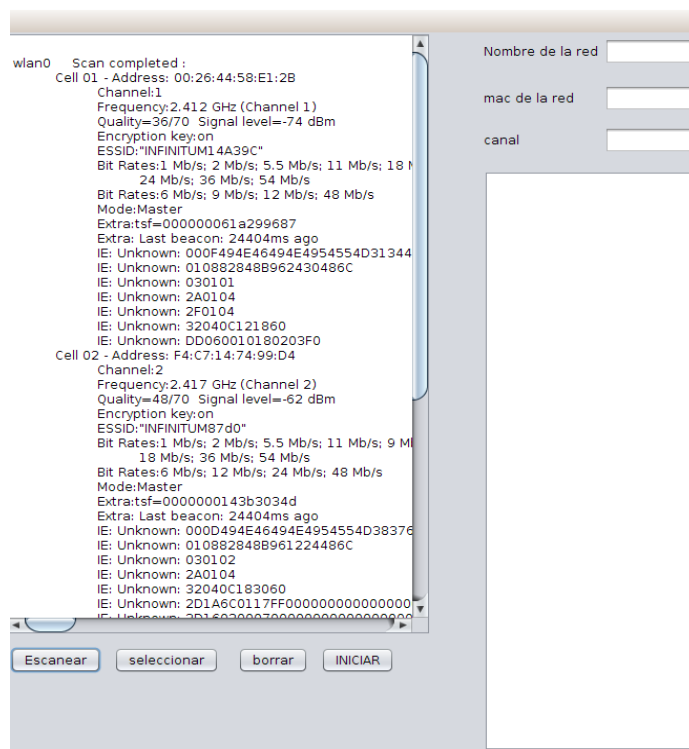


Figura 6.2 escaneo de redes

Una vez teniendo el listado de las redes al alcance, vamos a dar clic en el botón de “seleccionar”, para poder seleccionar la red a la cual vamos a tratar de ver su vulnerabilidad, la cual nos aparecerá en el lado derecho en los recuadros de nombre de red, MAC de red y el canal en el que transmite esta (figura 6.3).

SISTEMA DE ATAQUE DE DICCIONARIO PARA CLAVES WPA

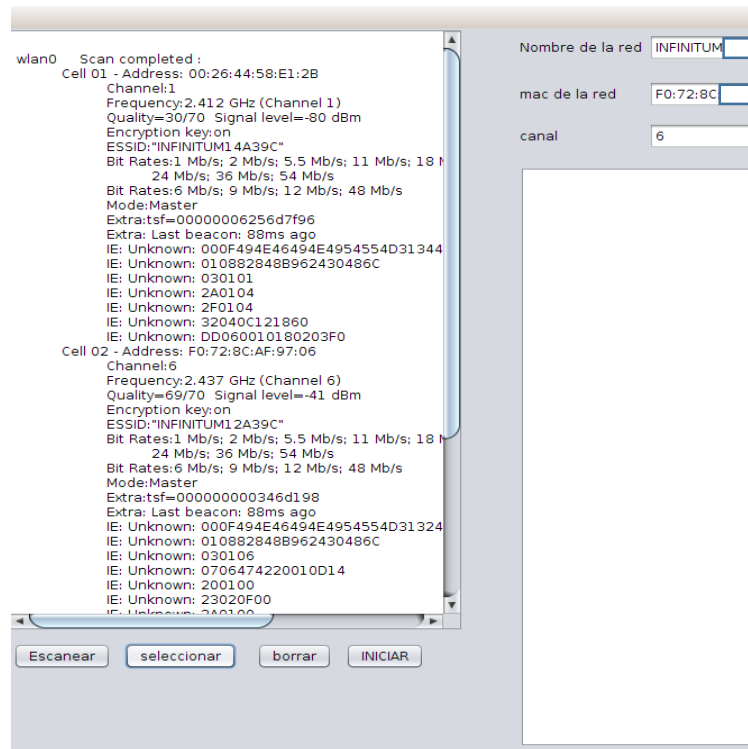


Figura 6.3 selección de red

Teniendo la red seleccionada, ahora podremos seleccionar el diccionario con el cual vamos a comenzar a averiguar la contraseña de esta red, y así darnos cuenta, si la podremos encontrar en un diccionario de contraseñas (figura 6.4).

