

3 Anforderungen an VPNs

Die Einsatzgebiete für virtuelle private Netzwerke sind sehr vielfältig. Je nach den gestellten Anforderungen an Sicherheit, Quality-of-Service sowie anderen Rahmenbedingungen kann man, entsprechend dem Angebot der Service Provider, die komplette Weitverkehrsinfrastruktur, die eigene Business-to-Business-Kommunikation (B2B) und den Remote Access als virtuelles privates Netzwerk aufbauen.

Bei der Auswahl der geeigneten Technologie muss man sehr genau untersuchen, welche Anforderungen an das VPN gestellt werden. In der Regel resultieren diese aus Sicherheitsbedürfnissen, gefolgt von Kostenaspekten, der Verfügbarkeit und – abhängig von den eingesetzten Applikationen – den benötigten Bandbreiten und tolerierbaren Verzögerungszeiten.

3.1 Sicherheit

Im Bereich der Datensicherheit gibt es eine ganze Reihe von Anforderungen, die sich in verschiedene Bereiche gliedern:

- ▶ Datenvertraulichkeit
- ▶ Schlüsselmanagement
- ▶ Paket-Authentifizierung
- ▶ Datenintegrität
- ▶ Benutzer-Authentifizierung
- ▶ Benutzer-Autorisierung
- ▶ Schutz vor Sabotage
- ▶ Schutz vor unerlaubtem Eindringen

3.1.1 Datenvertraulichkeit

Es muss sichergestellt werden, dass Unbefugte die Daten auf ihrem Weg durch das Internet nicht lesen können.

Vielfach wird auch gefordert, dass das interne Netzwerk mit seinen Verkehrsbeziehungen (Quell- und Zieladressen, Protokoll- und Portnummern) ebenfalls nicht ausgespäht werden kann. Dies wird im allgemeinen durch die Ver-

schlüsselung der Paketdaten erreicht. Falls die Verkehrsbeziehungen ebenfalls geschützt werden sollen, muss das originale Paket vollständig in den Datenbereich eines neuen Pakets eingepackt werden. Dies nennt man *Tunneling*.

Als Verschlüsselungsverfahren sollten unbedingt standardisierte, wohl bekannte Verfahren wie DES (Data Encryption Standard, ein standardisiertes, weltweit eingesetztes Verschlüsselungsverfahren) oder Triple-DES eingesetzt werden. Meist gibt die eingesetzte VPN-Technologie aus Gründen der Interoperabilität die Verfahren auch schon fest vor. In der Praxis werden wegen ihrer hohen Geschwindigkeit ausschließlich so genannte symmetrische Verfahren eingesetzt, bei denen Sender und Empfänger den gleichen Schlüssel zum Ver- und Entschlüsseln der Daten benötigen.

3.1.2 Schlüsselmanagement

Um auf sicherem und vor allem automatischem Wege eine Verteilung von symmetrischen Schlüsseln zu ermöglichen, benötigt man ein zuverlässiges Schlüsselmanagement. Dessen Aufgabe besteht im Erzeugen aller benötigten Schlüssel zur Verschlüsselung, Integritätsprüfung und Authentifizierung und in deren Verteilung zu den richtigen Gegenstellen in einem VPN (siehe Abbildung 3.1).

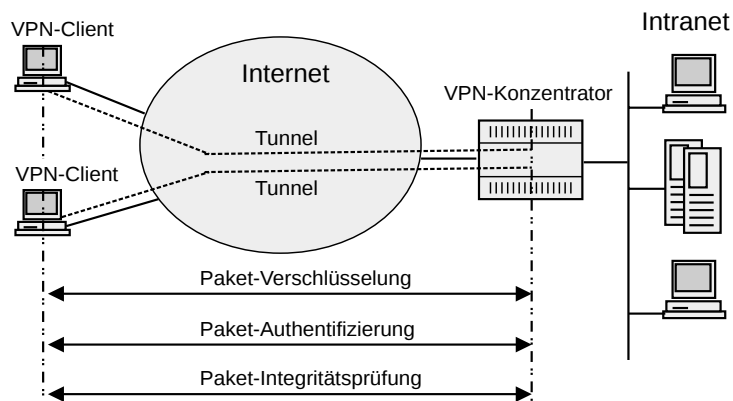


Abbildung 3.1: Beim Transport privater Daten über das Internet müssen verschiedene Sicherheitsaspekte berücksichtigt werden.

Gute Schlüssel, vor allem solche zur Datenverschlüsselung, haben eine relativ kurze Lebensdauer von meist einer Session oder wenigen Stunden und müssen deshalb sehr oft erzeugt und verteilt werden. Manuelle Verfahren scheiden aus diesem Grund, vor allem auch bei größeren Installationen aus. Da Out-of-Band-Verfahren, also die Verteilung der Schlüssel über ein anderes Kommunikationsmedium, praktisch den doppelten Aufwand bei der Auslegung eines Netzes erfordern, scheiden diese ebenfalls in den meisten Fällen aus.

Die heute bekannten Verfahren zum Schlüsselmanagement basieren meist auf so genannten asymmetrischen Verfahren, bei denen zum Ver- und Entschlüsseln jeweils unterschiedliche Schlüssel verwendet werden, von denen einer, der öffentliche Schlüssel (Public Key), allgemein bekannt sein darf. Diese Verfahren nennt man daher auch *Public-Key-Verfahren*.

3.1.3 Paket-Authentifizierung

Es muss garantiert werden, dass ankommende Pakete auch tatsächlich von dem authentischen Sender kommen und nicht von Dritten mit gefälschten Absenderadressen und neu berechneten Prüfsummen geschickt wurden.

Ähnlich wie bei einer Benutzer-Authentifizierung muss tatsächlich jedes ankommende Paket authentifiziert werden, was man durch symmetrische Schlüssel oder so genannte *Pre-Shared Secrets*, vertrauliche Daten, die nur dem Sender und dem Empfänger bekannt sind, erreicht. Aus Gründen der Geschwindigkeit und der Einfachheit kombiniert man dies meist mit Verfahren zur Prüfung der Datenintegrität.

3.1.4 Datenintegrität

Der Empfänger muss zuverlässig erkennen können, ob ein ankommendes Paket während des Transports verändert wurde.

Normale Prüfsummenverfahren reichen hierfür nicht aus, da ein Angreifer nach Änderung eines Datenpakets auch dessen Prüfsumme neu berechnen kann. Spezielle Verfahren auf Basis von symmetrischen Verschlüsselungsverfahren berechnen die Paketprüfsumme und fügen sie in das Paket mit ein. Die Schlüssel sind nur dem Sender und dem Empfänger bekannt. Ein Angreifer, der ein Paket ändern will, kann die Prüfsumme nicht korrekt berechnen.

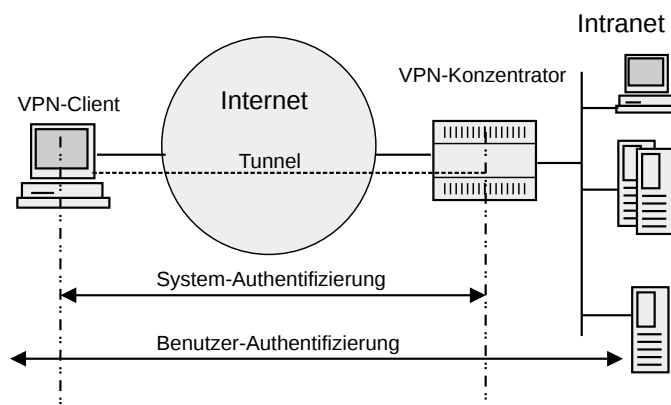


Abbildung 3.2: Neben der Verbindung selbst muss beim Remote Access auch der Benutzer authentifiziert werden.

3.1.5 Benutzer-Authentifizierung

Dies ist ganz wichtig bei Remote-Access-VPNs. Ein Benutzer, der über ein VPN-Gateway Zugriff auf das Intranet verlangt, muss seine Identität möglichst zuverlässig nachweisen.

Der Grad dieser Zuverlässigkeit ist von Fall zu Fall verschieden, so dass ein guter VPN-Konzentrator eine Reihe unterschiedlich starker Authentifizierungsverfahren unterstützen muss. Die Abstufung reicht dabei von einfachen Passwortverfahren bis hin zur Verwendung von Tokenkarten oder digitalen Zertifikaten. In diesem Bereich, zum Beispiel bei den Tokenkarten, gibt es keine verbindlichen Standards, so dass man hier immer proprietär ist. Für PKIs (Public Key Infrastructure, eine Infrastruktur zum Management von öffentlichen Schlüsseln und digitalen Zertifikaten) gibt es eine Arbeitsgruppe innerhalb der IETF, die auch schon eine Reihe von Standards hervorgebracht hat.

3.1.6 Benutzer-Autorisierung

Wird das VPN auch als Extranet verwendet, dann greifen auch externe Personen oder Organisationen, denen nur begrenzte Zugriffsrechte gewährt werden dürfen, auf Ressourcen des Unternehmensnetzes zu.

Die Autorisierung ist aber nur begrenzt ein Thema für Zugriffssysteme wie Router oder VPN-Konzentratoren. Denn diese Systeme sind nicht in der Lage, auf Benutzer- oder Gruppenebene rollenbasierende Zugriffe auf Verzeichnisse, Datenbanktabellen oder Drucker zu steuern. Diese Funktionen können nur von den Systemen wahrgenommen werden, die diese Ressourcen auch verwalten, also von Betriebssystemen und Datenbankmanagementsystemen.

Zugriffssysteme müssen aber die Fähigkeit besitzen, auf der Ebene von Netzwerkprotokollen eine Filterung vorzunehmen. Im Fall des IP-Protokolls heißt dies, dass aufgrund von Host- oder Netzwerkadressen, Protokollnummern, Portnummern usw. Entscheidungen getroffen werden können, ob ein Paket weitergeleitet oder verworfen wird. Aber dies ist keine Filterung auf Benutzerebene mehr, sondern auf Netzwerkebene, da Benutzer normalerweise nicht an feste IP-Adressen gebunden sind.

Falls die Filterungen auch aufgrund von Informationen in höheren Ebenen des OSI-Referenzmodells vorgenommen werden sollen, spricht man von Firewalls. Diese wurden bisher meist als Übergangspunkt vom Intranet in das Internet verwendet, sind aber auch als Schutzsystem im Intranet-Extranet-Übergang einsetzbar. Firewalls bieten eine Filterung von Netzwerkpaketen bis hin zu Dateninhalten, Gateways auf Applikationsebene und Integrationsmöglichkeiten von externen Applikationen wie z.B. Virenschutzprogramme. Manche Router oder VPN-Konzentratoren bieten als Option die Möglichkeit, Firewalls zu integrieren, jedoch ist es aus Gründen der Sicherheit und Performance oft besser, beide Funktionalitäten strikt voneinander zu trennen.

3.1.7 Schutz vor Sabotage

Das VPN-Gateway soll vor Angriffen sicher sein, die darauf abzielen, seine Funktionalität zu beeinträchtigen oder es funktionsunfähig zu machen.

Es gibt eine Reihe mehr oder weniger subtiler Arten von Angriffen auf Internetsysteme, denen auch VPN-Gateways ausgesetzt sein können. Vor der plumpsten Art einer solchen DoS-Attacke (Denial-of-Service, dabei wird verhindert, dass ein System seine Dienste erbringen kann), dem Überlasten einer Übertragungsstrecke oder einer Netzwerkschnittstelle, kann man sich nicht schützen, allerdings sind solche Angriffe mittlerweile sehr schnell zurückzuverfolgen, und der Angreifer bekommt in der Folge sehr viel Zeit, um über seine Untaten nachzudenken. Die Strafverfolgungsbehörden haben sich mittlerweile auf der ganzen Welt recht gut auf diese Art von Kriminalität eingestellt.

