



Redes de computadores e a Internet

A Internet de hoje é provavelmente o maior sistema de engenharia já criado pela humanidade, com centenas de milhões de computadores conectados, enlaces de comunicação e nós de comutação; bilhões de usuários que se conectam por meio de *notebooks*, *tablets* e *smartphones*; e com uma série de “coisas” conectadas à Internet, incluindo console para jogos, sistemas de vigilância, relógios, óculos, termostatos e automóveis. Dado que a Internet é tão ampla e possui inúmeros componentes e utilidades, há a possibilidade de compreender como ela funciona? Existem princípios de orientação e estrutura que forneçam um fundamento para a compreensão de um sistema surpreendentemente complexo e abrangente? Se a resposta for sim, é possível que, nos dias de hoje, seja interessante e divertido aprender sobre rede de computadores? Felizmente, as respostas para todas essas perguntas é um retumbante SIM! Na verdade, nosso objetivo neste livro é fornecer uma introdução moderna ao campo dinâmico das redes de computadores, apresentando os princípios e o entendimento prático necessários para utilizar não apenas as redes de hoje, como também as de amanhã.

O primeiro capítulo apresenta um panorama de redes de computadores e da Internet. Nosso objetivo é pintar um quadro amplo e estabelecer um contexto para o resto deste livro, para ver a floresta a partir das árvores. Cobriremos um terreno bastante extenso neste capítulo de introdução e discutiremos várias peças de uma rede de computadores, sem perder de vista o quadro geral.

O panorama geral de redes de computadores que apresentaremos neste capítulo será estruturado como segue. Após apresentarmos brevemente a terminologia e os conceitos fundamentais, examinaremos primeiro os componentes básicos de *hardware* e *software* que compõem uma rede. Partiremos da periferia da rede e examinaremos os sistemas finais e aplicações de rede executados nela. Consideraremos os serviços de transporte fornecidos a essas aplicações. Em seguida, exploraremos o núcleo de uma rede de computadores examinando os enlaces e os nós de comutação que transportam dados, bem como as redes de acesso e os meios físicos que conectam sistemas finais ao núcleo da rede. Aprenderemos que a Internet é uma rede de redes, e observaremos como essas redes se conectam umas com as outras.

Depois de concluirmos essa revisão sobre a periferia e o núcleo de uma rede de computadores, adotaremos uma visão mais ampla e mais abstrata na segunda metade deste capítulo. Examinaremos atraso, perda e vazão de dados em uma rede de computadores e

forneceremos modelos quantitativos simples para a vazão e o atraso fim a fim: modelos que levam em conta atrasos de transmissão, propagação e enfileiramento. Depois, apresentaremos alguns princípios fundamentais de arquitetura em redes de computadores, a saber: protocolos em camadas e modelos de serviço. Aprenderemos, também, que as redes de computadores são vulneráveis a diferentes tipos de ameaças; analisaremos algumas dessas ameaças e como a rede pode se tornar mais segura. Por fim, encerraremos este capítulo com um breve histórico da computação em rede.

1.1 O QUE É A INTERNET?

Neste livro, usamos a Internet pública, uma rede de computadores específica, como o veículo principal para discutir as redes de computadores e seus protocolos. Mas o que é a Internet? Há diversas maneiras de responder a essa questão. Primeiro, podemos descrever detalhadamente os aspectos principais da Internet, ou seja, os componentes de *software* e *hardware* básicos que a formam. Segundo, podemos descrever a Internet em termos de uma infraestrutura de redes que fornece serviços para aplicações distribuídas. Iniciaremos com a descrição dos componentes, utilizando a Figura 1.1 como ilustração para a nossa discussão.

1.1.1 Uma descrição dos componentes da rede

A Internet é uma rede de computadores que interconecta bilhões de dispositivos de computação ao redor do mundo. Há pouco tempo, esses dispositivos eram basicamente PCs (do inglês *personal computers*) de mesa, estações de trabalho Linux e os chamados servidores que armazenam e transmitem informações, como páginas da Web e mensagens de *e-mail*. No entanto, cada vez mais usuários se conectam à Internet com *smartphones* e *tablets* – hoje, cerca de metade da população mundial é composta por usuários ativos de Internet móvel, e espera-se que essa porcentagem salte para 75% até 2025 (Statista, 2019). Além disso, “coisas” não tradicionais conectadas à Internet, como TVs, consoles de videogame, termostatos, sistemas de segurança doméstica, eletrodomésticos, relógios, óculos, automóveis, sistemas de controle de trânsito e outras estão sendo conectadas à rede. Na verdade, o termo *rede de computadores* está começando a soar um tanto desatualizado, em razão dos muitos equipamentos não tradicionais que estão sendo conectados à Internet. No jargão da área de redes, todos esses equipamentos são denominados **hospedeiros** ou **sistemas finais**. Estima-se que havia cerca de 18 bilhões de dispositivos conectados à Internet em 2017, um número que chegará a 28,5 bilhões até 2022 (Cisco VNI, 2020).

Sistemas finais são conectados entre si por **enlaces (links) de comunicação e comutadores (switches) de pacotes**. Na Seção 1.2, veremos que há muitos tipos de enlaces de comunicação, que são constituídos de diferentes tipos de meios físicos, entre eles cabos coaxiais, fios de cobre, fibras ópticas e ondas de rádio. Enlaces diferentes podem transmitir dados em taxas diferentes, sendo a **taxa de transmissão** de um enlace medida em *bits* por segundo. Quando um sistema final possui dados para enviar a outro sistema final, o sistema emissor segmenta esses dados e adiciona *bytes* de cabeçalho a cada segmento. Os blocos de informação resultantes, conhecidos como **pacotes** no jargão de rede de computadores, são enviados através da rede ao sistema final de destino, onde são remontados na forma dos dados originais.

Um nó de comutação de pacotes encaminha o pacote que está chegando em um de seus enlaces de comunicação de entrada para um de seus enlaces de comunicação de saída. Há comutadores de pacotes de todos os tipos e formas, mas os dois mais proeminentes na Internet de hoje são **roteadores** e **switches**. Esses dois tipos de nós de comutação encaminham pacotes a seus destinos finais. Os *switches* geralmente são utilizados em redes locais, enquanto os roteadores são utilizados principalmente na parte mais interna da rede.