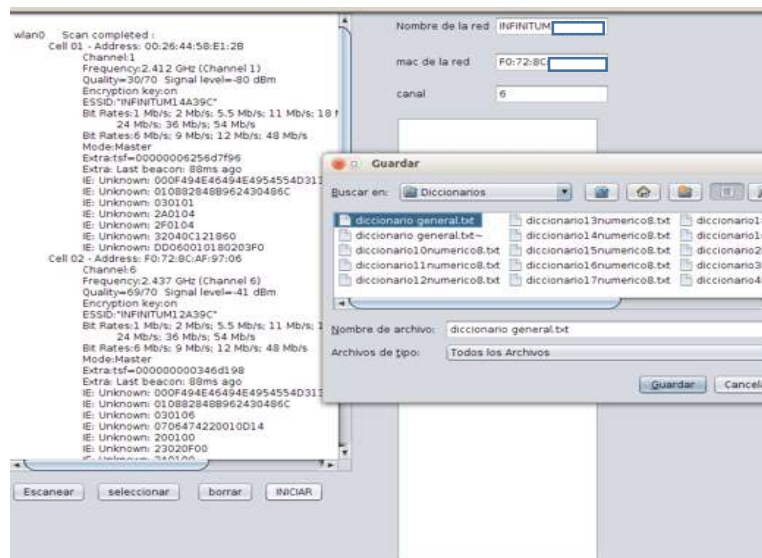


Figura 6.4 selección de diccionario y comienzo de ataque a red

Una vez encontrada la contraseña y establecida la conexión, el programa se detendrá, arrojándonos en un cuadro, la clave, ip obtenida, nombre de la red y el número de intentos (figura 6.5).

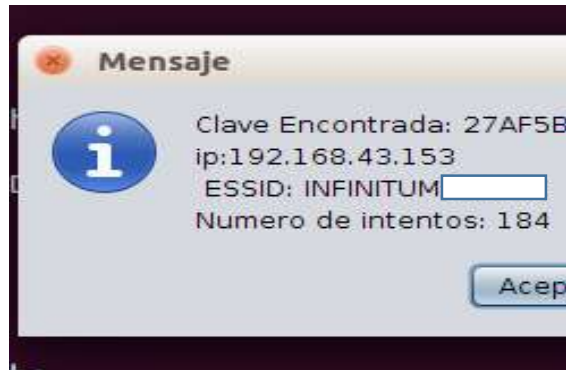


Figura 6.5 clave encontrada

Capítulo 6 Pruebas, resultados y recomendaciones

6.1 Introducción

Se realizó escaneo en diferentes puntos de accesos inalámbricos en distintos puntos de la ciudad de Morelia, donde se encontraron que en la mayoría de estos utilizan el cifrado wpa/wpa2.

6.2 Pruebas y Resultados

Pruebas

En distintos lugares de la ciudad de Morelia se realizaron pruebas a 10 redes inalámbricas, que se muestran en la tabla 6.1.

NUMERO	SSID	MAC	CANAL	PROTOCOLO
1	INFINITUM	F0:72:8C:	11	WPA
2	ARRIS-	CC:A4:62:	11	WPA
3	INFINITUM	58:98:35:	6	WPA
4	4627	E0:69:95:	5	WPA
5	MEGACABLE-	F4:DC:F9:	10	WPA2
6	HETZAMANY	E8:94:F6:	10	WPA2
7	INFINITUM	1C:8E:5C:	8	WPA2
8	val1um	00:8E:F2:	1	WPA2
9	megacable	14:AB:F0:	11	WPA2
10	INFINITUM	F4:C7:14:	1	WPA

Tabla 6.1 Redes inalámbricas probadas

Resultados

En las pruebas realizadas resultó que el 60% de los usuarios utilizan claves en los puntos de

acceso que pueden ser encontradas en un diccionario de palabras, tal como se muestra en la tabla 6.2.