Subtilere Angriffe verschleiern, teilweise recht effizient, ihren eigentlichen Ursprung und schicken wenige, unverdächtig scheinende Pakete zu den Zielsystemen. Sie versuchen sie damit zu umfangreichen Aktivitäten zu bewegen, aufgrund derer ihnen für ihre eigentlichen Dienstleistungen immer weniger Ressourcen zur Verfügung stehen. So genannte DDoS-Angriffe (Distributed DoS) gehen noch einen Schritt weiter und legen ihren Angriff zweistufig an. Im ersten Schritt wird auf einer großen Anzahl von Rechnern unbemerkt ein Programm aufgespielt, das die eigentliche Attacke durchführen soll. Als Nächstes fangen alle diese Rechner zu einem bestimmten Zeitpunkt an, gleichzeitig mit unverdächtig scheinenden Paketen das Zielsystem regelrecht zu bombardieren. Die Zurückverfolgung des Angriffs ist sehr schwer, da er von Hunderten oder Tausenden weltweit verteilter Systeme auszugehen scheint, diese aber selbst auch nur »Opfer« sind. Der wirkliche Urheber hat inzwischen genug Zeit gehabt, seine Spuren zu verwischen.

3.1.8 Schutz vor unerlaubtem Eindringen

Ein VPN-Gateway muss verhindern, dass Unbefugte die Möglichkeit haben, über seine öffentlichen Schnittstellen in das Unternehmensnetzwerk zu gelangen.

Denn dies ist der direkte Weg zu den Informationen, die ein Angreifer ausspionieren oder manipulieren will. Der Hacker begnügt sich nicht mit dem, was über Netzwerkverbindungen übertragen wird, sondern er will Zugang zu den Systemen, auf denen die Daten gespeichert und verarbeitet werden. Ganz beliebte Systeme sind Authentifizierungssysteme, auf denen Passwörter und andere kritische Informationen gespeichert werden. Wer solch ein System »geknackt« hat, dem steht ein ganzes Unternehmensnetzwerk offen.

3 Anforderungen an VPNs

Da VPN-Gateways in vielen Fällen die einzige Schnittstelle zwischen einem Intranet und dem Internet sind, auf das Millionen Unbekannte Zugriff haben, müssen hier besondere Maßnahmen zum Zugriffsschutz getroffen werden. Man unterscheidet dabei:

- ▶ Physische Sicherheit
- ▶ Interface-Sicherheit
- ▶ Betriebssicherheit

Physische Sicherheit

Die Systeme müssen in sicheren Umgebungen betrieben werden. Die Sicherheit des Betriebsraums ist wesentlich, Security-Gateways dürfen nicht in Büros oder normalen EDV-Räumen betrieben werden, sondern nur in entsprechend sicheren Zonen. Sehr gute Systeme sind durch spezielle Maßnahmen gegen unerlaubtes Öffnen des Geräts abgesichert und erzeugen im Zugriffsfall einen entsprechenden Alarm im Netzwerkmanagementsystem.

Interface-Sicherheit

Die Interfaces, die mit dem Internet verbunden sind, sollten spezielle Mechanismen zum Schutz gegen die verschiedenen Arten von Angriffen aufweisen. Als besonders geeignet erweisen sich so genannte *gehärtete* IP-Stacks, die eines großen Teils ihrer normalen Funktionalität beraubt und damit gegen eine ganze Reihe von Angriffen immun sind, da die meisten Angriffspunkte überhaupt nicht mehr vorhanden sind.

Als schlecht im Sinne von unsicher sind im Allgemeinen VPN-Systeme einzustufen, die als Programm oder Prozess auf nicht sicheren Betriebssystemen laufen. Trotz anders lautender Behauptungen kann die Sicherheit eines Programms nicht größer als die Sicherheit des Betriebssystems sein, auf dem es läuft. Insbesondere die IP-Stacks vieler Betriebssysteme sind beliebte Ziele von Angriffen verschiedenster Art.

Betriebssicherheit

Hier gilt es dafür Sorge zu tragen, dass sich durch den Betrieb des VPN-Gateways keine Hintertüren für Angreifer öffnen. Da verschiedenen Studien zufolge der weitaus größte Teil von Angriffen von innen heraus erfolgt und nicht vonseiten des Internets, sind hier spezielle Maßnahmen nötig. Insbesondere sollte die Administration der Systeme im LAN verschlüsselt erfolgen, also zum Beispiel mit SSL oder noch besser mit IPSec. Überhaupt sollte man sich ganz genau vergewissern, welche Zugriffsmöglichkeiten im privaten Netz gegeben sind und wie damit umzugehen ist. Ein cleverer Angreifer – und viele davon sind sehr clever – wird, sobald er sieht, dass ein VPN bei-

spielsweise IPSec mit starker Verschlüsselung benutzt, sofort alle Gedanken an einen Angriff auf die IPSec-Verbindung selbst aufgeben und nach anderen Schwachstellen im System suchen. Und eben diese gilt es zu vermeiden.

3.2 Verfügbarkeit

Ein virtuelles privates Netzwerk soll traditionelle Weitverkehrs- oder Remote-Access-Lösungen ergänzen oder sogar ganz ersetzen. Dies bedeutet aber auch, dass ein VPN eine Verfügbarkeit bieten muss, die nicht unter der von herkömmlichen WAN-Infrastrukturen liegt. Die »alten« Lösungen, die auf Standardfestverbindungen, Frame Relay und ISDN basieren, weisen in der Regel eine hohe Verfügbarkeit auf.

3.2.1 Die Verfügbarkeit von Wählverbindungen

Der Remote Access wird über entsprechende PC-Karten oder externe Geräte mittels Wählverbindungen über das analoge oder digitale Fernsprechnet, bei Bedarf auch über das Mobilfunknetz, aufgebaut und in einem Remote-Access-Konzentrator terminiert. Diese Netze bieten in der Regel Verfügbarkeitszeiten von mindestens 99,999%. Mit anderen Worten, das Netzwerk zwischen der ISDN-Karte im Notebook eines Außendienstmitarbeiters und dem Primärmultiplexanschluss seines Unternehmens fällt pro Jahr maximal 5,2 Minuten aus!

An dieser Stelle auch gleich eine Anmerkung zum Begriff der Verfügbarkeit: Der obige Wert bedeutet nicht, dass das Netz im Jahr im Schnitt 5,2 Minuten ausfällt, sondern dass es garantiert nicht länger als 5,2 Minuten außer Betrieb ist. Bei den in solchen Netzen eingesetzten Hochverfügbarkeitssystemen ist es nicht selten der Fall, dass sie jahrelang überhaupt nicht ausfallen.

Für viele Remote-Access-Infrastrukturen ist die so genannte »Niemals besetzt«-Eigenschaft eines Telefonnetzes ebenso wichtig. Haben Sie schon jemals am soeben abgehobenen Hörer eines am öffentlichen Fernsprechnet angeschlossenen, funktionsfähigen Telefonapparates kein Freizeichen gehört? Vermutlich nicht. Wenn nun mehr Zugangskanäle in ein Unternehmensnetzwerk gelegt werden, als es potenzielle Benutzer gibt, kann man sich damit eine Remote-Access-Lösung mit tatsächlich garantiertem Zugriff aufbauen, da man selbst Einfluss auf die Auslegung seiner Systeme hat. Dies ist in vielen Fällen ein entscheidendes Kriterium, zum Beispiel beim Einsatz von Applikationen für Buchungs- oder Reservierungssysteme.

3.2.2 Die Verfügbarkeit von permanenten Verbindungen

Das Gleiche gilt für die meisten Festverbindungen und anderen WAN-Services. In ihren Verträgen garantieren die Provider entsprechende Verfügbarkeiten. Die für Standardfestverbindungen eingesetzten Fernvermittlungssysteme der großen Carrier bieten Verfügbarkeiten ähnlich der von ISDN, und auch die Frame-Relay- und ATM-Technologie ist mittlerweile mit so genanntem »Carrier-grade«-Equipment aufgebaut, das auch ein mittleres Erdbeben unbeschadet übersteht.

3.2.3 Die Verfügbarkeit von IP-VPNs

Diese Verfügbarkeiten muss ein virtuelles privates Netzwerk ebenfalls bieten können, soll es ein traditionell aufgebautes Netzwerk ergänzen oder gar ersetzen. Denn üblicherweise investiert man nicht in eine Technologie, die qualitativ schlechter ist als die aktuelle. Es sei denn, die neue Technologie bietet enorme Kostenvorteile. Dann ist man schon eher geneigt, einen Kompromiss zu machen oder sich zu überlegen, welche Verfügbarkeit denn wirklich benötigt wird.

Es gilt also, überhaupt erst einmal zu ermitteln, welche Verfügbarkeit tatsächlich benötigt wird. In sehr vielen Fällen arbeiten Unternehmen beispielsweise mit Festverbindungen nicht wegen ihrer 99,99% Verfügbarkeit, sondern einfach, weil es zum Zeitpunkt der Entscheidung keine wirklich besseren Alternativen gab. Um die notwendige Mindestverfügbarkeit zu ermitteln, muss man sich Klarheit verschaffen, welche Applikationen in welchem Maße das VPN benutzen werden und welche Implikationen Verbindungsabbrüche verursachen können. Es ist auch wichtig zu wissen, welcher Art der Datenverkehr ist und welches Zeitprofil er aufweist. Wird das VPN zum Beispiel nur während der Bürostunden benutzt, z.B. für Online-Anwendungen, kann man getrost eine nur halb so gute Verfügbarkeit akzeptieren, ohne eine Qualitätseinbuße zu haben. Statistisch gesehen fällt mehr als die Hälfte der Ausfallzeit in Zeiten, in denen das VPN ohnehin nicht benutzt wird.

Die Verfügbarkeit eines Internet-VPN kann man nicht allgemein beurteilen. Man muss dabei drei Fälle unterscheiden:

1. Die Verbindungen, auf denen das VPN basiert, gehen zu einem einzigen Provider.
2. Die Verbindungen, auf denen das VPN basiert, gehen zu zwei oder mehreren Providern, die miteinander kooperieren und entsprechende Durchleitungsverträge und/oder Service Level Agreements (SLA) abgeschlossen haben.
3. Die Verbindungen, auf denen das VPN basiert, gehen zu verschiedenen Providern, die keine Verträge miteinander abgeschlossen haben. Es ist nicht nachvollziehbar oder vorherbestimmbar, welchen Weg die Pakete nehmen und wie sie dort behandelt werden.

In den beiden ersten Fällen kann man mit dem oder den Service Providern geeignete Service Level Agreements abschließen, in denen neben anderen Eckwerten auch die Mindestverfügbarkeitszeit vertraglich geregelt ist. Hier werden mittlerweile in den meisten Fällen bereits Verfügbarkeitszeiten garantiert, die deutlich über denen der meisten lokalen Kundennetze liegen. Im dritten Fall ist dies nicht möglich; man kann zwar mit jedem einzelnen Provider SLAs vereinbaren, diese gelten aber nur für ihre eigenen Netze, und es ist keine durchgehende Verfügbarkeit garantiert.

Letztendlich mündet das Ganze wieder in eine Kostenkalkulation: Kosten die theoretisch möglichen Ausfälle mehr, als ich durch die neue VPN-Technologie einspare?

3.3 Performance

Mit geeigneten herkömmlichen WAN-Lösungen kann man die meisten Verbindungen mit einer garantierten, festen Bandbreite und kurzen Verzögerungszeiten betreiben.

3.3.1 Die Performance von Wählverbindungen

Eine Wählverbindung eines Remote-Access-Benutzers arbeitet, wenn er einem ISDN-B-Kanal benutzt, mit einer festen, garantierten, nicht schwankenden Übertragungsbandbreite von 64 Kbit/s und unmerklichen Verzögerungszeiten.

Modemverbindungen sind mit maximal 56 Kbit/s etwas langsamer und können bei schlechter Leitungsqualität sogar mit langsameren Geschwindigkeiten arbeiten. Dies tun sie leider in der Praxis auch meistens: Die wenigsten haben schon einmal ein V.90-Modem mit 56 Kbit/s an einem öffentlichen Fernsprechnetz arbeiten sehen, meist pegelt sich die Geschwindigkeit zwischen 40 und 50 Kbit/s ein.