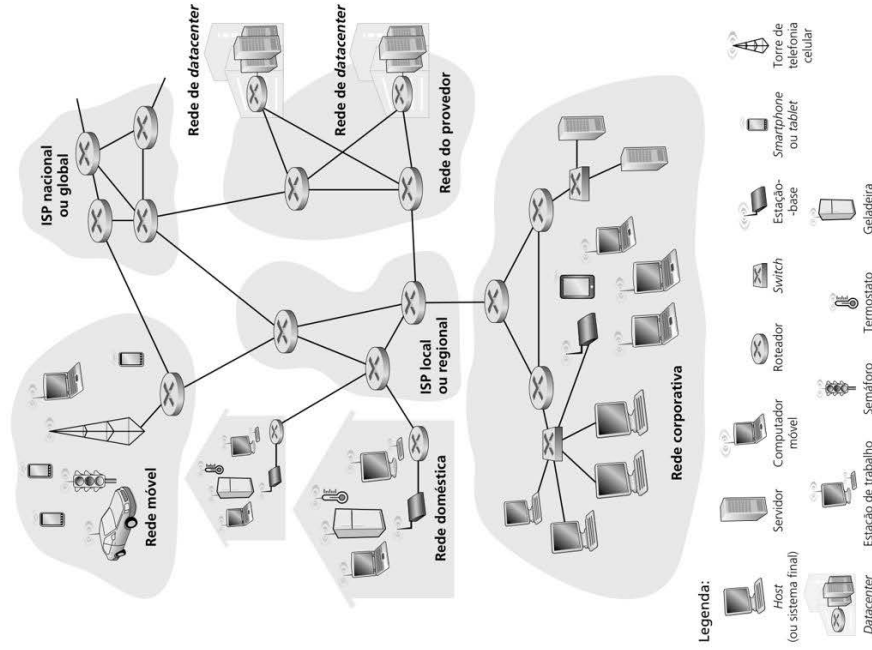


Figura 1.1 Alguns componentes da Internet.

A sequência de enlaces de comunicação e nós de comutação de pacotes que um pacote percorre desde o sistema final remetente até o sistema final receptor é conhecida como **rota** ou **caminho** através da rede. A Cisco prevê que o tráfego de IP global anual atingirá cerca de cinco *zettabytes* (1.021 *bytes*) até 2022 (Cisco VNI, 2020).

As redes comutadas por pacotes (que transportam pacotes) são, de muitas maneiras, semelhantes às redes de transporte de rodovias, estradas e cruzamentos (que são percorridas por veículos). Considere, por exemplo, uma fábrica que precise transportar uma quantidade de carga muito grande a algum depósito localizado a milhares de quilômetros. Na fábrica, a carga é dividida e carregada em uma frota de caminhões. Cada caminhão viaja, de modo independente, pela rede de rodovias, estradas e cruzamentos ao depósito de destino. No depósito, a carga é descarregada e agrupada com o resto da carga pertencente à mesma

remessa. Deste modo, os pacotes se assemelham aos caminhões, os enlaces de comunicação representam rodovias e estradas, os nós de comutação seriam os cruzamentos, e cada sistema final se assemelha aos depósitos. Assim como o caminhão faz o percurso pela rede de transporte, o pacote utiliza uma rede de computadores.

Sistemas finais acessam a Internet por meio de **ISPs**, (do inglês *Internet Service Providers* – **Provedores de Serviços de Internet**), entre eles: ISPs residenciais como empresas de TV a cabo ou empresas de telefonia; corporativos, de universidades e que fornecem acesso sem fio em aeroportos, hotéis, cafés e outros locais públicos; e ISPs de dados móveis, que oferecem acesso aos nossos *smartphones* e a outros dispositivos. Cada ISP é uma rede de nós de comutação e enlaces de comunicação. ISPs oferecem aos sistemas finais uma variedade de tipos de acesso à rede, incluindo acesso residencial de banda larga como *modem* a cabo ou DSL (do inglês *digital subscriber line* – linha digital de assinante), acesso por LAN de alta velocidade e acesso sem fio móvel. Os ISPs também fornecem acesso a provedores de conteúdo, conectando servidores diretamente à Internet. A Internet trata fundamentalmente da conexão entre os sistemas finais; portanto, os ISPs que fornecem acesso a esses sistemas também devem se interconectar. Esses ISPs de nível mais baixo são interconectados por meio de ISPs de nível mais alto, nacionais e internacionais. Um ISP de nível mais alto consiste em roteadores de alta velocidade interconectados com enlaces de fibra ótica de alta velocidade. Cada rede ISP, seja de nível mais alto ou mais baixo, é gerenciada de forma independente, executa o protocolo IP (ver adiante) e obedece a certas convenções de nomeação e endereçamento. Examinaremos ISPs e sua interconexão com mais detalhes na Seção 1.3.

Os sistemas finais, os nós de comutação e outras peças da Internet executam **protocolos** que controlam o envio e o recebimento de informações. O **TCP** (do inglês *Transmission Control Protocol* – **Protocolo de Controle de Transmissão**) e o **IP** (do inglês *Internet Protocol* – **Protocolo da Internet**) são dois dos mais importantes da Internet. O protocolo IP especifica o formato dos pacotes que são enviados e recebidos entre roteadores e sistemas finais. Os principais protocolos da Internet são conhecidos como **TCP/IP**. Começaremos a examinar protocolos neste capítulo de introdução. Mas isso é só um começo – grande parte deste livro trata de protocolos de redes de computadores!

Dada a importância de protocolos para a Internet, é adequado que todos concordem sobre o que cada um deles faz, de modo que as pessoas possam criar sistemas e produtos que operem entre si. É aqui que os padrões entram em ação. **Padrões da Internet** são desenvolvidos pela IETF (do inglês *Internet Engineering Task Force* – Força de Trabalho de Engenharia da Internet) (IETF, 2020). Os documentos padronizados da IETF são denominados **RFCs** (do inglês *Request For Comments* – **pedido de comentários**). Os RFCs começaram como solicitações gerais de comentários (daí o nome) para resolver problemas de arquitetura que a precursora da Internet enfrentava (Allman, 2011). Os RFCs costumam ser bastante técnicos e detalhados. Definem protocolos como TCP, IP, HTTP (para a Web) e SMTP (para *e-mail*). Hoje, existem quase 9.000 RFCs. Outros órgãos também especificam padrões para componentes de rede, principalmente para enlaces. O IEEE 802 LAN Standards Committee (IEEE 802, 2020), por exemplo, especifica os padrões Ethernet e WiFi sem fio.

1.1.2 Uma descrição do serviço

A discussão anterior identificou muitos dos componentes que compõem a Internet. Mas também podemos descrevê-la partindo de um ângulo completamente diferente – ou seja, como *uma infraestrutura que provê serviços a aplicações*. Além de aplicações tradicionais como correio eletrônico e navegação na Web, as aplicações de Internet incluem aplicativos para *smartphones* e *tablets*, incluindo serviços de mensagem instantânea, mapeamento com informações em tempo real sobre condições de trânsito, *streaming* de músicas, filmes e televisão, mídias sociais on-line, videoconferência, *jogos multiplayer* e sistemas de recomendação baseados em localização. Essas aplicações são conhecidas como **aplicações distribuídas**, uma vez que envolvem diversos sistemas finais que trocam informações mutuamente.

De forma significativa, as aplicações da Internet são executadas em sistemas finais – e não em nós de comutação no núcleo da rede. Embora os nós de comutação facilitem a troca de dados entre os sistemas finais, eles não estão relacionados com a aplicação, que é a origem ou o destino dos dados.