NUMERO	SSID	TIEMPO	CONTRASEÑA	DICCIONARIO UTILIZADO
1	INFINITUM	120 min	27AF5BD45F	ALFANUMERICO DE 10 DIGITOS
2	ARRIS- [REDACTED]	280 min	8C8C45FAA1BA3E2D	ALFANUMERICO DE 16 DIGITOS
3	INFINITUM [REDACTED]	300 min	-	ALFANUMERICO DE 10 DIGITOS
4	4627 [REDACTED]	67 min	234151947	NUMERICO DE 8 A 9 DIGITOS
5	MEGACABLE- [REDACTED]	325 min	4435331040	NUMERICO DE 10 DIGITOS
6	HETZAMANY	300 min	-	ALFANUMERICO DE 10 DIGITOS
7	INFINITUM [REDACTED]	165 min	15E1CBE1BA	ALFANUMERICO DE 10 DIGITOS
8	val1um	300 min	-	ALFANUMERICO DE 10 DIGITOS
9	megacable [REDACTED]	300 min	-	NUMERICO DE 10 DIGITOS
10	INFINITUM [REDACTED]	217 min	3964636634	NUMERICO DE 10 DIGITOS

Tabla 6.2 Resultado de redes inalámbricas probadas

6.3 Recomendaciones

De las pruebas realizadas se puede recomendar el uso de una Metodología para la Generación de Contraseñas Robustas.

Con forme a [Redhasec4], existen varios métodos que las personas usan para crear

contraseñas seguras. Uno de los métodos más populares incluye acrónimos. Este procedimiento consiste en usar una frase que nos sea muy recordada para no olvidarla, puede ser algún acontecimiento importante para nosotros o simplemente algún hecho que nos haya pasado en nuestro día a día, por ejemplo:

sobre el rio y entre el bosque, a la casa de la abuela vamos.

Ahora, lo podemos convertir en un acrónimo (incluyendo la puntuación), para ello vamos a tomar la primera letra de cada palabra usada en nuestra frase memorable, donde como resultado la siguiente palabra:

seryeeb,alcdlav.

Ahora le podemos añadir complejidad sustituyendo números y símbolos por letras en el acrónimo, por ejemplo, sustituir “e” por “3” y “a” por el símbolo “@”.

s3ry33b,@lcdl@v.

Para hacer nuestro ejemplo de contraseña aún más seguro, podemos cambiar una letra de esta frase de una minúscula por una mayúscula, por ejemplo la letra C, lo que nos daría como resultado la siguiente frase:

S3ry33b,@lCdI@v.

Por último, es importante mencionar que no se debe utilizar la contraseña de ejemplo, puede crear sus propias contraseñas con un grado de dificultad acorde a como quiera, pero este solo es un ejemplo.

Capítulo 7 Conclusiones

Conclusiones

Se logró desarrollar la herramienta propuesta para el ataque de diccionario en redes WPA, utilizando el sistema operativo Linux y la programación en Java con el entorno de desarrollo Netbeans. Se aplicó la herramienta en un entorno real para valorar el nivel de seguridad de las contraseñas empleadas en puntos de acceso inalámbricos. En la ciudad de Morelia, Mich., se realizaron pruebas a 10 diferentes puntos de acceso inalámbricos con claves WPA, de lo cual resultó que el 60% de los puntos de acceso se emplean contraseñas débiles, las cuales ser identificadas mediante un ataque de diccionario. De igual manera en este trabajo se hace la recomendación de emplear contraseñas robustas que pueden ser generadas mediante la metodología propuesta en el capítulo 6, para asegurar en los recursos de hardware y software los atributos de confidencialidad, integridad y disponibilidad.

Bibliografía

REDCOMTITTEL. Redes de computadores, Ed Tittel, McGraw-hill/Interamericana de España, SAU, primera edición, ISBN 84-481-4280-2

Smith1ed. Linux Network Security, Peter G. Smith Charles, River Medina, first edition, ISBN 1-58450-396-3

Deitel7ma. Como programar en Java, Deitel, Paul J. y Harvey M. Deitel, Pearson Prentice Hall, séptima edición, ISBN 13:978-970-26-1190-5

Aircrack. <http://www.aircrack-ng.org/>

Wikijonh. http://es.wikipedia.org/wiki/John_the_Ripper

Compredes. <http://recursostic.educacion.es/observatorio/web/es/component/content/article/961-monografico-redes-wifi?start=7>