3.3.2 Die Performance von permanenten Verbindungen

Standardfestverbindungen

Auch die Standardfestverbindungen, wie sie in Deutschland zum Beispiel von der Deutschen Telekom angeboten werden, bieten festgelegte, garantierte Übertragungsraten mit minimalen Verzögerungen, keinem Jitter und einer sehr hohen Verfügbarkeit.

Frame Relay

Bei Frame-Relay-Netzen kann man neben garantierten Mindestbandbreiten (CIR, Committed Information Rate) auch Spitzengeschwindigkeiten (CBR, Committed Burst Rate) vereinbaren, die meist der Geschwindigkeit der Zugangsleitungen entsprechen. So kann man sich zum Beispiel mit einer 64-Kbit/s-Standardfestverbindung auf das Frame-Relay-Netzwerk eines Providers aufschalten lassen und vereinbaren, dass jederzeit eine Übertragung mit 24 Kbit/s garantiert ist – dass man aber, falls es die Auslastung des Netzwerks zulässt, auch bis zur maximalen Leitungsgeschwindigkeit von 64 Kbit/s arbeiten kann. Die Kosten erhöhen oder reduzieren sich durch die hohe Geschwindigkeit nicht, da bei Frame Relay volumenabhängig abgerechnet wird.

ATM

Für ATM gelten die gleichen Kriterien wie auch für Frame Relay, jedoch sind je nach Provider Geschwindigkeiten bis 155 Mbit/s möglich.

3.3.3 Die Performance von IP-VPNs

Bei all diesen Technologien hat man feste oder kalkulierbare Performancewerte, die ein VPN ebenfalls bieten muss, wenn es diese ablösen soll. Auch hier muss man sich beim konkreten Anforderungsprofil die gleichen Gedanken machen, die man sich auch bei der Frage der Verfügbarkeit gemacht hat. Was brauche ich wirklich? Wenn eine Standardfestverbindung von 64 Kbit/s durch VPN-Technologie ersetzt werden soll, kann man dies beispielsweise auf eine ganz einfache Weise tun, indem man mehrere Tage oder Wochen rund um die Uhr den Datenverkehr analysiert und daraus ein Verkehrsprofil erzeugt. Daraus leitet man seine tatsächlichen Anforderungen hinsichtlich der benötigten Performance ab.

Garantien über Durchsatzwerte in Internet-VPNs kann man auch hier in einem Service Level Agreement mit einem oder mehreren Service Providern abschließen. Das ist allerdings bis zu einem gewissen Maße Augenwischerei. Denn in Wirklichkeit wird nichts garantiert, das ist den meisten Providern technisch auch gar nicht möglich, sondern es wird festgelegt, was passiert (Vertragsstrafe), wenn die Bandbreite nachgewiesenermaßen nicht zur Verfügung steht.

3.4 Quality-of-Service (QoS)

Bevor wir uns mit dem Thema Quality-of-Service (QoS) in IP-VPNs auseinander setzen, ist es sinnvoll, die QoS-Thematik selbst etwas eingehender zu beleuchten. Gerade in diesem Bereich, obwohl zur Zeit viel diskutiert, gibt es reichlich Begriffsverwirrung und vor allem auch eine Reihe von unterschiedlichen Konzepten zur Umsetzung des Gedankens.

Quality-of-Service wird manchmal als Maßnahme gegen eine zu kleine Bandbreite angesehen, aber auch das cleverste QoS-Konzept macht aus einer 64-Kbit/s-Leitung keine 2-Mbit/s-Leitung. Was QoS in Wirklichkeit tut, ist, die zur Verfügung stehende Bandbreite ungerecht zu verteilen. Die Systematik, die hinter dieser gezielten Ungerechtigkeit steckt, ist das Thema des folgenden Abschnitts.

3.4.1 Einführung in QoS-Konzepte

Die verschiedenen Konzepte, um QoS in Übertragungsnetzen einzuführen, resultieren aus den verschiedenen Applikationen, die solche Netze gleichzeitig benutzen, und aus der im Weitverkehrsbereich oft limitierten Übertragungskapazität. In lokalen Netzen ist es, im Vergleich zu Weitverkehrsnetzen, nicht so aufwendig, hohe Bandbreiten zur Verfügung zu stellen. In den letzten sechs Jahren hat sich die mittlere Bandbreite in Unternehmens-LANs etwa ver Hundertfacht, im WAN-Bereich ist solch ein Wachstum offensichtlich nicht zu beobachten. Aus einer 64-Kbit/s-Verbindung ist vielleicht eine 128-Kbit/s- oder 2-Mbit/s-Verbindung geworden, aber bestimmt keine 6,4-Mbit/s-Leitung. Andererseits operieren Unternehmen immer dezentraler, und die eingesetzten Applikationen werden immer hungriger nach Bandbreite. So fand zunehmend eine Überbuchung der vorhandenen WAN-Bandbreiten statt, und man begann zu untersuchen, ob denn alle Datenpakete die gleiche Priorität besitzen müssen und welche Anforderungen bestimmte Applikationen an die Netzwerkverbindung stellen. Denn das mittlerweile immer häufiger benutzte IP-Protokoll behandelt alle Pakete gleich, leitet sie weiter, so gut es geht (Best-Effort-Transport), und stellt somit ohne zusätzliche Maßnahmen ein Problem dar, falls man den Datenverkehr differenzieren will.

Qualitätsparameter

Welche Parameter bestimmen nun die Qualität eines Datenflusses oder einer Verbindung? Es sind, neben der Bandbreite, vor allem drei Werte, die darüber entscheiden, ob eine bestimmte Qualität gegeben ist – oder ob eine Übertragung überhaupt erfolgreich ist:

- ▶ Verzögerungszeit (Delay)
- ▶ Variation der Verzögerungszeit (Jitter)
- ▶ Mittlere Fehlerrate

Die Verzögerungszeit ist die Zeit, die vergeht, bis ein Datenpaket vom Sender zum Empfänger gelangt. Eine negative Beeinflussung dieses Wertes erfolgt üblicherweise in Vermittlungssystemen.

3 Anforderungen an VPNs

Eine Variation der Verzögerungszeit, der so genannte Jitter, liegt dann vor, wenn die Verzögerungszeit kein fester Wert und damit vorherbestimmbar ist, sondern sich laufend ändert.

Die mittlere Fehlerrate ist der Wert, der sich aus dem Verhältnis von korrekt übertragenen und zerstörten oder verlorenen Datenpaketen ergibt. Ein Paket gilt definitionsgemäß auch dann als verloren, wenn es aus Sicht einer Applikation zu spät ankommt. In Tabelle 3.1 sehen Sie, welche Anforderungen verschiedene Applikationen ihrer Natur gemäß an die verschiedenen Quality-of-Service-Parameter stellen.

	Verzögerung	Variation der Verzögerung	Fehlerrate
E-Mail	Niedrig	Niedrig	Hoch (0 Fehler)
Dateiübertragung (FTP)	Niedrig	Niedrig	Hoch (0 Fehler)
Videostrom (MPEG)	Hoch	Mittel	Niedrig
Videokonferenz	Hoch	Hoch	Niedrig
Telefonie	Hoch	Hoch	Niedrig

Tab. 3.1: Die Priorität der Anforderungen verschiedener Applikationen an die QoS

Zusätzlich zu diesen Qualitätsparametern ist die verfügbare Bandbreite ein Faktor, der einen indirekten Einfluss auf die Verbindungsqualität hat. Die Bandbreite selbst wird in der Regel von der Physik bestimmt, also von elektrischen und optischen Parametern der Leitungen und von der Verarbeitungsgeschwindigkeit der Vermittlungssysteme. An diesem Wert kann meistens auch nichts verändert werden: Die Bits werden in einer Leitung in der Reihenfolge transportiert, in der sie eingegeben wurden, und zwar mit der Geschwindigkeit bzw. dem Takt des Mediums. Eingreifen kann man aber in den Vermittlungssystemen, wenn man bestimmte Pakete oder bestimmte Datenflüsse dort unterschiedlich behandelt.

Qualitätsanforderungen verschiedener Applikationen

In Tabelle 3.1 haben Sie gesehen, dass unterschiedliche Arten von Anwendungen auch sehr unterschiedliche Anforderungen an die verschiedenen Qualitätsparameter stellen. So muss die mittlere Fehlerrate bei einer Dateiübertragung gleich 0 (Null) sein, denn sobald auch nur ein einziges Bit fehlerhaft ist, ist die übertragene Datei nicht mehr integer. Bei der Sprachkommunikation hingegen kann sogar das eine oder andere kleine Paket ganz verloren gehen, man ist ja von »schlechten« Telefonleitungen schon einmal das eine oder andere »Knacken in der Leitung« gewohnt – von dem, was ein Mobiltelefon-Benutzer über sich ergehen lassen muss, einmal ganz zu schweigen. Das Gleiche gilt auch für Videoübertragungen: Ein paar kleine schwarze Punkte

nimmt man in einem bewegten Bild subjektiv ohnehin nicht wahr. Was sich aber extrem schlecht auf solche Verbindungen auswirkt, sind große Verzögerungszeiten bei interaktiver Sprach- oder Videokommunikation. Bei einem Mail-Transfer wiederum ist dieser Wert nicht so tragisch, denn ob die E-Mail nach zwei oder nach acht Sekunden ankommt, ändert an deren Qualität überhaupt nichts. Welche Ansätze gibt es nun, die für bestimmte Applikationen erforderliche Qualität einzuhalten?

Vorherbestimmte Qualität

Ein Paradebeispiel hierfür ist eine Telefonverbindung, zum Beispiel ISDN. Hier wird für die Dauer der Verbindung zwischen den beiden Gesprächspartnern ein Datenpfad mit einem Vielfachen der erforderlichen Bandbreite permanent durch das Netzwerk geschaltet. Die Verzögerung ist minimal, subjektiv garantiert nicht wahrnehmbar und konstant. Somit sind alle drei Qualitätskriterien voll erfüllt.

Eine andere Idee ist es, das Transportnetz einfach so überzudimensionieren, dass die gesamte Bandbreite höher ist als die Summe der von allen Datenströmen benötigten Bandbreiten. Somit haben auch im schlimmsten Fall alle Applikationen ihre benötigte Bandbreite, eine relativ kurze, vorherbestimmbare und feste Verzögerungszeit und keine überlastbedingten Fehlerraten.

Allerdings hat dieses Konzept einen unschönen Nachteil, denn man schießt hierbei mit Kanonen auf Spatzen. Man stellt viel mehr Ressourcen zur Verfügung, als im Durchschnitt wirklich benötigt werden – und diese müssen auch bezahlt werden.

Flussbasierende Qualität

Hinter diesem Ansatz steht der Gedanke, eine Reihe von Netzwerkressourcen für einen bestimmten Datenfluss zu reservieren. Dieser Datenfluss kann zum Beispiel eine Verbindung zwischen Netzen, Rechnern oder Applikationen sein. Eine Voraussetzung dafür, dass dies funktioniert, ist sowohl ein Signalisierungsprotokoll als auch die Notwendigkeit, dass alle beteiligten Systeme – vor allem auch alle Vermittlungssysteme – diese Technologie implementiert haben.

Sobald ein Datenfluss beginnen soll, erfolgt durch das Signalisierungsprotokoll eine sukzessive Signalisierung an alle beteiligten Systeme, mit der Aufforderung, die entsprechenden Ressourcen für diesen Datenfluss zur Verfügung zu stellen und zu reservieren. Wenn dies erfolgreich war, also alle positiven Rückmeldungen erfolgten, dann kann die Übertragung beginnen.