Vamos explorar um pouco mais o significado de uma infraestrutura que fornece serviços a aplicações. Nesse sentido, suponha que você tenha uma grande ideia para uma aplicação distribuída para a Internet, uma que possa beneficiar bastante a humanidade ou que simplesmente o enriqueça e o torne famoso. Como transformar essa ideia em uma aplicação real da Internet? Já que as aplicações são executadas em sistemas finais, você precisará criar programas que sejam executados em sistemas finais. Você poderia, por exemplo, criar seus programas em Java, C ou Python. Agora, já que você está desenvolvendo uma aplicação distribuída para a Internet, os programas executados em diferentes sistemas finais precisarão enviar dados uns aos outros. E, aqui, chegamos ao assunto principal – o que leva ao modo alternativo de descrever a Internet como uma plataforma para aplicações. De que modo um programa, executado em um sistema final, orienta a Internet a enviar dados a outro programa executado em outro sistema final?

Os sistemas finais ligados à Internet oferecem uma **interface socket** que especifica como o programa que é executado no sistema final solicita à infraestrutura da Internet que envie dados a um programa de destino específico, executado em outro sistema final. Essa interface *socket* da Internet é um conjunto de regras que o *software* emissor deve cumprir para que a Internet seja capaz de enviar os dados ao programa de destino. Discutiremos a interface *socket* da Internet mais detalhadamente no Capítulo 2. Agora, vamos traçar uma simples comparação, que será utilizada com frequência neste livro. Suponha que Alice queira enviar uma carta para Bob utilizando o serviço postal. Alice, é claro, não pode apenas escrever a carta (os dados) e atirá-la pela janela. Em vez disso, o serviço postal necessita que ela coloque a carta em um envelope; escreva o nome completo de Bob, endereço e CEP no centro do envelope; feche; coloque um selo no canto superior direito; e, por fim, leve o envelope a uma agência de correio oficial. Dessa maneira, o serviço postal possui sua própria “interface de serviço postal”, ou conjunto de regras, que Alice deve cumprir para que sua carta seja entregue a Bob. De modo semelhante, a Internet possui uma interface *socket* que o *software* emissor de dados deve seguir para que a Internet envie os dados para o *software* receptor.

O serviço postal, claro, fornece mais de um serviço a seus clientes: entrega expressa, aviso de recebimento, carta simples e muitos mais. De modo semelhante, a Internet provê diversos serviços a suas aplicações. Ao desenvolver uma aplicação para a Internet, você também deve escolher um dos serviços que a rede oferece. Uma descrição dos serviços será apresentada no Capítulo 2.

Acabamos de apresentar duas descrições da Internet: uma delas diz respeito a seus componentes de *hardware* e *software*, e a outra, aos serviços que ela oferece a aplicações distribuídas. Mas talvez você ainda esteja confuso sobre o que é a Internet. O que é comutação de pacotes e TCP/IP? O que são roteadores? Quais tipos de enlaces de comunicação estão presentes na Internet? O que é uma aplicação distribuída? Como um termostato ou uma balança podem ser ligados à Internet? Se você está um pouco assustado com tudo isso agora, não se preocupe – a finalidade deste livro é lhe apresentar os mecanismos da Internet e também os princípios que determinam como e por que ela funciona. Explicaremos esses termos e questões importantes nas seções e nos capítulos subsequentes.

1.1.3 O que é um protocolo?

Agora que já entendemos um pouco sobre o que é a Internet, vamos considerar outra palavra fundamental usada em redes de computadores: *protocolo*. O que é um protocolo? O que um protocolo faz?

Uma analogia humana

Talvez seja mais fácil entender a ideia de um protocolo de rede de computadores considerando primeiro algumas analogias humanas, já que seguimos protocolos o tempo todo. Considere o que você faz quando quer perguntar as horas a alguém. Um diálogo comum é ilustrado na Figura 1.2. O protocolo humano (ou as boas maneiras, ao menos) dita que, ao iniciarmos uma comunicação com outra pessoa, primeiro a cumprimentamos (o primeiro “Oi” da Figura 1.2). A resposta comum para um “Oi” é um outro “Oi”. Implicitamente, tomamos a resposta cordial “Oi” como uma indicação de que podemos prosseguir e perguntar as horas. Uma reação diferente ao “Oi” inicial (tal como “Não me perturbe!”; “Eu não falo português!” ou alguma resposta atravessada) poderia indicar falta de vontade ou incapacidade de comunicação. Nesse caso, o protocolo humano seria não perguntar que horas são. Às vezes, não recebemos nenhuma resposta para uma pergunta, caso em que em geral desistimos de perguntar as horas ao interlocutor. Note que, no nosso protocolo humano, *há mensagens específicas que enviamos e ações específicas que realizamos em reação às respostas recebidas ou a outros eventos* (como nenhuma resposta após certo tempo). É claro que mensagens transmitidas e recebidas e ações realizadas quando essas mensagens são enviadas ou recebidas ou quando ocorrem outros eventos desempenham um papel central em um protocolo humano. Se as pessoas seguirem protocolos diferentes (p. ex., se uma pessoa tem boas maneiras, mas a outra não; se uma delas entende o conceito de horas, mas a outra não), as pessoas não interagem e nenhum trabalho útil pode ser realizado. O mesmo é válido para redes – é preciso que duas (ou mais) entidades comunicantes sigam o mesmo protocolo para que uma tarefa seja realizada.

Vamos considerar uma segunda analogia humana. Suponha que você esteja assistindo a uma aula (p. ex., sobre redes de computadores). O professor está falando monotonicamente sobre protocolos e você está confuso. Ele para e pergunta: “Alguma dúvida?” (uma mensagem que é transmitida a todos os alunos e recebida por todos os que não estão dormindo). Você levanta a mão (transmitindo uma mensagem implícita ao professor). O professor percebe e,

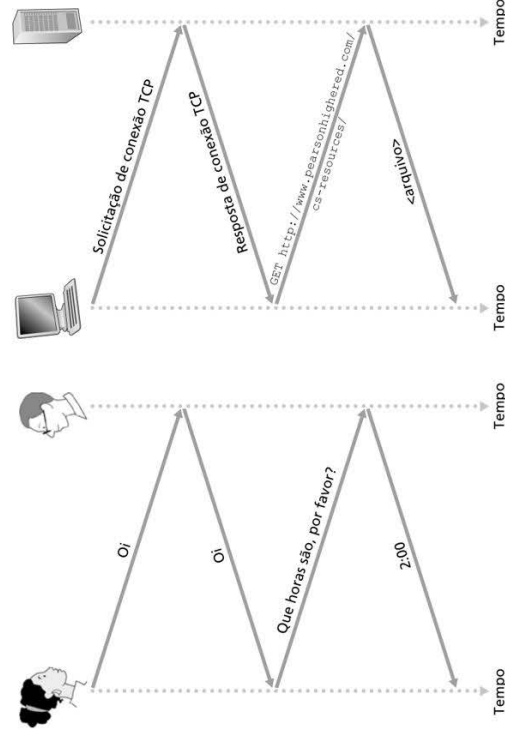


Figura 1.2 Um protocolo humano e um protocolo de rede de computadores.

com um sorriso, diz “Sim...” (uma mensagem transmitida, incentivando-o a fazer sua pergunta — professores *adoram* perguntas), e você então faz a sua (i.e., transmite sua mensagem ao professor). Ele ouve (recebe sua mensagem) e responde (transmite uma resposta a você). Mais uma vez, percebemos que a transmissão e a recepção de mensagens e um conjunto de ações convencionais, realizadas quando as mensagens são enviadas e recebidas, estão no centro desse protocolo de pergunta e resposta.