Leakedmail .<http://www.acunetix.com/blog/news/statistics-from-10000-leaked-hotmail-passwords/>

Realpass2006. https://www.schneier.com/blog/archives/2006/12/realworld_passw.html

Redhassec4. https://access.redhat.com/documentation/en-US/Red_Hat_Enterprise_Linux/4/pdf/Security_Guide/Red_Hat_Enterprise_Linux-4-Security_Guide-en-US.pdf

WikiAtakdic. http://es.wikipedia.org/wiki/Ataque_de_diccionario

Archinetman. [https://wiki.archlinux.org/index.php/NetworkManager_\(Espa%C3%B1ol\)](https://wiki.archlinux.org/index.php/NetworkManager_(Espa%C3%B1ol))

Linuxcom. <https://www.hscripts.com/es/tutoriales/linux-commands/ifconfig.html>

Linuxpass. http://linux.die.net/man/8/wpa_passphrase

Linuxiwlist. <http://linux.die.net/man/8/iwlist>

Archsupp. [https://wiki.archlinux.org/index.php/WPA_suppllicant_\(Espa%C3%B1ol\)](https://wiki.archlinux.org/index.php/WPA_suppllicant_(Espa%C3%B1ol))

Lindhc. <http://linux.die.net/man/8/dhclient>

Wikisnif. http://es.wikipedia.org/wiki/Anexo:Tipos_de_packet_sniffers

Toxen00. Toxen, B. Real World Linux Security. PTR, 2000.

Group92. Group, N. W. Rfc 1321 the md5 message-digest algorithm. Inf. téc., April 1992. [Http://www.faqs.org/rfcs/rfc1321.html](http://www.faqs.org/rfcs/rfc1321.html).

ofStandards93. of Standards, N. I. y Technology. Fips pub 180-1 secure hash standard. Inf. téc., May 11, 1993. [Http://www.itl.nist.gov/fipspubs/fip180-1.htm](http://www.itl.nist.gov/fipspubs/fip180-1.htm).

Lucas Ballard06. Lucas Ballard, D. L., Fabian Monroe. Biometric authentication revisited: Understanding the impact of wolves in sheep's clothing. 2006. 15th USENIX Security Symposium.

Sonia Chiasson06. Sonia Chiasson, P. v. O. y Biddle, R. A usability study and critique of two password managers. 2006. 15th USENIX Security Symposium.

Drimer07. Drimer, S. y Murdoch, S. J. Keep your enemies close: Distance bounding against smartcard relay attacks. 2007. 16th USENIX Security Symposium.

Valentin Sgarciu06. Valentin Sgarciu, M. S. V. Smart card technology used in secured personal identification systems. Bucharest, Romania, 2006. Proceedings of the 5th WSEAS Int. Conf. on DATA NETWORKS, COMMUNICATIONS & COMPUTERS.

cpf09. Código penal federal. Inf. tec., January 23, 2009. [Http://www.cddhcu.gob.mx/LeyesBiblio/pdf/9.pdf](http://www.cddhcu.gob.mx/LeyesBiblio/pdf/9.pdf).

cpd07. Código penal para el distrito federal. Inf. tec., Febreaury 2, 2007. [Http://www.paot.org.mx/centro/codigos/df/pdf/cpdfn.pdf](http://www.paot.org.mx/centro/codigos/df/pdf/cpdfn.pdf).

cps06. Código penal para el estado de Sinaloa. Inf. tec., August 11, 2006. [Http://www.stj-sin.gob.mx/Leyes/CODPENAL.html](http://www.stj-sin.gob.mx/Leyes/CODPENAL.html).

lpd03. Ley de protección de datos personales del estado de colima. Inf. téc., June 14, 2003. [Http://www.ucol.mx/radio/textos/sip-4753.pdf](http://www.ucol.mx/radio/textos/sip-4753.pdf).

Gasser88. Gasser, M. Building a Secure Computer System. Van Nostrand Reinhold, 1988.