Allerdings ist dieses Konzept mit einigen Problemen behaftet, insbesondere dem der mangelnden Skalierbarkeit in großen Netzen wie etwa dem Internet. In kleineren Netzen oder bei langer Lebensdauer der Datenflüsse ist dies kein Problem, aber die meisten Übertragungen im Internet haben nur eine kurze

Lebensdauer. Die Folge ist die, dass die permanent notwendige Signalisierung sowie die Reservierung und Freigabe bestimmter Ressourcen und deren Verwaltung bei Hunderttausenden oder Millionen von Datenflüssen auf einem Internet-Router nicht praktikabel sind.

Es wurden bisher in den RFCs 2211, 2212 und 2205 hinsichtlich der flussbasierenden Qualität verschiedene Standardisierungsbemühungen unternommen, jedoch setzen sie sich in großen Netzen wie dem Internet nicht durch, da die notwendigen Ressourcen, insbesondere für das Resource Reservation Protocol (RSVP) auf den Internet-Routern zu groß sind.

Klassenbasierende Qualität

Ein anderer Ansatz, der dieses Problem umgehen soll, ist die so genannte klassenbasierende Qualität oder Class-of-Service (CoS). Die zu übertragenden Netzwerkpakete werden in bestimmte, wohl bekannte Klassen eingeteilt und im Paket-Header entsprechend markiert. Die an diesem Konzept beteiligten Vermittlungssysteme kennen entsprechende Strategien, wie Pakete einer bestimmten Klasse zu behandeln sind. Abhängig von der Art der Vermittlungssysteme (ATM, Frame Relay, LAN-Switch usw.), werden diese Klassenparameter in geeigneter Weise auf die jeweilige Technologie angewendet. Hier ist keine Signalisierung mehr notwendig, denn die für die Verarbeitung benötigten Ressourcen kennt das Vermittlungssystem bereits.

Pakete verschiedener Quellen oder Applikationen werden, falls sie der gleichen Klasse angehören, innerhalb ihrer Klasse nach dem Best-Effort-Prinzip transportiert.

Aber auch dieses verbindungslose Prinzip weist einige Nachteile auf. Denn ohne eine Signalisierung ist für den Sender oder bestimmte Vermittlungssysteme nicht erkennbar, ob die erforderlichen Ressourcen für die benötigte Service-Klasse auf allen beteiligten Systemen überhaupt verfügbar sind.

Ein weiteres Problem kann entstehen, wenn verschiedene Netzwerke – und das Internet ist ein Konglomerat aus verschiedenen Netzen – verschiedene Service-Klassen definiert haben. Hier kann es im Extremfall dazu kommen, dass sich die Pakete nach dem Durchlaufen verschiedener Netze bei konservativer Umsetzung (im Zweifelsfall in die nächsthöhere Klasse) der CoS-Markierungen am Ende in der höchsten Service-Klasse wiederfinden. »Höchste Service-Klasse« klingt zwar gut, bedeutet aber in letzter Konsequenz tatsächlich einen Best-Effort-Transport der Pakete, da diese sich alle in einer einzigen Klasse drängeln.

3.4.2 Quality-of-Service bei Wählverbindungen

Die kritischen Qualitätsparameter bei Wählverbindungen hängen vom verwendeten Medium ab, also davon, ob man ISDN, analoge Telefonie oder Mobilfunk benutzt.

Im Fall von ISDN liegen alle Werte »im grünen Bereich«, die mittlere Fehler-rate ist praktisch null, die Verzögerung ist minimal, und es gibt keinen Jitter. Darüber hinaus ist die Bandbreite ebenfalls garantiert.

Analoge Verbindungen sind ebenfalls fast verzögerungsfrei, die mittlere Fehler-rate ist aus Sicht der höheren Protokolle ebenfalls null.

Mobilfunkverbindungen sind aus Datenübertragungssicht katastrophal, denn was dem menschlichen Gehirn an Verarbeitungsleistung abverlangt wird, um den gesprochenen und während der Übertragung oft verstümmelten Text zu entschlüsseln, das leisten Protokolle wie IP im Datenbereich nicht. Hier müssen entsprechende höhere Protokolle unter Umständen den Paket-verlust ausgleichen, und die Übertragung ist insgesamt nicht sehr effizient.

3.4.3 Quality-of-Service bei permanenten Verbindungen

Standardfestverbindungen

Die kritischen Qualitätsparameter bei Standardfestverbindungen sind alle sehr gut, die mittlere Fehler-rate ist praktisch null, es gibt keinen Jitter, und die Verzögerung kann insgesamt vernachlässigt werden. Darüber hinaus ist die Bandbreite ebenfalls garantiert.

ATM

ATM wurde, wie Sie vermutlich wissen, primär zum Transport von Sprache und Video entwickelt, aber auch zum Transport anderer Daten. Aus diesem Grund wurde von Anfang an auch auf das Thema Quality-of-Service geachtet. Die Technik wurde auf hohe Geschwindigkeit und geringe Verzögerungen getrimmt. ATM definiert fünf unterschiedliche Service-Klassen, die in Tabelle 3.2 beschrieben sind. Diese unterschiedlichen Klassen sind für verschiedene Arten von Applikationen wie Video, interaktive Sprachübertragung (Telefonie) oder Datenpaketübertragung geeignet.

Kategorie	ATM-Parameter	Applikationstyp
Constant Bit Rate (CBR)	PCR, CTD, CDV, CLR	Telefonie, Sprache
Real-time Variable Bit Rate (rtVBR)	PCR, CTD, CDV, CLR, SCR	Sprache/Video (Paket)
Non-real-time VBR (Nrt-VBR)	PCR, SCR, CLR	Frame Relay over ATM
Available Bit Rate (AVB)	PCR, MCR	Daten (IP, IPX)
Unspecified Bit Rate (UBR)	(keine)	Daten (IP, IPX)

Tab. 3.2: Die fünf ATM-Service-Klassen

3 Anforderungen an VPNs

In ATM definieren die folgenden sechs verschiedenen Parameter die fünf Service-Klassen:

- ▶ Peak Cell Rate (PCR), die maximal mögliche Rate, mit der ATM-Zellen übertragen werden können
- ▶ Cell Delay Variation (CDV), die Abweichung der Verzögerung zweier beliebiger ATM-Zellen
- ▶ Cell Loss Ratio (CLR), das Verhältnis von fehlerfreien zu fehlerhaft übertragenen oder verworfenen Zellen
- ▶ Sustainable Cell Rate (SCR), der mittlere, erreichbare Durchsatz
- ▶ Minimum Cell Rate (MCR), die minimale Zellübertragungsrate

Obwohl ATM zwischen verschiedenen ATM-Switches kleine Zellen überträgt, ist es verbindungsorientiert und basiert auf virtuellen Verbindungen (VC, Virtual Circuit). Es werden entweder feste virtuelle Verbindungen (PVC, Permanent Virtual Circuit) durch ein ATM-Netz konfiguriert, oder sie werden für die Dauer einer Verbindung geschaltet (SVC, Switched Virtual Circuit) und anschließend wieder deaktiviert. Im Internet, in dem ATM vor allem in den Backbones der Service Provider und Carrier eingesetzt wird, werden vor allem PVCs eingesetzt. ATM bietet die Möglichkeit, die Konfiguration eines ATM-Netzwerks zu vereinfachen, indem man verschiedene VCs zu einem virtuellen Pfad (Virtual Path, VP) zusammenfasst.

Frame Relay

Frame Relay bietet ebenfalls die Möglichkeit, verschiedene Service-Klassen zur Verfügung zu stellen. Diese unterscheiden sich hauptsächlich durch die zugesicherte Bitübertragungsrate (CIR, Committed Information Rate) und die maximal mögliche Übertragungsrate (EIR, Excess Information Rate).

Ähnlich wie ATM ist Frame Relay verbindungsorientiert und transportiert seine Daten in kleinen Paketen (Frames). Auch hier gibt es permanente (PVC) und geschaltete (SVC) virtuelle Verbindungen. Aufgrund seiner Architektur und seiner Ausrichtung auf die Datenübertragung ist nicht so gut für Sprach- oder Videoübertragungen geeignet und wird nicht für Geschwindigkeiten über 45 Mbit/s eingesetzt.

ATM und Frame Relay sind im Augenblick die einzige Möglichkeit, mit der die Service Provider den Kunden eine Quality-of-Service anbieten können. Der Provider stellt für jede Verbindung einen entsprechenden PVC mit den geeigneten Serviceparametern zur Verfügung. Allerdings existieren in einem großen Netz auf diese Art und Weise sehr viele PVCs, denn wenn ein Kunde beispielsweise zehn Niederlassungen sternförmig vernetzen will, muss der Service Provider $N \times (N-1) = 10 \times 9 = 90$ PVCs konfigurieren.

3.4.4 Quality-of-Service im IP-Protokoll

Das Internet-Protokoll wurde ursprünglich als verbindungsloses Verfahren zum asynchronen Transport von Datenpaketen entwickelt. Es garantiert keine Paketzustellung oder bestimmte Zeiten, in denen ein Paket übertragen wird. Die Vermittlungssysteme (Router) verarbeiten die Pakete in der Reihenfolge, in der sie eintreffen, und machen sich keine Gedanken über deren Wichtigkeit oder die mit den Paketen verknüpfte Anwendung. Falls auf einem schnellen Medium wie Ethernet mehr Pakete in einem Router ankommen, als dieser auf einem langsameren Interface ausgeben kann, kommt es zu Verzögerungen oder im schlimmsten Fall zu Paketverlusten. Dieses Verhalten ist kennzeichnend für IP und das Internet.

Dies war lange Zeit auch kein Problem, da über dieses Netz in den Anfängen meist nur E-Mail, FTP und Telnet betrieben wurden – alles Applikationen, die sehr gutmütig auf durch Bandbreitenengpässe bedingte Verzögerungen reagieren. Auch das World Wide Web mit seinem grafischen Interface war zunächst an die Bedingungen im Internet angepasst. Da aber zunehmend Applikationen, die auf den LAN-Betrieb zugeschnitten sind und Anwendungen wie Voice- oder Video-über-IP eingesetzt werden, ist das Best-Effort-Prinzip von IP nicht mehr adäquat. Und genau hier liegt auch das Problem von IP, denn im Gegensatz zu ATM oder Frame Relay, die von vornherein für die Möglichkeit, Quality-of-Service zu bieten, entwickelt wurden, war dies nie die Idee von IP, und man hatte auch nichts dafür vorgesehen. »Nichts« ist allerdings nicht ganz richtig. Es gibt doch ein Byte im IP-Header, das viele Jahre ziemlich stiefmütterlich behandelt und von den meisten Systemen schlicht ignoriert wurde: Das Type-of-Service-Byte (TOS-Byte), das eigentlich für Quality-of-Service-Zwecke gedacht war. Dieses in der Vergangenheit sehr vernachlässigte Feld im IP-Header erlebt zur Zeit eine Renaissance, verbunden mit seiner Umbenennung, denn es heißt jetzt DSCP (Differentiated Services Code Point) und dient dazu, die Zugehörigkeit eines Pakets zu einer bestimmten Service-Klasse anzuzeigen.

3.4.5 Die IP-Differentiated-Services-Architektur (DiffServ)

DiffServ ist eine klassenbasierende Architektur (RFC2475) für die Dienstqualität in IP. Ihr werden aufgrund ihrer hohen Skalierbarkeit die besten Chancen als dominierendes Verfahren eingeräumt. Die Grundidee ist, dass die IP-Pakete beim Eintritt in ein Netzwerk einer bestimmten Klasse zugeordnet und entsprechend markiert werden. Die Vermittlungssysteme eines DiffServ-Netzwerks werten die Markierungen aus und behandeln die Pakete entsprechend. Die Systeme kennen also nur Service-Klassen und unterscheiden weder bestimmte Datenflüsse noch Adressen. Sie sind für jede Service-Klasse auf ein bestimmtes Verhalten (PHB, Per Hop Behaviour) konfiguriert, das heißt sie lei-

ten eingehende Pakete in entsprechende Warteschlangen verschiedener Priorität oder verwerfen sie unter bestimmten Umständen (z.B. Überlast) auch.