Protocolos de rede

Um protocolo de rede é semelhante a um protocolo humano: a única diferença é que as entidades que trocam mensagens e realizam ações são componentes de *hardware* ou *software* de algum dispositivo (p. ex., computador, *smartphone*, *tablet*, roteador ou outro equipamento habilitado para rede). Todas as atividades na Internet que envolvem duas ou mais entidades remotas comunicantes são governadas por um protocolo. Por exemplo, protocolos executados no *hardware* de dois computadores conectados fisicamente controlam o fluxo de *bits* no “cabo” entre as duas placas de interface de rede; protocolos de controle de congestionamento em sistemas finais controlam a taxa com que os pacotes são transmitidos entre a origem e o destino; protocolos em roteadores determinam o caminho de um pacote da origem ao destino. Eles estão em execução por toda a Internet e, em consequência, grande parte deste livro trata de protocolos de rede de computadores.

Como exemplo de um protocolo de rede de computadores com o qual você provavelmente está familiarizado, considere o que acontece quando fazemos uma requisição a um servidor Web, isto é, quando digitamos o URL de uma página Web no cliente da Web, que é chamado *browser*. Isso é mostrado no lado direito da Figura 1.2. Primeiro, o computador enviará uma mensagem de requisição de conexão ao servidor Web e aguardará uma resposta. O servidor receberá essa mensagem de requisição de conexão e retornará uma mensagem de resposta de conexão. Sabendo que agora está tudo certo para requisitar o documento da Web, o computador envia então o nome da página Web que quer buscar naquele servidor com uma mensagem GET. Por fim, o servidor retorna a página (arquivo) para o computador. Dados o exemplo humano e o exemplo de rede anteriores, as trocas de mensagens e as ações realizadas quando essas mensagens são enviadas e recebidas são os elementos fundamentais para a definição de um protocolo:

Um protocolo define o formato e a ordem das mensagens trocadas entre duas ou mais entidades comunicantes, bem como as ações realizadas na transmissão e/ou no recebimento de uma mensagem ou outro evento.

A Internet e as redes de computadores em geral fazem uso intenso de protocolos. Diferentes tipos são usados para realizar diferentes tarefas de comunicação. A medida que for avançando na leitura deste livro, você perceberá que alguns protocolos são simples e diretos, enquanto outros são complexos e intelectualmente profundos. Dominar a área de redes de computadores equivale a entender o que são, por que existem e como funcionam os protocolos de rede.

1.2 A PERIFERIA DA INTERNET

Nas seções anteriores, apresentamos uma descrição de alto nível da Internet e dos protocolos de rede. Agora, passaremos a tratar com um pouco mais de profundidade os componentes da Internet. Nesta seção, começamos pela periferia de uma rede e examinamos os componentes com os quais estamos mais familiarizados — a saber, computadores, *smartphones* e outros equipamentos que usamos diariamente. Na seção seguinte, passaremos na periferia para o núcleo da rede e estudaremos comutação e roteamento em redes de computadores.

Como descrito na seção anterior, no jargão de rede de computadores, os computadores e outros dispositivos conectados à Internet são frequentemente chamados de sistemas finais, pois se encontram na periferia da Internet, como mostrado na Figura 1.3. Os sistemas finais da Internet incluem computadores de mesa (p. ex., PCs de mesa, Macs e sistemas Linux), servidores (p. ex., servidores Web e de e-mails) e computadores móveis (p. ex., *notebooks*, *smartphones* e *tablets*). Além disso, diversas “coisas” alternativas estão sendo ligadas à Internet e transformadas em sistemas finais (veja o quadro Histórico do caso, a seguir).

Sistemas finais são denominados *hospedeiros* (ou *hosts*) porque hospedam (i.e., executam) programas de aplicação, tais como um navegador (*browser*) da Web, um programa servidor da Web, um programa leitor de e-mail ou um servidor de e-mail. Neste livro, utilizaremos os termos *hospedeiros* e *sistemas finais* como sinônimos; ou seja, *hospedeiro* = *sistema final*. Às vezes, sistemas finais são ainda subdivididos em duas categorias: **clientes** e **servidores**. Informalmente, clientes costumam ser PCs de mesa ou portáteis, *smartphones* e assim por diante, ao passo que servidores tendem a ser máquinas mais poderosas, que armazenam e distribuem páginas Web, vídeo em tempo real, retransmissão de e-mails e assim por diante. Hoje, a maioria dos servidores dos quais recebemos resultados de busca, e-mail, páginas, vídeos e conteúdo de aplicativos móveis reside em grandes **datacenters**.

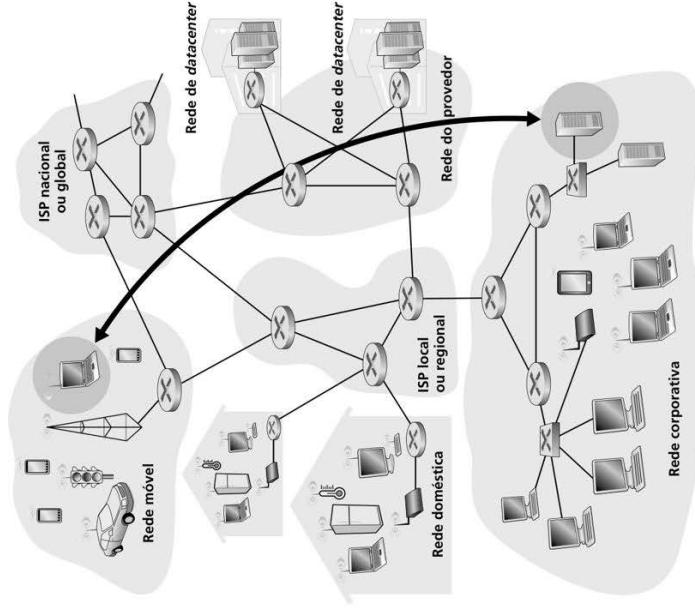


Figura 1.3 Interação entre sistemas finais.

HISTÓRICO DO CASO

DATACENTERS E COMPUTAÇÃO EM NUVEM

Empresas de Internet como Google, Microsoft, Amazon e Alibaba construíram grandes *datacenters*, cada um abrigando dezenas a centenas de milhares de servidores. Além de estarem conectados à Internet, como mostra a Figura 1.1, esses *datacenters* também incluem, internamente, redes complexas de computadores que interconectam os servidores do *datacenter*. Os *datacenters* são os motores por trás das aplicações da Internet que utilizamos diariamente.

Em linhas gerais, os *datacenters* servem três propósitos, que descrevemos aqui no contexto da Amazon para tornar o exemplo mais concreto. Primeiro, eles fornecem as páginas de comércio eletrônico da Amazon para os usuários; por exemplo, páginas que descrevem os produtos e apresentam informações de compra. Segundo, funcionam como infraestruturas de computação massivamente paralela para tarefas de processamento de dados específicas à Amazon. Terceiro, fornecem serviços de **computação em nuvem** para outras empresas. Na verdade, a maior tendência atual na computação

é que empresas usem um servidor de nuvem como a Amazon para cuidar de praticamente todas as suas necessidades de TI. Por exemplo, a Airbnb e muitas outras empresas de Internet não têm nem administram seus próprios *datacenters*, preferindo executar todos os seus serviços baseados na Web na nuvem da Amazon, chamada Amazon Web Services (AWS).