Anderson72. Anderson, J. P. Computer security technology planning study. Inf. téc., 1972. Volume I,II.

ofDefense05. of Defense, D. The defense advanced research projects agency, 2005.
[Http://www.darpa.mil](http://www.darpa.mil).

Mukkamala02. Mukkamala, S., Janoski, G., y Sung, A. Intrusion detection using neural networks and support vector machines. 2002. Institute of Mining and Technology Socorro New Mexico.

Lincolnd Laboratory99. Lincolnd Laboratory, M. Intrusion detection attacks database, 1999.
[Http://www.ll.mit.edu/IST/ideval/docs/1999/attackDB.html](http://www.ll.mit.edu/IST/ideval/docs/1999/attackDB.html).

CERT05. CERT. Computer emergency response team, 2005.
[Http://www.cert.org](http://www.cert.org).

Tanenbaum97. Tanenbaum, A. y Woodhull, A. Operating Systems Deasign and Implementation. Prentice Hall, 1997.

CERT03. CERT. Cert/cc overview. Inf. téc., Carnegie Mellon University, 2003.

Schiffman00. Schiffman,M. D. Libnet 101, part1: The primer. Inf. téc., Guardent, 2000. [Http://www.insecure.org/tools.html](http://www.insecure.org/tools.html).

nem05. Nemesis, 2005. [Http://nemesis.sourceforge.net/](http://nemesis.sourceforge.net/).

nma05. Nmap, 2005. [Http://www.insecure.org/nmap/](http://www.insecure.org/nmap/).

sel05. Selinux, 2005. [Http://www.nsa.gov/selinux](http://www.nsa.gov/selinux).

lid05. Lids, 2005. [Http://www.lids.org/](http://www.lids.org/).

Apéndice A. Código Fuente de la Herramienta desarrollada en Java.

```
import java.io.BufferedReader;
import java.io.DataInputStream;
import java.io.File;
import java.io.FileInputStream;
import java.io.FileReader;
import java.io.InputStreamReader;
import java.io.IOException;
import java.util.logging.Level;
import java.util.logging.Logger;
import java.util.StringTokenizer;
import javax.swing.JFileChooser;
import javax.swing.JOptionPane;

private void jButton2ActionPerformed(java.awt.event.ActionEvent evt) {
    StringTokenizer tokens = new StringTokenizer(redes[i]);
    String aux = "";
    ESSID.setText("");
    while(tokens.hasMoreTokens())
    {
        aux = tokens.nextToken();
        if(aux.equals("Address:"))
        {
            bssid = tokens.nextToken();
            ADDRESS.setText(bssid);
        }
        if(aux.startsWith("ESSID:"))
        {
            aux = aux.substring(0, aux.length()-1);
            aux = aux.substring(7, aux.length());

            essid = aux ;
            ESSID.setText(aux);
        }
        if(aux.startsWith("Channel:"))
        {
            aux = aux.substring(8, aux.length());
            canal = aux;
            Canal.setText(aux);
        }
    }
}
```

```

        }
        if(aux.equals("Address"))
        {
            tokens.nextToken();
            tokens.nextToken();
            tokens.nextToken();
            ADDRESS.setText(tokens.nextToken());
        }
    }
    i++;
    if(i== redes.length)
        i=0;
}

private void jButton3ActionPerformed(java.awt.event.ActionEvent evt) {
1  listado.setText("");
1  listado1.setText("");
    ADDRESS.setText("");
    Canal.setText("");
    ESSID.setText("");
}

public static void main(String args[]) {
    java.awt.EventQueue.invokeLater(new Runnable() {
        public void run() {
            new Proyecto().setVisible(true);
        }
    });
}

private javax.swing.JTextField ADDRESS;
private javax.swing.JTextField Canal;
private javax.swing.JTextField ESSID;
private javax.swing.JButton jButton1;
private javax.swing.JButton jButton2;
private javax.swing.JButton jButton3;
private javax.swing.JButton jButton4;
private javax.swing.JLabel jLabel1;
private javax.swing.JLabel jLabel2;
private javax.swing.JLabel jLabel3;
private javax.swing.JScrollPane jScrollPane1;
private javax.swing.JScrollPane jScrollPane2;
private javax.swing.JTextArea listado;
private javax.swing.JTextArea listado1;

```