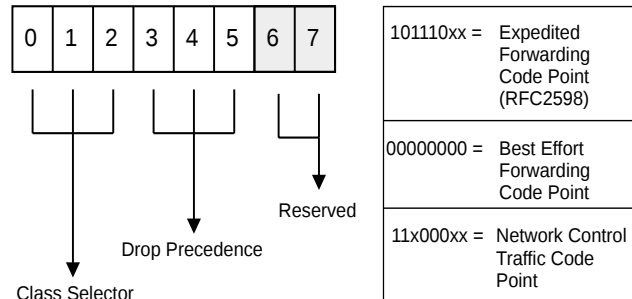
DiffServ verlegt die Komplexität der Klassifizierung des Datenverkehrs auf die Grenzen des Netzwerks. Diese Systeme müssen die zu übertragenden Datenpakete beim Eintritt in das Netz klassifizieren, markieren und bereits ihrer Service-Klasse entsprechend behandeln. Die Router im Netzwerk müssen die Entscheidung, wie ein Paket zu behandeln ist, lediglich aufgrund des DSCP-Felds im IP-Header treffen. Es erfolgt weder eine Signalisierung zwischen den Routern noch müssen die Zustände einer Vielzahl verschiedener Datenflüsse verwaltet werden. Viele dieser Router haben bereits Mechanismen und Ressourcen zur Paket-Priorisierung implementiert, so dass eine Erweiterung auf die DiffServ-Architektur nur ein logischer Schritt ist, der meist keine teuren Erweiterungen der Systeme nach sich zieht.

Die Architektur von DiffServ legt fest, dass die Klassifizierung und Markierung der Pakete in einem so genannten Edge-Router erfolgt, also einem Router an der Grenze eines Netzwerks. Endsysteme oder Applikationen müssen diese Architektur nicht unterstützen, können dies aber und tun es auch zunehmend. Wo die Klassifizierung und Markierung letztendlich stattfindet, ist daher unterschiedlich und hängt neben der Qualitätsstrategie des Netzwerks auch von den verschiedenen beteiligten Komponenten, also Applikationen, Betriebssystemen und Netzwerksystemen, ab. Es kann sowohl eine Klassifizierung in den Endsystemen stattfinden, z.B. durch Applikationen, oder dies kann auch auf dedizierten Netzwerkkomponenten nach entsprechenden Regeln (Policies) erfolgen. Diese Regeln können zentral festgelegt und verwaltet und dann auf die beteiligten Systemkomponenten verteilt werden. Es kann auch nötig sein, durch geeignete Policies bestimmte Klassifizierungen an Netzwerkgrenzen zu ändern, da unter Umständen in Endgeräten Markierungen durch Applikationen oder Benutzer vorgenommen wurden, die nicht mit der netzwerkweiten QoS-Strategie konform sind.

DiffServ ist ein Protokoll in der Schicht 3, also der Netzwerkschicht. Dadurch, dass es somit unabhängig von Layer-2-Protokollen (wie Frame Relay oder ATM) ist, kann es mit einer Vielzahl von Infrastrukturen genutzt werden. Die Router oder Switches müssen natürlich entsprechend konfiguriert werden, um die entsprechende Service-Klasse korrekt zu behandeln, also muss beispielsweise festgelegt werden, welche ATM-Service-Parameter wie zu konfigurieren sind. Aber dies erfolgt nur in eine Richtung: aus der Sicht des Schichtenmodells von oben nach unten, und es findet keine Wechselwirkung in der Art statt, dass Layer-2-Spezifika die Service-Klasse beeinflussen. Dies ist ein weiterer wichtiger Pluspunkt zugunsten des DiffServ-Modells, denn sehr viele größere Netze benutzen zum Transport von IP-Paketen eine Vielzahl verschiedener Infrastrukturen (ATM, IEEE802.1q, Frame Relay usw.), die unterschiedliche Quality-of-Service- und Priorisierungsmechanismen bieten.

Quality-of-Service (QoS)

Diese können den DSCP auswerten, ihn aber nicht ändern, so dass die Klasseninformation während der vollständigen Übertragung erhalten bleibt.



Assured Forwarding Code Points (RFC2597)

	Class 1	Class 2	Class 3	Class 4
Low Drop Precedence	001010	010010	011010	100010
Medium Drop Precedence	001100	010100	011100	100100
High Drop Precedence	001110	010110	011110	100110

Class Selector Code Points (RFC2474)

Class 0	Class 1	Class 2	Class 3	Class 4	Class 5	Class 6	Class 7
000000	001000	010000	011000	100000	101000	110000	111000

Abbildung 3.3: Der Differentiated Services Code Point (DSCP) im IP-Header

Die DiffServ-Service-Klassen

Zur Zeit sind in DiffServ drei verschiedene Service-Klassen festgelegt:

- Premium
- Tiered
- Best Effort

Zu jeder dieser Klassen gibt es ein festgelegtes *Per Hop Behaviour* (PHB), das in den Routern oder Switches eines DiffServ-Netzwerks unterstützt werden muss. Das PHB unterscheidet zwischen unmittelbarem Weiterleiten (EF, Expedited forward), garantiertem Weiterleiten (AF, Assured forward) und normalem Weiterleiten (DF, Default forward) der IP-Pakete. Wie das die verschiedenen Systeme technisch realisieren, ist nicht im Standard festgelegt.

3 Anforderungen an VPNs

Für den so genannten Premium-Service muss jeder Router permanent einen gewissen Teil seiner Ressourcen reservieren, unabhängig von seiner tatsächlichen Auslastung. Für den Tiered-Service gilt das Gleiche, jedoch werden diese Ressourcen in verschiedene Prioritätsstufen untergliedert. Was an Ressourcen übrig bleibt, wird vom Best-Effort-Service benutzt.

In Abbildung 3.3 sehen Sie den Aufbau des DSCP-Felds im IP-Header. DiffServ benutzt derzeit, wie im RFC2474 festgelegt, die ersten sechs Bits davon. Davon legen die ersten drei Bits die Klasse fest, und die folgenden drei Bits geben die Vorrangigkeit innerhalb der jeweiligen Klasse an, mit der Pakete in Überlastsituationen verworfen werden können. Der Codepunkt 101110 markiert den Premium-Service und der Codepunkt 000000 den Best-Effort-Service.

Der DiffServ-Edge-Router

Der Klassifizierer ist die Funktion, die in einem so genannten Edge-Router ein eingehendes Paket daraufhin prüft, welche Bits im DSCP-Feld zur Klassifizierung des Pakets gesetzt werden müssen. Er teilt die IP-Pakete somit in verschiedene Klassen ein, die der Betreiber des Transportnetzwerks unterstützt. Nach welchen Kriterien eine solche Klassifizierung erfolgt, richtet sich entweder nach dem Service-Level-Agreement (SLA), das mit dem Provider abgeschlossen wurde, und/oder nach der Qualitätsstrategie des Netzwerks. Die Entscheidung des Klassifizierers erfolgt zum Beispiel aufgrund von Quell- oder Zieladressen, Port- oder Protokollnummern oder aufgrund des DSCP-Felds. Denn es kann durchaus sein, dass ein Paket verschiedene Netze unterschiedlicher Provider durchläuft oder dass in einem Kundennetz bereits schon eine DiffServ-Klassifizierung erfolgt ist. Zur Zeit erfolgt die Markierung statisch nach fest im Edge-Router vorgegebenen Regeln. Anschließend wird das markierte Paket seiner Markierung entsprechend in den passenden Ausgangspuffer geschrieben und von dort in das Transportnetz weitergeleitet. In Abbildung 3.4 sehen Sie in dem Edge-Router neben dem Markierer und dem Klassifizierer noch einen so genannten Shaper. Dieses Modul sorgt dafür, dass der Datenverkehr in das Transportnetz einen bestimmten vereinbarten Wert oder die mögliche Bandbreite einer Verbindung nicht überschreitet.

Der DiffServ-PHB-Router

Der PHB-Router braucht sich keine Gedanken mehr über die Klassifizierung oder Markierung eines Pakets zu machen, er muss es nur aufgrund der Einträge im DSCP-Feld verarbeiten und in die korrekte Warteschlange stellen. In Abbildung 3.5 können Sie erkennen, dass der PHB-Router eine Teilmenge der Funktionen des Edge-Routers aufweist und somit viel weniger belastet wird. Die Verarbeitung ist demzufolge relativ schnell, denn es muss nur das DSCP-Feld ausgewertet werden. Da das PHB für die EF-Klasse bedingt, dass ein jitterfreies, schnelles Weiterleiten der Pakete garantiert ist, muss der Provider umfangreiche Maßnahmen treffen, um dies zu erreichen – und dies verur-

sacht natürlich entsprechende Kosten. Zukünftige Modelle, die auf dem Einsatz so genannter dynamischer Policy-Server basieren, ermöglichen eine dynamische Zuteilung von Ressourcen für die EF-Klasse aufgrund der aktuellen Netzlast und befreien den Provider damit von der statischen Reservierung der notwendigen Betriebsmittel.

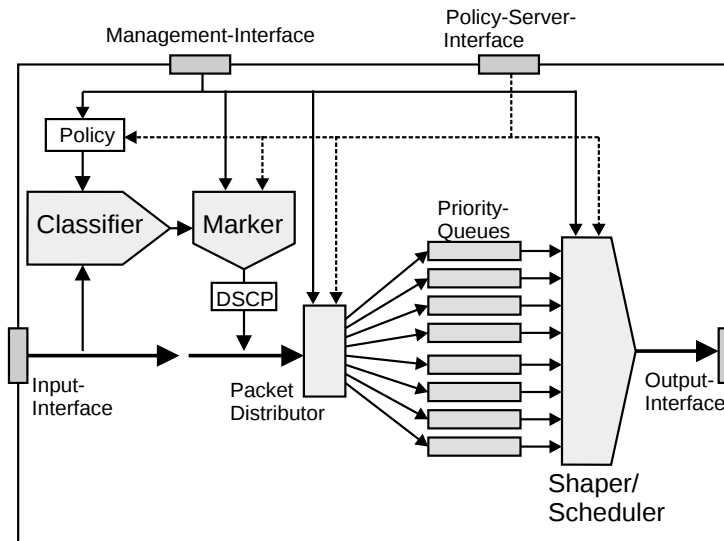


Abbildung 3.4: Der DiffServ-Edge-Router klassifiziert und markiert die IP-Pakete.

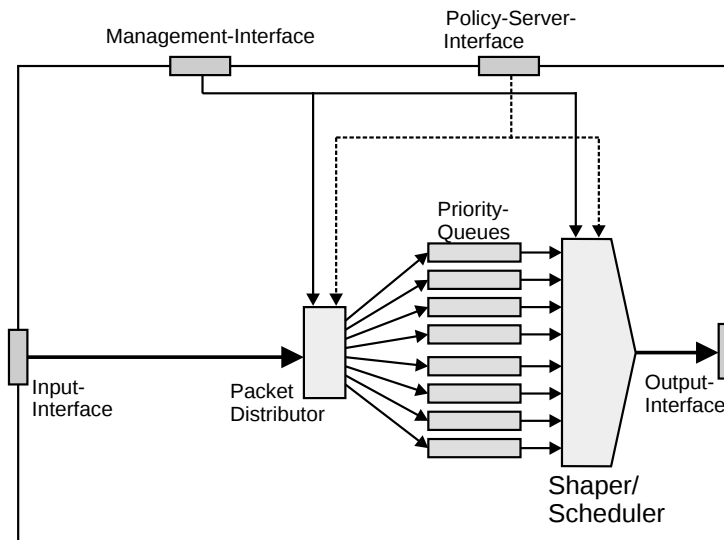


Abbildung 3.5: Der DiffServ-PHB-Router verarbeitet die IP-Pakete anhand des DSCP-Feldes.