As abelhas trabalhadoras em um *datacenter* são os servidores. Eles servem o conteúdo (p. ex., páginas Web e vídeos), armazenam mensagens de correio eletrônico e documentos e realizam coletivamente cálculos massivamente distribuídos. Os servidores nos *datacenters*, chamados lâminas e semelhantes a embalagens de pizza, são em geral servidores básicos incluindo CPU, memória e armazenamento de disco. Os servidores são empilhados em estantes, com cada uma normalmente tendo de 20 a 40 lâminas. As estantes são então interconectadas usando projetos de rede de *datacenter* sofisticados e sempre em evolução. As redes de *datacenter* são discutidas em mais detalhes no Capítulo 6.

Por exemplo, em 2020, a Google tinha 19 *datacenters* em quatro continentes, que coletivamente continham vários milhões de servidores. A Figura 1.3 inclui dois desses *datacenters*, e o quadro Histórico do caso descreve os *datacenters* em mais detalhes.

1.2.1 Redes de acesso

Tendo considerado as aplicações e os sistemas finais na “periferia da Internet”, vamos agora considerar a rede de acesso – a rede física que conecta um sistema final ao primeiro roteador (também conhecido como “roteador de borda”) de um caminho partindo de um sistema final até outro qualquer. A Figura 1.4 apresenta diversos tipos de redes de acesso com linhas espessas, linhas cinzas e os ambientes (doméstico, corporativo e móvel sem fio) em que são usadas.

Acesso doméstico: DSL, cabo, FTTH e sem fio fixo 5G

Em 2020, mais de 80% dos lares europeus e americanos tinham acesso à Internet (Statista, 2019). Em razão desse uso disseminado das redes de acesso doméstico, vamos começar nossa introdução às redes de acesso considerando como os lares se conectam à Internet.

Os dois tipos de acesso residencial banda larga predominantes são a **linha digital de assinante (DSL)** e o cabo. Normalmente, uma residência obtém acesso DSL à Internet da mesma empresa que fornece acesso telefônico local com fio (p. ex., a operadora local). Assim, quando a DSL é utilizada, uma operadora do cliente é também seu ISP. Como ilustrado na Figura 1.5, o *modem* DSL de cada cliente utiliza a linha telefônica existente para trocar dados com um multiplexador digital de acesso à linha do assinante (DSLAM, do inglês *digital subscriber line access multiplexer*), em geral localizado nas dependências da operadora. O *modem* DSL da casa recebe dados digitais e os transforma em tons de alta frequência,

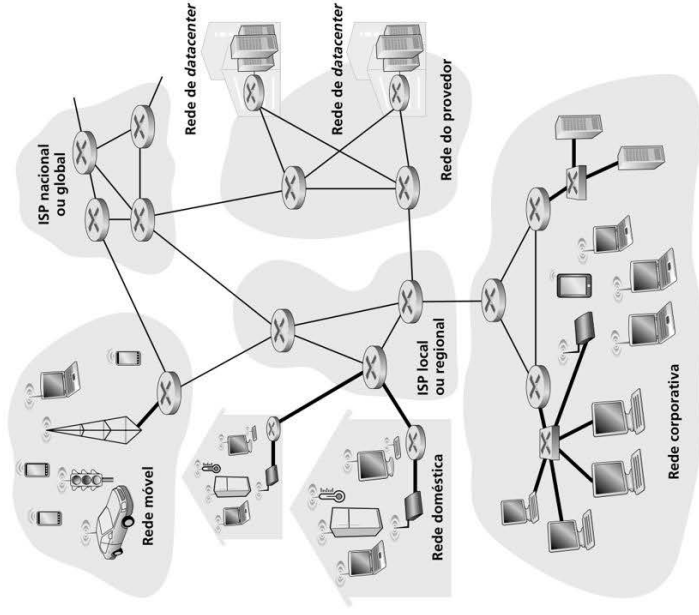


Figura 1.4 Redes de acesso.

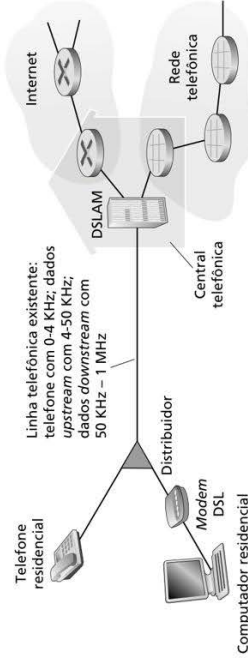


Figura 1.5 Acesso à Internet por DSL.

para transmissão pelos fios de telefone até a central telefônica (CT); os sinais analógicos de muitas dessas residências são traduzidos de volta para o formato digital no DSLAM.

A linha telefônica conduz, simultaneamente, dados e sinais telefônicos tradicionais, que são codificados em frequências diferentes:

- um canal *downstream* de alta velocidade, com uma banda de 50 kHz a 1 MHz;
- um canal *upstream* de velocidade média, com uma banda de 4 kHz a 50 kHz;
- um canal de telefone bidirecional comum, com uma banda de 0 a 4 kHz.

Essa abordagem faz a conexão DSL parecer três conexões distintas, de modo que um telefonema e a conexão com a Internet podem compartilhar a DSL ao mesmo tempo. (Descreveremos essa técnica de multiplexação por divisão de frequência na Seção 1.3.1.) Do lado do consumidor, para os sinais que chegam até sua casa, um distribuidor separa os dados e os sinais telefônicos e conduz o sinal com os dados para o *modem* DSL. Na operadora, na CT, o DSLAM separa os dados e os sinais telefônicos e envia aqueles para a Internet. Centenas ou mesmo milhares de residências se conectam a um único DSLAM.

Os padrões DSL definem múltiplas taxas de transmissão, incluindo 24 e 52 *Mbits/s* *downstream* e 3,5 e 16 *Mbits/s* *upstream*; o mais novo padrão estabelece taxas *downstream* e *upstream* somadas de 1 *Gbits/s* (ITU, 2014). Em razão de as taxas de transmissão e recebimento serem diferentes, o acesso é conhecido como assimétrico. As taxas reais alcançadas podem ser menores do que as indicadas anteriormente, pois o provedor de DSL pode, de modo proposital, limitar uma taxa residencial quando é oferecido o serviço em categorias (diferentes taxas, disponíveis a diferentes preços). A taxa máxima também pode ser limitada pela distância entre a residência e a CT, pela bitola da linha de par trançado e pelo grau de interferência elétrica. Os engenheiros projetaram a DSL expressamente para distâncias curtas entre a residência e a CT; quase sempre, se a residência não estiver localizada dentro de um limite de 8 a 16 quilômetros da CT, ela precisa recorrer a uma forma de acesso alternativa à Internet.

Embora a DSL utilize a infraestrutura de telefone local da operadora, o **acesso à Internet a cabo** utiliza a infraestrutura de TV a cabo da operadora de televisão. Uma residência obtém acesso à Internet a cabo da mesma empresa que fornece a televisão a cabo. Como ilustrado na Figura 1.6, as fibras ópticas conectam o terminal de distribuição às junções da região, sendo o cabo coaxial tradicional utilizado para chegar às casas e aos apartamentos de maneira individual. Cada junção costuma suportar de 500 a 5.000 casas. Em razão de a fibra e o cabo coaxial fazerem parte desse sistema, a rede é denominada híbrida fibra-coaxial (HFC).

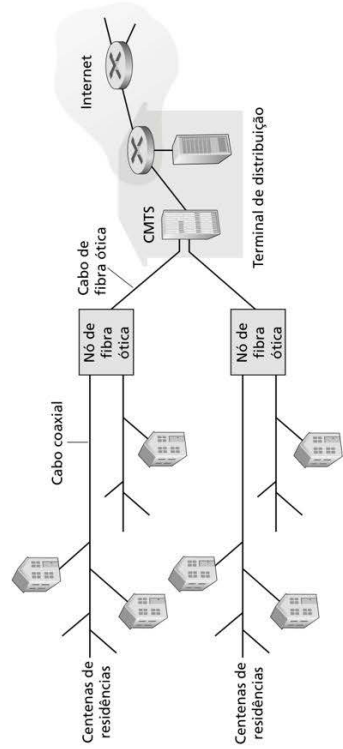


Figura 1.6 Uma rede de acesso híbrida fibra-coaxial.

O acesso à Internet a cabo necessita de *modems* especiais, denominados *modems* a cabo. Como a DSL, o *modem* a cabo é, em geral, um aparelho externo que se conecta ao computador residencial pela porta Ethernet. (Discutiremos Ethernet em detalhes no Capítulo 6.) No terminal de distribuição, o sistema de terminação do *modem* a cabo (CMTS, do inglês *cable modem termination system*) tem uma função semelhante à do DSLAM da rede DSL – transformar o sinal analógico enviado dos *modems* a cabo de muitas residências *downstream* para o formato digital. Os *modems* a cabo dividem a rede HFC em dois canais, um de transmissão (*downstream*) e um de recebimento (*upstream*). Assim como a tecnologia DSL, o acesso costuma ser assimétrico, com o canal *downstream* recebendo uma taxa de transmissão maior do que a do canal *upstream*. Os padrões DOCSIS 2.0 e 3.0 definem taxas *downstream* de 40 *Mbits/s* e 1,2 *Gbits/s*, e taxas *upstream*, de 30 *Mbits/s* e 100 *Mbits/s*, respectivamente. Como no caso das redes DSL, a taxa máxima possível de ser alcançada pode não ser observada em virtude de taxas de dados contratadas inferiores ou problemas na mídia.

Uma característica importante do acesso a cabo é o fato de ser um meio de transmissão compartilhado. Em especial, cada pacote enviado pelo terminal viaja pelos enlaces *downstream* até cada residência, e cada pacote enviado por uma residência percorre o canal *upstream* até o terminal de transmissão. Por essa razão, se diversos usuários estiverem fazendo o *download* de um arquivo em vídeo ao mesmo tempo no canal *downstream*, cada um receberá o arquivo a uma taxa bem menor do que a taxa de transmissão a cabo agregada. Por outro lado, se há somente alguns usuários ativos navegando, então cada um poderá receber páginas da Web a uma taxa de *downstream* máxima, pois esses usuários raramente solicitarão uma página ao mesmo tempo. Como o canal *upstream* também é compartilhado, é necessário um protocolo de acesso múltiplo distribuído para coordenar as transmissões e evitar colisões. (Discutiremos a questão de colisão no Capítulo 6.)

Embora as redes DSL e a cabo representem hoje a maior parte do acesso de banda larga residencial nos Estados Unidos, uma tecnologia que oferece velocidades ainda mais altas é a chamada **fiber to the home (FTTH)** (Fiber Broadband, 2020). Como o nome indica, o conceito da FTTH é simples – oferece um caminho de fibra ótica da CT diretamente até a residência. A FTTH tem o potencial de oferecer taxas de acesso à Internet na faixa de *gigabits* por segundo.

Existem várias tecnologias concorrentes para a distribuição ótica das CTs às residências. A rede mais simples é chamada fibra direta, para a qual existe uma fibra saindo da CT para cada casa. Em geral, uma fibra que sai da central telefônica é compartilhada por várias residências; ela é dividida em fibras individuais do cliente apenas após se aproximar relativamente das casas. Duas arquiteturas concorrentes de rede de distribuição ótica apresentam essa divisão: redes óticas ativas (AONs, do inglês *active optical networks*) e redes óticas passivas (PONs, do inglês *passive optical networks*). A AON é, na essência, a Ethernet comutada, assunto discutido no Capítulo 6.

Aqui, falaremos de modo breve sobre a PON, que é utilizada no serviço FIOS da Verizon. A Figura 1.7 mostra a FTTH utilizando a arquitetura de distribuição de PON. Cada residência possui um terminal de rede ótica (ONT, do inglês *optical network terminator*), que é conectado por uma fibra ótica dedicada a um distribuidor da região. O distribuidor combina certo número de residências (em geral, menos de 100) a uma única fibra ótica compartilhada, que se liga a um terminal de linha ótica (OLT, do inglês *optical line terminator*) na CT da operadora. O OLT, que fornece conversão entre sinais ópticos e elétricos, se conecta à Internet por meio de um roteador da operadora. Na residência, o usuário conecta ao ONT um roteador residencial (quase sempre sem fio) pelo qual acessa a Internet. Na arquitetura de PON, todos os pacotes enviados do OLT ao distribuidor são nele replicados (semelhante ao terminal de distribuição a cabo).

Além da DSL, cabo e FTTH, estamos começando a implantar o **acesso sem fio fixo 5G**. Além do acesso residencial de alta velocidade, o acesso sem fio fixo 5G também promete que isso ocorra sem a instalação de cabeamento caro e sujeito a falhas que se estende da CT da operadora até a residência. Com o acesso sem fio fixo 5G, usando uma tecnologia de *beam-forming* (formação de feixe), os dados são enviados da estação-base do provedor até o

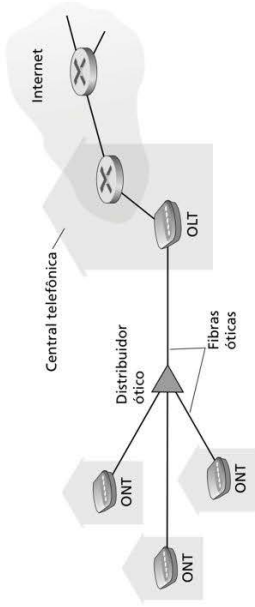


Figura 1.7 Acesso à Internet por FTTH.

modem residencial sem o uso de fios. Um roteador sem fio WiFi é conectado ao *modem* (possivelmente conjugados), semelhante ao modo como um roteador sem fio WiFi é conectado a um *modem* DSL ou a cabo. As redes celulares 5G são tratadas no Capítulo 7.