3.4.6 Differentiated Services in IP-VPNs

Hier müssen zwei unterschiedliche Szenarien betrachtet werden: Setzt der Kunde in seinem Netzwerk bereits DiffServ ein und möchte er, dass der Provider die Pakete entsprechend behandelt, oder soll eine entsprechende Klassifizierung der IP-Pakete erst beim Eintritt in das Provider-Netzwerk erfolgen?

Im ersten Fall muss eine Vereinbarung zwischen Kunde und Provider darüber geschlossen werden, wie die Entsprechung der Service-Klassen des Kunden und den mit ziemlicher Sicherheit unterschiedlichen Service-Klassen des Providers auszusehen hat. Weiterhin muss dafür Sorge getragen werden, dass die DSCP-Informationen beim Einkapseln des Pakets durch ein Tunnelingprotokoll nicht verloren gehen. In IP-VPNs benutzt man meist das IPSec-Protokoll, das genau dies tut: Es kopiert den Inhalt des DSCP-Felds bei ausgehenden Paketen in den neu erzeugten IP-Header und macht es damit für den Provider verfügbar. Beim Entkapseln des IP-Pakets auf der Empfängerseite wird das DSCP-Feld im äußeren Header verworfen. Dadurch ergibt sich auch die Möglichkeit für den Service Provider, das DSCP-Feld des IPSec-Pakets beim Eintritt in sein Netzwerk zu ändern, also eine neue Klassifizierung vorzunehmen, ohne dass das originale DSCP-Feld davon betroffen ist. Bei den Layer-2-Tunneling-Protokollen fehlt diese Funktionalität zurzeit.

Im zweiten Fall ist das Problem, je nach Reihenfolge der Verarbeitungsschritte und der Systeme möglicherweise überhaupt nicht lösbar. Nehmen wir einmal den heute üblichen Fall, ein Internet-VPN mit dem IPSec-Protokoll im Tunnelmodus aufzubauen. In diesem Fall wird das originale IP-Paket vollständig verschlüsselt und in ein anderes IP-Paket eingekapselt. Dies erfolgt meist in einem IPSec-Gateway des Kunden. Das IPSec-Paket wird nun zum Router des Service Providers geschickt, der eine Klassifizierung vornehmen soll. Nur wonach? Alle dafür notwendigen Parameter wie originale IP-Adressen, Port- und Protokollnummern usw. sind verschlüsselt. Einzig die äußeren IP-Adressen sind noch erkennbar, aber diese ermöglichen nur eine Priorisierung nach Quell- und Absenderstandort – etwas, was normalerweise nicht zur Paketklassifizierung benutzt wird.

Mögliche Auswege sehen zum Beispiel so aus, dass die IPSec-Funktionalität erst nach der Klassifizierung im Edge-Router des Providers greift oder dass das VPN-Gateway des Kunden bereits eine Klassifizierung vornimmt, die nach den Angaben des Service Providers konfiguriert wurde. In diesem Fall muss das IPSec-Gateway die Funktionalität eines Edge-Routers bieten.

3.5 Skalierbarkeit und Migrationsfähigkeit

Wenn Sie heute über den Einsatz eines VPN nachdenken, dann sollten Sie auch bereits an übermorgen denken. Langfristige Voraussagen sind aber ziemlich schwierig und treffen manchmal auch nicht zu, also strebt man eine möglichst hohe Offenheit seines Systems an – Offenheit sowohl in Hinblick auf die Einhaltung von Standards, um gegebenenfalls einen Hersteller wechseln zu können, als auch in Hinblick auf die Modularität und Erweiterbarkeit der Systemkomponenten.

»Think big – start small«

Mit diesem Leitsatz liegt man genau richtig, denn für VPNs gelten die gleichen Maßstäbe hinsichtlich Skalierbarkeit und Migrationsfähigkeit wie für andere Netzwerktechnologien auch. Während der Planungsphase sollte man sich bereits Gedanken machen, wie das Netz in der Zukunft aussehen kann. Denn sehr viele Rahmenbedingungen sind stetiger Änderung unterworfen, wie zum Beispiel die Zahl der Benutzer, benötigte Bandbreiten und Qualitätsmerkmale durch neue Applikationen, rechtliche Gesichtspunkte und neue Geschäftsfelder oder Akquisitionen.

Beim Thema Skalierbarkeit muss man aber unbedingt zwischen Leistung und Sicherheit differenzieren. Denn die Sicherheitsstrategie eines Unternehmens richtet sich nach völlig anderen Kriterien. Ein VPN-Gateway in einem kleinen Zweigbüro benötigt beispielsweise keinen so hohen Datendurchsatz wie seine Gegenstelle in der Unternehmenszentrale, auf der Hunderte von Verbindungen konzentriert werden, jedoch muss es den gleichen Sicherheitsanforderungen genügen. Hier gilt es sehr genau zu prüfen, ob Endgeräte für kleine Außenstellen, die vom Systemdesign her für eine kleine Übertragungsbandbreite ausgelegt sind, diese auch dann noch bieten, wenn eine starke Verschlüsselung eingesetzt wird.

Ein Faktor, der sich in den meisten Netzwerken laufend ändert, ist die steigende Zahl von mobilen Anwendern und Heimbüros. Hier sollte, wie es sich in der Praxis gezeigt hat, von vornherein mit einer entsprechend großen Zahl kalkuliert werden.

Die Skalierbarkeit im Bereich der Systemleistung ist ein ganz wesentliches Entscheidungskriterium bei der Planung eines VPN. Meist sind Standorte verschiedener Größe, Heimbüros und mobile Mitarbeiter mit der notwendigen Technologie auszustatten. Je nach Einsatzgebiet sind unterschiedliche Datendurchsätze und Anschlusstechnologien notwendig, vom redundanten VPN-Konzentrator bis hinunter zur VPN-Clientsoftware für PCs. Idealerweise sollen diese verschiedenen VPN-Systeme aber auch eine möglichst einheitliche Konfigurations- und Managementoberfläche bieten.

3.6 Integration in existierende Netze

Selten bekommt ein Netzwerkplaner die Gelegenheit, ein VPN auf der grünen Wiese zu planen. Er hat viel häufiger die Aufgabe, dass ein VPN in bestehende Infrastrukturen zu integrieren. Diese Infrastrukturen sind bestehende lokale Netze, Weitverkehrsnetze und Remote-Dienste sowie die dazugehörigen Management-, Accounting- und Überwachungssysteme. Das zunehmende Verlangen vieler Unternehmen nach Kostentransparenz auch im Netzwerkbereich führt zu nutzer- oder kostenstellenbezogenen Abrechnungssystemen, die vor allem in Verbindung mit WAN- und Remote-Access-Diensten eingesetzt werden, da diese durch die zusätzlich an die Carrier zu entrichtenden Gebühren sehr kostenintensiv sind.

3.6.1 Management

Insbesondere sind bei der Integration der VPN-Komponenten die bereits vorhandenen Managementsysteme zur Konfiguration und Überwachung zu unterstützen, die auch zu diesem Zweck für traditionelles Netzwerkequipment eingesetzt werden.

Die wichtigsten Funktionalitäten sind die zur Konfiguration, zur Überwachung und zur Abrechnung der Netzwerkdienste.

Konfiguration

Die Konfiguration von Netzwerkkomponenten erfolgt mittlerweile fast ausschließlich über SNMP (Simple Network Management Protocol, eine Menge von Funktionen zum Konfigurieren und Überwachen von Netzwerkgeräten und einem Übertragungsprotokoll hierfür). Dieses Protokoll wurde entwickelt, um eine problemlose Interoperabilität zwischen Netzwerkmanagementprodukten unterschiedlicher Hersteller zu ermöglichen. Eine Netzwerkmanagementstation, meist eine Grafikworkstation, dient als zentrales Element zur Steuerung und Überwachung, während auf den Netzwerkkomponenten selbst so genannte SNMP-Agents residieren. Diese Agents kommunizieren über das Netzwerk mit der Managementstation, um Konfigurationsparameter zu setzen oder abzufragen. Eine weitere wichtige Funktionalität der Agents ist ihre Fähigkeit, bei kritischen Zuständen auf dem Netzwerkgerät einen so genannten SNMP-Trap zu erzeugen, der unmittelbar zur Managementstation geschickt wird, um diese über das aufgetretene Problem zu informieren.

Die Kommunikation der verschiedenen SNM-Komponenten erfolgt über UDP-Pakete, in denen verschiedene Kommandotypen zum Setzen und Lesen von MIB-Objekten oder -Variablen eingekapselt werden. Eine MIB (Management Information Base) ist eine Datenstruktur, die die Summe aller managbaren Objekte der Netzwerkkomponenten darstellt. Leider ist die Sicherheit,

wenn man es überhaupt so nennen kann, von SNMP auf einem nicht allzu hohen Stand, so dass die Kommunikation mit einfachsten Methoden aufgezeichnet oder unbefugt verändert werden kann. Aus diesem Grund werden sicherheitskritische Systeme im Allgemeinen nicht über SNMP konfiguriert. Auch die Überwachung mit SNMP beschränkt man auf nicht sicherheitskritische Parameter, wie Zähler für übertragene Bytes oder Prüfsummenfehler und Ähnliches.

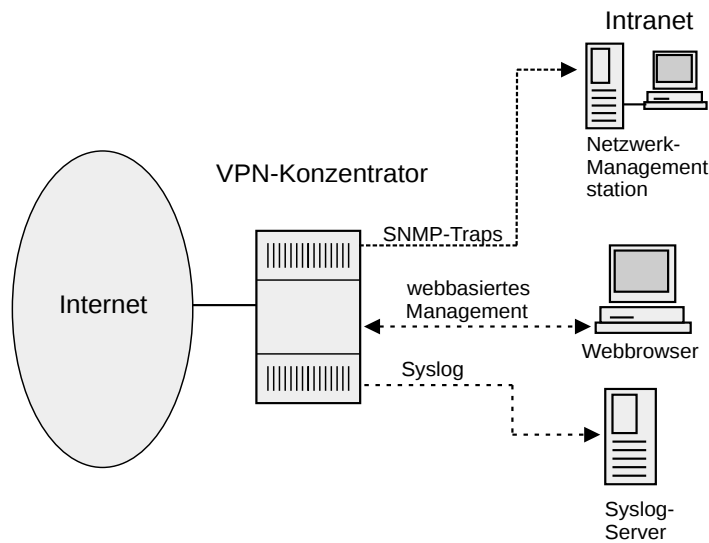


Abbildung 3.6: Ein VPN-System muss mehrere Managementprotokolle unterstützen.

Bei sicherheitskritischen Netzwerkkomponenten benutzt man andere Protokolle zu deren Konfiguration. Neben herstellerspezifischen, oft nicht offen gelegten Protokollen haben sich hier das SSL-Protokoll (Secure Socket Layer) und IP Security (IPSec) durchgesetzt. Diese Protokolle sind standardisiert und mittlerweile weit verbreitet.

Die Anforderungen an das Management, die aus dem hier Gesagten resultieren, sind genau genommen ein Kompromiss aus dem Wunsch nach reibungsloser Integration in bestehende SNMP-basierende Managementsysteme und der Forderung nach ausreichender Sicherheit. Es kann also auf einem VPN-Gerät durchaus SNMP unterstützt werden, jedoch mit der Einschränkung, dass die Konfiguration des Geräts abweichend vom Protokoll auf keinen Fall mit SNMP-Set-Befehlen durchgeführt werden darf, sondern über ein sicheres Protokoll erfolgen muss.