Acesso na empresa (e na residência): Ethernet e WiFi

Nos campi universitários e corporativos, e cada vez mais em residências, uma rede local (LAN, do inglês *local area network*) costuma ser usada para conectar sistemas finais ao roteador da periferia. Embora existam muitos tipos de tecnologia LAN, a Ethernet é, de longe, o padrão predominante nas redes universitárias, corporativas e domésticas. Como mostrado na Figura 1.8, os usuários utilizam um par de fios de cobre trançado para se conectarem a um *switch* Ethernet, uma tecnologia tratada com mais detalhes no Capítulo 6. O *switch* Ethernet, ou uma rede desses *switches* interconectados, é, por sua vez, conectado à Internet como um todo. Com o acesso por uma rede Ethernet, os usuários normalmente têm acesso entre 100 *Mbits/s* e dezenas de *Gbits/s* com o *switch* Ethernet, enquanto os servidores possuem um acesso de 1 *Gbit/s* ou 10 *Gbits/s*.

Está cada vez mais comum as pessoas acessarem a Internet sem fio, seja por *notebooks*, *smartphones*, *tablets* ou por outros dispositivos. Em uma LAN sem fio, os usuários transmitem/recebem pacotes para/de um ponto de acesso que está conectado à rede da empresa (quase sempre incluindo Ethernet com fio) que, por sua vez, é conectada à Internet com fio. Um usuário de LAN sem fio deve estar no espaço de alguns metros do ponto de acesso. O acesso à LAN sem fio baseado na tecnologia IEEE 802.11, ou seja, WiFi, hoje está

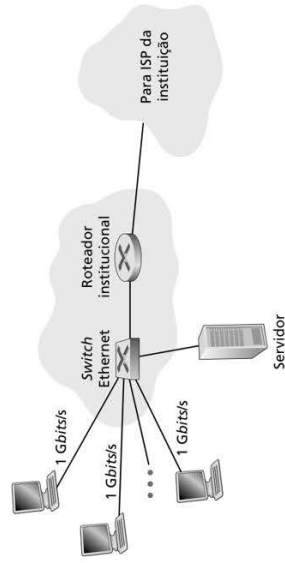


Figura 1.8 Acesso à Internet por Ethernet.

presente em todos os lugares – universidades, empresas, cafés, aeroportos, residências e, até mesmo, em aviões. Como discutido com detalhes no Capítulo 7, hoje o 802.11 fornece uma taxa de transmissão compartilhada de mais de 100 *Mbits/s*.

Embora as redes de acesso por Ethernet e WiFi fossem implantadas no início em ambientes corporativos (empresas, universidades), elas há pouco se tornaram componentes básicos comuns das redes residenciais. Muitas casas unem o acesso residencial banda larga (ou seja, *modem* a cabo ou DSL) com a tecnologia LAN sem fio a um custo acessível para criar redes residenciais potentes. A Figura 1.9 mostra um esquema de uma rede doméstica típica. Ela consiste em um *notebook* móvel, múltiplos eletrodomésticos conectados à Internet e um computador com fio; uma estação-base (o ponto de acesso sem fio), que se comunica com o computador sem fio e com outros dispositivos sem fio pela casa; e um roteador, que interconecta o ponto de acesso sem fio, assim como outros dispositivos domésticos com fio, à Internet. Essa rede permite que os moradores tenham acesso banda larga à Internet com um usuário se movimentando da cozinha ao quintal e até os quartos.

Acesso sem fio em longa distância: 3G e LTE 4G e 5G

Dispositivos móveis como iPhones e dispositivos Android estão sendo usados para enviar mensagens, compartilhar fotos em redes sociais, realizar pagamentos móveis, assistir filmes, fazer *streaming* de música e muito mais, sempre em movimento. Esses dispositivos empregam a mesma infraestrutura sem fio usada pela telefonia celular para enviar/receber pacotes por uma estação-base que é controlada pela operadora da rede celular. Diferente do WiFi, um usuário só precisa estar dentro de um raio de algumas dezenas de quilômetros (e não de algumas dezenas de metros) da estação-base.

As empresas de telecomunicação têm investido enormemente na chamada quarta geração (4G) sem fio, que oferece velocidades de *download* de até 60 *Mbits/s* no mundo real. Porém, tecnologias de acesso remotas de maior velocidade – uma quinta geração (5G) de redes sem fio de longa distância – já estão sendo implantadas. Veremos os princípios básicos das redes sem fio e mobilidade, além de tecnologias WiFi, 4G e 5G (e mais!) no Capítulo 7.

1.2.2 Meios físicos

Na subseção anterior, apresentamos uma visão geral de algumas das mais importantes tecnologias de acesso à Internet. Ao descrevê-las, indicamos também os meios físicos utilizados por elas. Por exemplo, dissemos que a HFC usa uma combinação de cabo de fibra ótica e cabo coaxial. Dissemos que DSL e Ethernet utilizam fios de cobre. Dissemos também que redes de acesso móveis usam o espectro de rádio. Nesta subseção, damos uma visão geral desses e de outros meios de transmissão empregados na Internet.

Para definir o que significa meio físico, vamos pensar na curta vida de um *bit*. Considere um *bit* saindo de um sistema final, transitando por uma série de enlaces e roteadores e chegando a outro sistema final. Esse pobre e pequeno *bit* é transmitido muitas e muitas

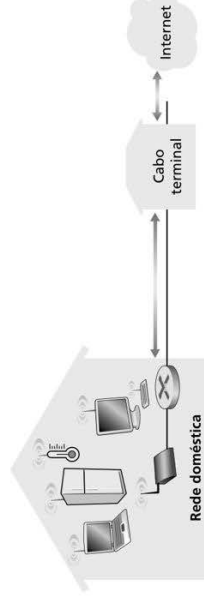


Figura 1.9 Esquema de uma rede doméstica típica.

vezes. Primeiro, o sistema final originador transmite o *bit* e, logo em seguida, o primeiro roteador da série recebe-o; então, o primeiro roteador envia-o para o segundo roteador e assim por diante. Assim, nosso *bit*, ao viajar da origem ao destino, passa por uma série de pares transmissor-receptor, que o recebem por meio de ondas eletromagnéticas ou pulsos óticos que se propagam por um **meio físico**. Com muitos aspectos e formas possíveis, o meio físico não precisa ser obrigatoriamente do mesmo tipo para cada par transmissor-receptor ao longo do caminho. Alguns exemplos de meios físicos são: par de fios de cobre trançado, cabo coaxial, cabo de fibra ótica multimodo, espectro de rádio terrestre e espectro de rádio por satélite. Os meios físicos se enquadram em duas categorias: **meios guiados** e **meios não guiados**. Nos meios guiados, as ondas são dirigidas ao longo de um meio sólido, tal como um cabo de fibra ótica, um par de fios de cobre trançado ou um cabo coaxial. Nos meios não guiados, as ondas se propagam na atmosfera e no espaço, como é o caso de uma LAN sem fio ou de um canal digital de satélite.