In letzter Zeit hat sich das so genannte Web-based-Management immer weiter verbreitet. Hierbei greift man über einen Webbrowser direkt auf die Netzwerkgeräte zu, um sie zu konfigurieren oder abzufragen. Der Vorteil dieser

Methode besteht darin, mit dem Webbrowser eine grafische Benutzeroberfläche zu haben, mit der mittlerweile fast jeder umgehen kann und die praktisch auf jedem Arbeitsplatzrechner installiert ist. Ein weiterer Pluspunkt ist, dass neue Funktionen auf dem Gerät auch gleich in der grafischen Oberfläche verfügbar sind. In traditionellen SNMP-Umgebungen mit speziellen Netzwerkmanagement-Applikationen müssen die Funktionen dem Programm erst bekannt gemacht werden, bevor man sie benutzen kann. Vor allem mit grafischen Darstellungen, insbesondere in heterogenen Umgebungen, tun sich leider die meisten SNMP-Managementsysteme und die Gerätehersteller etwas schwer. Das Web-Management kann auch in Teilen kompatibel zu SNMP sein, indem der Browser intern die MIB dieses Geräts modifiziert oder abfragt und indem SNMP-Traps verwendet werden, um Alarmierungen abzusetzen.

Allerdings ist das HTTP (Hyper Text Transfer Protocol), genau wie SNMP, alles andere als sicher: Man kann es relativ problemlos mitlesen oder verändern. Die Lösung dieses Sicherheitsproblems ist auch hier die Verwendung des SSL-Protokolls, das speziell zur sicheren Browser-Server-Kommunikation entwickelt wurde, oder – noch sicherer – der Einsatz von IPSec. Bei Einsatz und Kombination der richtigen, sicheren Übertragungsprotokolle steht also einem zeitgemäßen webbasierten Management der VPN-Komponenten nichts mehr im Wege.

Überwachung

Auch für die Überwachung von traditionellen Netzwerkkomponenten hat sich neben dem SNMP-Protokoll zunehmend das webbasierte Management durchgesetzt. Unter Überwachung versteht man in diesem Zusammenhang die gezielte, durch eine Managementstation oder einen Browser ausgelöste Abfrage von Systemparametern und das Versenden so genannter Traps, die von einem Netzwerksystem aufgrund eines kritischen Systemzustands ausgelöst wurden.

Das gezielte Abfragen von MIB-Variablen in den Netzwerkkomponenten muss sowohl durch SNMP-Get-Befehle als auch durch einen Browser unterstützt werden. Auch hier gilt es Vorsicht walten zu lassen, falls man damit auch sicherheitsrelevante Parameter abfragt, denn jedermann, der Zugriff auf die Netzwerkverbindung zwischen VPN-Komponente und Managementstation hat, kann SNMP und HTTP mitlesen! Der Einsatz von sicheren Übertragungsprotokollen ist auch hier dringend angeraten.

Die Erzeugung von SNMP-Traps zur Signalisierung von Alarmzuständen muss in jedem Fall unterstützt werden. Hier werden keine sicherheitsrelevanten Parameter ungeschützt übertragen. Es wird lediglich der Managementstation mitgeteilt, welches Problem wo aufgetreten ist. Spezielle sichere Übertragungsprotokolle sind hierfür nicht notwendig.

Aber ein anderer Punkt verdient in diesem Zusammenhang Beachtung: SNMP definiert im Standard nur eine kleine Anzahl von Traps. Will man aber Traps erzeugen, die bei Eindringversuchen in das Netzwerk oder bei anderen einsatzspezifischen Vorfällen oder Situationen generiert werden, muss das VPN-System eine derart erweiterte SNMP-Funktionalität bieten. Generell sollte es möglich sein, für alle kritischen Zustände – und damit sind nicht nur Hardwareprobleme gemeint – die Generierung eines Traps konfigurieren zu können.

Accounting

Während in lokalen Netzwerken das Accounting, also die auf Benutzer, Organisationseinheiten oder Kostenstellen heruntergebrochene Abrechnung von Netzwerkdiensten, fast niemals eingesetzt wird, ist dies im WAN-Bereich immer öfter und beim Remote Access fast immer zu finden. Dies hat eine Reihe von Gründen. Im LAN-Bereich fallen keine zusätzlichen Übertragungskosten an, man kann eine Weiterverrechnung auf Basis von Netzwerkports, Plattenplatz, Geschwindigkeiten usw. durchführen. Des Weiteren bieten die meisten LAN-Komponenten zurzeit nur sehr rudimentäre oder gar keine Accounting-Möglichkeiten, so dass nicht selten die notwendigen technischen Möglichkeiten gar nicht gegeben sind.

Im Bereich der öffentlichen Netze addieren sich zu den Kosten für das Equipment und den Betrieb die teilweise recht hohen Gebühren, die an die Carrier auf zeit- oder volumenabhängiger Basis zu entrichten sind. Diese Kosten will man in der Regel nicht in einem festen Verhältnis auf die Kostenstellen umlegen, da nicht jede Organisationseinheit die Dienste in gleichem Maße in Anspruch nimmt. Es sollen vielmehr die tatsächlich verursachten Kosten ermittelt und weiter verrechnet werden. Somit findet man insbesondere im Remote-Access-Bereich Accounting-Funktionalitäten, die eine zeit- und volumenabhängige Abrechnung der Dienste ermöglichen. Auch im Bereich der traditionellen WAN-Router findet man zunehmend Abrechnungsfunktionen auf IP-Ebene.

Löst man diese herkömmlichen Strukturen nun durch die VPN-Technologie ab, müssen diese Accounting-Systeme ebenfalls unterstützt werden, auch im parallelen Betrieb während der Migrationsphase. RADIUS (Remote Authentication Dial In User Service, ein Standard, der Protokolle und Funktionen zur Authentifizierung, Autorisierung und zum Accounting von Remote-Access-Benutzern beschreibt) hat sich im Laufe der Jahre zu einem weit verbreiteten Standard mit hoher Interoperabilität zwischen verschiedenen Herstellern entwickelt. Praktisch alle Remote-Access-Systeme, sowohl im Enterprise- als auch im Carrier-Umfeld, arbeiten mit RADIUS zur Benutzer-Authentifizierung und zum Accounting. Um die Datensätze auszuwerten, die von RADIUS-Servern erzeugt werden, gibt es mittlerweile eine Fülle von Programmen. Das Format eignet sich auch zum direkten Importieren in verbreitete Applikationen wie Datenbanken oder Tabellenkalkulationsprogramme.

Aus diesen Gründen liegt es auf der Hand, dass VPN-Systeme RADIUS-Accounting unterstützen müssen. Dadurch können sie einfach und schnell in eine bestehende Infrastruktur integriert werden.

3.6.2 Sicherheit

Sicherheitsstrategie

Ein VPN muss sich reibungslos in die Sicherheitsstrategie einer Organisation integrieren lassen. Die Sicherheitsstrategie, auch als *Security Policy* bezeichnet, ist eine unternehmensweite Definition von Sicherheitsanforderungen, die eine ganze Reihe von Bereichen abdecken müssen, wie zum Beispiel die Behandlung von Kennwörtern, physikalischer Zugangsschutz, Zugriffsregelung auf kritische Ressourcen usw. Sie beschreibt dabei sowohl die Anforderungen als auch die Zuständigkeitsbereiche für deren Umsetzung. Eine gute Sicherheitsstrategie definiert dabei nicht, welche Technologien eingesetzt werden, oder bestimmte Schlüssellängen, dies ist Aufgabe der für die Umsetzung verantwortlichen Organisationseinheiten. Die Security Policy legt lediglich fest, welche Daten über welchen Zeitraum, vor welchen potenziellen Angriffen und durch wen zu schützen sind. Dies klingt einfach und ist es auch. Was komplex ist, ist ihre konsistente Umsetzung.

Eine Sicherheitsstrategie legt zum Beispiel fest, dass die Konstruktionsdaten eines Unternehmens mindestens zwanzig Jahre auch vor Zugriffen durch fremde Geheimdienste sicher sein müssen, und benennt optional noch die dafür verantwortlichen Organisationseinheiten. Das war es auch schon, denn die Umsetzung obliegt den betroffenen Organisationseinheiten, die mit diesen Daten umgehen. Im Netzwerkbereich leiten sich daraus die entsprechenden Kriterien ab, um die Übertragungen sowohl im Intranet als auch im Weitverkehrsnetz oder VPN durch starke Verschlüsselung zu schützen.

Authentifizierung

Üblicherweise sind für Benutzer, die sich per Remote Access in das Unternehmensnetzwerk einwählen, in der Sicherheitsstrategie bestimmte Anforderungen hinsichtlich der Authentifizierung festgelegt. Im lokalen Netzwerk kann man sich ohne besondere Authentifizierung anschließen, denn hier sieht die Security Policy des Unternehmens in der Regel einen physikalischen Zugangsschutz vor, meist per Magnetkarte am Eingang der Gebäude. Dieser Zugangsschutz entfällt beim Remote Access natürlich und muss durch besondere Verfahren zur Identitätsfeststellung nachgebildet werden. Dies trifft auf den traditionellen Remote Access und Remote-Access-VPNs gleichermaßen zu.

Firewalls

In den meisten Organisationen wird heute ein Internetanschluss betrieben. Da man aus Sicherheits- und anderen Gründen (z.B. wegen der Verwendung privater, nicht registrierter IP-Adressen) sein Intranet nicht direkt mit dem Internet verbinden kann, wird diese Funktion in der Regel von so genannten Firewalls vorgenommen.

Eine Firewall kontrolliert den gesamten Datenverkehr zwischen Systemen im Internet und dem Intranet. Sie soll einerseits verhindern, dass sich Unbefugte Zugang zum Intranet verschaffen, und andererseits den Datenverkehr kontrollieren, den Systeme im Intranet in das Internet leiten. Hierbei werden die Datenpakete sogar bis auf die Ebene der Dateninhalte geprüft, um zum Beispiel Webseiten mit moralisch fragwürdigem Inhalt zu sperren. Der wichtige Punkt hierbei ist, dass eine Firewall den Verkehr in das Internet kontrolliert.

Ein Internet-VPN greift hingegen auf kein einziges Internet-System zu, sondern es benutzt das Internet lediglich als Transportmedium für IP-Pakete. Dies ist eine völlig andere Situation, da die IP-Pakete ausschließlich im privaten Bereich versendet und empfangen werden. Querverbindungen in das Internet sind bei einigen sehr guten, dedizierten VPN-Systemen überhaupt nicht möglich.

Aus diesem Grund macht es auch wenig Sinn (und kann von Fall zu Fall sogar Probleme bereiten), wenn man beide Funktionalitäten auf einem System zusammenführt, egal ob man VPN-Funktionen in eine Firewall integriert oder umgekehrt eine Firewall in einen VPN-Konzentrator. Vielmehr ist zu fordern, dass der VPN-Konzentrator als Ablösung klassischer Netzwerkkomponenten wie Router oder Remote-Access-Konzentratoren und die Firewall gleichzeitig nebeneinander zu betreiben sind.

PKI

Technologien wie E-Business oder E-Commerce und neu geschaffene rechtliche Rahmenbedingungen für digitale Unterschriften haben zusammen mit einem gestiegenen Sicherheitsbedürfnis bei der Benutzung öffentlicher Netze eine Reihe von Applikationen hervorgebracht, die mit elektronischen, digitalen Schlüsseln Daten verschlüsseln oder Dokumente signieren. Bei großen Organisationen fallen dabei eine ganze Reihe verschiedener Schlüssel an, die erzeugt, verwaltet, regelmäßig erneuert und vor allem eindeutig bestimmten Personen zugeordnet werden müssen.

Eine Public-Key-Infrastruktur (PKI) leitet ihren Namen von der so genannten Public-Key-Kryptographie (vgl. Kapitel 4) ab, bei der mit zwei verschiedenen Schlüsseln, einem öffentlichen und einem privaten, gearbeitet wird. Mit einem Schlüssel (Public Key, dem öffentlichen Schlüssel) wird verschlüsselt und mit dem anderen (Private Key, dem privaten, geheimen Schlüssel) wieder

entschlüsselt. Die beiden Schlüssel bilden ein Paar; der öffentliche Schlüssel ist jedermann zugänglich und ist fest an eine Person gebunden, und der private wird von der Person geheim gehalten. Etwas, das mit dem einen Schlüssel verschlüsselt wurde, kann nur mit dem korrespondierenden anderen Schlüssel wieder entschlüsselt werden.

Die Bindung einer Person an einen öffentlichen Schlüssel erfolgt über ein so genanntes digitales Zertifikat. Die Erzeugung und Verwaltung dieser Zertifikate, die Registrierung der Benutzer und das Erzeugen und Speichern von Schlüsselpaaren zum Zweck der Datenverschlüsselung auf Rechnern ist die Hauptfunktion einer PKI. Diese kann sowohl von einem Unternehmen selbst betrieben werden, oder man benutzt eine öffentliche PKI.

Falls nun digitale Zertifikate zum digitalen Signieren eingesetzt werden, liegt es auf der Hand, diesen Mechanismus auch zur Authentifizierung von Benutzern oder VPN-Systemen einzusetzen. Hier sollten die VPN-Geräte und die Clientsoftware eine entsprechende Unterstützung bieten.

3.7 Koexistenz zu traditionellen WAN-Strukturen

Üblicherweise findet nicht immer sofort eine Ablösung von traditionellen Netzen durch virtuelle private Netze statt. Vielmehr müssen beide für eine gewisse Migrationsphase gleichzeitig zu betreiben sein. Typische Fälle sind der gleichzeitige Einsatz von VPNs und Frame-Relay- oder ATM-Netzen, da es oft der Fall ist, dass die beiden letztgenannten Technologien nicht an jedem benötigten Standort zur Verfügung stehen.

Auch normaler Remote Access per Einwahl über das Telefonnetz wird sehr oft parallel zum VPN-Remote-Access betrieben. Hier sollen aus Kostengründen zwar möglichst viele Verbindungen über das VPN laufen, jedoch nutzt man den Einwahlteil oft als Backup oder für bestimmte kritische Applikationen, die eine garantierte Einwahl und bestimmte Qualitätsmerkmale benötigen.

Um das Ratiopotenzial eines VPN voll ausschöpfen zu können, dürfen Einsparungen, die auf der einen Seite durch geringere Gebühren und günstigeres Equipment erzielt wurden, nicht durch andere, durch die VPN-Technologie anfallende, neue Kosten zunichte gemacht werden. Solche Kosten sind oft versteckt und auf den ersten Blick vielleicht gar nicht mit dem Einsatz eines VPN in Zusammenhang zu bringen.

Einer der negativen Kostenfaktoren, vielleicht der schwerwiegendste, sind erhöhte Managementkosten. Diese können erzeugt werden, wenn sich die VPN-Systeme nicht in das vorhandene Netzwerkmanagement integrieren lassen und ein paralleles Management aufgebaut wird.

Auch Zusatzkosten durch möglicherweise erforderliche IP-Adressänderungen in einem großen Netzwerk sollten vermieden werden, indem man die VPN-Technologie in bestehende Strukturen integriert und nicht umgekehrt bestehende Strukturen an das VPN anpasst.

3.8 Adressmanagement

Das IP-Adressmanagement ist in vielen Unternehmen zunehmend zu einem kritischen Faktor geworden. Dies hat historische Gründe, die mit der Entwicklung und Einführung des IP-Protokolls zu tun haben.

IP war viele Jahre ein Protokoll, das fast ausschließlich im Internet eingesetzt wurde. Die Folge war, dass alle beteiligten Geräte offizielle IP-Adressen und etliche Organisationen und Firmen offizielle IP-Netze zugeteilt bekamen. Als TCP/IP praktisch Bestandteil von Unix wurde, es als kostenloser Zusatz für Windows 3.x verfügbar war und Betriebssysteme wie VM oder OS/400 von IBM ebenfalls damit ausgerüstet wurden, gab es ein paar Entwicklungen, an denen viele Unternehmen heute noch zu knabbern haben. Denn in vielen Firmen begannen sich regelrechte »IP-Inseln« zu bilden, oft mit selbst vergebenen, nicht registrierten Adressbereichen. Zuerst wollte man überhaupt keinen Internetzugriff, im Zeitalter des World Wide Web wurde dieser Wunsch von den Anwendern jedoch immer dringlicher vorgetragen. Als Ausweg boten sich Firewalls mit NAT (Network Address Translation, Umsetzung von nicht registrierten in offizielle IP-Adressen) an. Im Intranet der Unternehmen wurden bis dato die IP-Adressen manuell vergeben und statisch auf den Systemen eingetragen. Dokumentiert wurde meist nichts, und wenn doch, dann waren die Daten oft nicht aktuell und unvollständig. Im Verlauf dieser und anderer Entwicklungen wurde der Ruf nach einem einfachen und automatischen Management der IP-Adressen immer lauter.

Das Resultat war das DHCP (Dynamic Host Configuration Protocol), ein Verfahren mit dem IP-Schnittstellen automatisch und dynamisch von einem speziellen Server konfiguriert werden. DHCP ist eine Client-Server-Applikation, bei der die Clients beim Start eines Rechners mit einer IP-Schnittstelle von einem DHCP-Server eine IP-Adresse und weitere optionale Konfigurationsparameter anfordern. Es hat mittlerweile eine recht breite Verwendung gefunden. Neben Implementierungen in PC-Betriebssystemen gibt es auch professionelle, hoch skalierbare Systeme, die mit relationalen Datenbanken arbeiten, und einen dynamischen DNS-Server (Domain Name System, ein Verfahren, um Systeme statt mit ihrer IP-Adresse über einen Namen anzusprechen) integriert haben, der die sich ändernden IP-Adressen immer dem korrekten Domainnamen zuordnet.

3 Anforderungen an VPNs

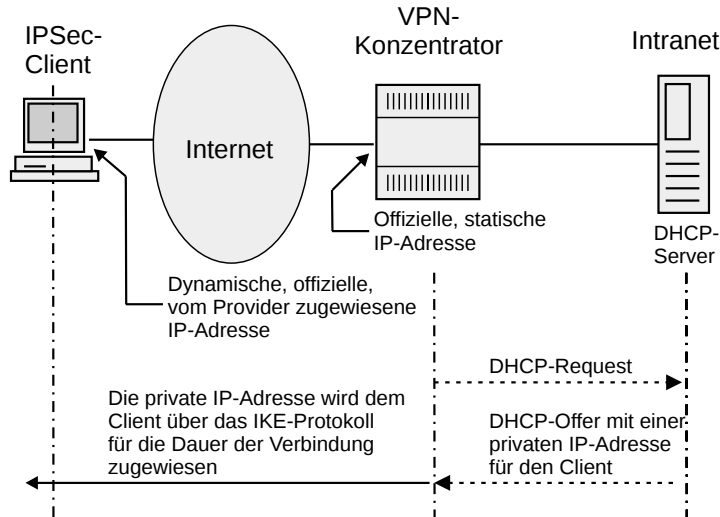


Abbildung 3.7: Das IP-Adressmanagement in IP-VPNs muss gut geplant werden.

Viele heute eingesetzte Remote-Access-Konzentratoren verwenden DHCP, um den Clientrechnern eine IP-Adresse dynamisch zuzuweisen. Diese Funktionalität muss ein VPN-Konzentrator ebenfalls unterstützen, um in das unternehmensweite IP-Adressmanagement integrierbar zu sein. Idealerweise ist sowohl ein gerichtetes DHCP möglich, bei dem sich die Systeme an dedizierte DHCP-Server wenden, als auch ein ungerichtetes, bei dem per Netzwerkbroadcast ein passender Server gesucht wird.

Die Anforderungen im Bereich der IP-Adressverwaltung beschränken sich jedoch keinesfalls nur auf DHCP, es müssen auch andere Mechanismen zur Zuteilung von IP-Adressen unterstützt werden. Hierbei resultieren die Anforderungen sowohl aus der aktuellen Managementumgebung als auch aus zukünftig einzusetzenden Systemen. Insbesondere arbeiten viele Remote-Access-Konzentratoren mit RADIUS-Servern, auf denen neben der Authentifizierung und dem Accounting auch IP-Adressen userbezogen oder aus einem so genannten Pool heraus vergeben werden. Auch eine gerätespezifische Vergabe von IP-Adressen auf Basis von Benutzernamen oder aus einem Pool sollte von einem modernen VPN-Konzentrator unterstützt werden.

Directory Services (Verzeichnisdienste) erfreuen sich in der Fachwelt einer immer größeren Beliebtheit. Nachdem in der Vergangenheit einige proprietäre Verfahren wie Novell NDS oder Banyan Streetwork und Standards wie X.500 miteinander konkurriert hatten, konzentriert sich die Entwicklung momentan auf das Lightweight Directory Access Protocol (LDAP). Es wird mittlerweile von allen Größen der Netzwerkindustrie sowohl im Bereich der Betriebssysteme als auch im Bereich der Netzwerkkomponenten und Managementsys-

teme eingesetzt. LDAP ist sehr flexibel und bietet einen fest spezifizierten Satz von Funktionen und Formaten zum Erzeugen, Modifizieren, Abfragen und Löschen von Einträgen in der Verzeichnisdatenbank. Die Struktur eines Verzeichnisses ist frei wählbar und kann somit maßgerecht an unterschiedliche Bedürfnisse angepasst werden.

Ein VPN-System, das auch noch in der Zukunft einsetzbar sein soll, muss daher ebenfalls LDAP unterstützen können, um in unternehmensweite Verzeichnisdienste integrierbar zu sein.

3.9 Interoperabilität

Interoperabilität ist eine wichtige und kritische Anforderung an heutige VPN-Konzentratoren und VPN-Gateways. Dies hat gleich mehrere Gründe. Je nach Auswahl eines geeigneten Tunneling-Modells (vgl. Kapitel 6) können mehrere Partner an einem VPN beteiligt sein (zum Beispiel ein Internet Service Provider und der Endkunde), die unter Umständen Equipment verschiedener Hersteller einsetzen.

Ein anderes Beispiel für die Notwendigkeit der Interoperabilität verschiedener VPN-Systeme ist die Dynamik im Bereich von Firmenzusammenschlüssen oder Kooperationen, die in der heutigen Zeit zu beobachten ist. Hierbei können in den unterschiedlichen Unternehmen Anforderungen der Art entstehen, dass die eingesetzten VPN-Geräte miteinander kommunizieren müssen. Das *Automotive Network Exchange* (ANX) ist beispielsweise ein VPN, das für Automobilhersteller und Zulieferer mit unterschiedlichsten Geräten auf der Basis des IPSec-Protokolls betrieben wird.

Dies erfordert, dass das jeweils verwendete Equipment mit den Gegenstellen interoperabel ist. Die Lösung besteht darin, ausschließlich solche Protokolle einzusetzen, die standardisiert und allgemein akzeptiert sind. Leider hat die Vergangenheit gezeigt, dass auch standardkonforme Implementierungen bestimmter Datenkommunikationsprotokolle keineswegs ein Garant für eine reibungslose Kommunikation sind. Dies liegt an der Tatsache, dass viele Protokolle sehr flexibel sein müssen und gegebenenfalls so konfiguriert werden können, dass eine Kommunikation mit anders konfigurierten Gegenstellen unmöglich ist.

Wichtig ist daher, dass es eine oder mehrere mögliche Konfigurationen der Protokolle gibt, bei denen eine Interoperabilität gewährleistet oder sogar nachgewiesen ist. Insbesondere im Bereich des IPSec-Protokolls ist dies zu fordern. ICSA.net ist eine sich selbst als unabhängig bezeichnende Firma, die sich mit Internet-Sicherheit beschäftigt. Die ICSA Labs führen Tests durch, die die Funktionalität und Interoperabilität von IPSec-Implementierungen ermitteln sollen. Geräte oder Programme, die erfolgreich getestet wurden, dürfen

3 Anforderungen an VPNs

das ICSA-IPSec-Logo tragen und sind, was viel wichtiger ist, mit einer sehr hohen Wahrscheinlichkeit zu anderen zertifizierten Systemen interoperabel.