Contudo, antes de examinar as características dos vários tipos de meios, vamos discutir um pouco seus custos. O custo real de um enlace físico (fio de cobre, cabo de fibra ótica e assim por diante) costuma ser insignificante em comparação a outros. Em especial, o custo da mão de obra de instalação do enlace físico pode ser várias vezes maior do que o do material. Por essa razão, muitos construtores instalam pares de fios trançados, fibra ótica e cabo coaxial em todas as salas de um edifício. Mesmo que apenas um dos meios seja usado inicialmente, há uma boa probabilidade de outro ser usado no futuro próximo – portanto, poupa-se dinheiro por não ser preciso instalar fiação adicional depois.

Par de fios de cobre trançado

O meio de transmissão guiado mais barato e mais usado é o par de fios de cobre trançado, que vem sendo empregado há mais de 100 anos nas redes de telefonia. De fato, mais de 99% da fiação que conecta aparelhos telefônicos a centrais locais utilizam esse meio. Quase todos nós já vimos um em casa (ou na casa dos nossos pais e avós!) ou no local de trabalho: esse par constituído de dois fios de cobre isolados, cada um com cerca de 1 milímetro de espessura, enrolados em espiral. Os fios são trançados para reduzir a interferência elétrica de pares semelhantes que estejam próximos. Normalmente, uma série de pares é conjugada dentro de um cabo, isolando-se os pares com blindagem de proteção. Um par de fios constitui um único enlace de comunicação. O **par trançado sem blindagem (UTP, do inglês *unshielded twisted pair*)** costuma ser usado em redes de computadores de edifícios, isto é, em LANs. Hoje, as taxas de transmissão de dados para as LANs de pares trançados estão na faixa de 10 *Mbits/s* a 10 *Gbits/s*. As taxas de transmissão de dados que podem ser alcançadas dependem da bitola do fio e da distância entre transmissor e receptor.

Quando a tecnologia da fibra ótica surgiu na década de 1980, muitos depreciaram o par trançado por suas taxas de transmissão de *bits* relativamente baixas. Alguns até acharam que a tecnologia da fibra ótica o substituiria por completo. Mas ele não desistiu assim tão facilmente. A moderna tecnologia de par trançado, tal como o cabo de categoria 6a, pode alcançar taxas de transmissão de dados de 10 *Gbits/s* para distâncias de até algumas centenas de metros. No final, o par trançado firmou-se como a solução dominante para LANs de alta velocidade.

Como vimos, o par trançado também é usado para acesso residencial à Internet. Vimos que a tecnologia do *modem* discado possibilita taxas de acesso de até 56 *kbits/s* com pares trançados. Vimos também que a tecnologia DSL (linha digital de assinante) permitiu que usuários residenciais acessem a Internet em dezenas de *Mbits/s* com pares de fios trançados (quando as residências estão próximas à CT do ISP).

Cabo coaxial

Como o par trançado, o cabo coaxial é constituído por dois condutores de cobre, porém concêntricos e não paralelos. Com essa configuração, isolamento e blindagem especiais,

pode alcançar taxas altas de transmissão de dados. Cabos coaxiais são muito comuns em sistemas de televisão a cabo. Como já comentamos, recentemente sistemas de televisão a cabo foram acoplados com *modems* a cabo para oferecer aos usuários residenciais acesso à Internet a velocidades de centenas de *Mbits/s*. Em televisão a cabo e acesso a cabo à Internet, o transmissor passa o sinal digital para uma banda de frequência específica, e o sinal analógico resultante é enviado do transmissor para um ou mais receptores. O cabo coaxial pode ser utilizado como um **meio compartilhado** guiado. Vários sistemas finais podem ser conectados diretamente ao cabo, e todos eles recebem qualquer sinal que seja enviado pelos outros sistemas finais.

Fibras óticas

A fibra ótica é um meio delgado e flexível que conduz pulsos de luz, cada um deles representando um *bit*. Uma única fibra ótica pode suportar taxas de transmissão elevadíssimas, de até dezenas ou mesmo centenas de *gigabits* por segundo. Fibras óticas são imunes à interferência eletromagnética, têm baixíssima atenuação de sinal até 100 quilômetros, e são muito difíceis de criar derivações de sinal. Essas características fizeram da fibra ótica o meio preferido para a transmissão guiada de grande alcance, em especial para cabos submarinos. Hoje, muitas redes telefônicas de longa distância dos Estados Unidos e de outros países usam exclusivamente fibras óticas, que também predominam no *backbone* da Internet. Contudo, o alto custo de equipamentos óticos – como transmissores, receptores e nós de comutação – vem atrasando sua utilização para transporte a curta distância, como em LANs ou em redes de acesso residenciais. As velocidades de conexão do padrão Optical Carrier (OC) variam de 51,8 *Mbits/s* a 39,8 *Gbits/s*; essas especificações são frequentemente denominadas OC-*n*, em que a velocidade de conexão se iguala a $n \times 51,8 \text{ Mbits/s}$. Os padrões usados hoje incluem OC-1, OC-3, OC-12, OC-24, OC-48, OC-96, OC-192 e OC-768.

Canais de rádio terrestres

Canais de rádio carregam sinais dentro do espectro eletromagnético. São um meio atraente, porque sua instalação não requer cabos físicos, podem atravessar paredes, dão conectividade ao usuário móvel e, potencialmente, conseguem transmitir um sinal a longas distâncias. As características de um canal de rádio dependem muito do ambiente de propagação e da distância pela qual o sinal deve ser transmitido. Condições ambientais determinam perda de sinal no caminho e atenuação por efeito de sombra (que reduz a intensidade do sinal quando ele transita por distâncias longas e ao redor/através de objetos interferentes), atenuação por caminhos múltiplos (devido à reflexão do sinal quando atinge objetos interferentes) e interferência (por outras transmissões ou sinais eletromagnéticos).

Canais de rádio terrestres podem ser classificados, de modo geral, em três grupos: os que operam sobre distâncias muito curtas (p. ex., com 1 ou 2 metros); os de pequeno alcance, que funcionam em locais próximos, normalmente abrangendo de 10 a algumas centenas de metros; e os de longo alcance, que abrangem dezenas de quilômetros. Dispositivos pessoais, como fones sem fio, teclados e dispositivos médicos, operam por curtas distâncias; as tecnologias LAN sem fio, descritas na Seção 1.2.1, utilizam canais de rádio local de média distância; as tecnologias de acesso em telefone celular utilizam canal de rádio de longo alcance. Abordaremos canais de rádio detalhadamente no Capítulo 7.

Canais de rádio por satélite

Um satélite de comunicação liga dois ou mais transmissores-receptores de micro-ondas baseados na Terra, denominados estações terrestres. Ele recebe transmissões em uma faixa de frequência, gera novamente o sinal usando um repetidor (sobre o qual falaremos a seguir) e o transmite em outra frequência. Dois tipos de satélites são usados para comunicações: **satélites geostacionários** e **satélites de órbita baixa (LEO, do inglês *low-earth orbiting*)**.

Os satélites geostacionários ficam de modo permanente sobre o mesmo lugar da Terra. Essa presença estacionária é conseguida colocando-se o satélite em órbita a 36 mil quilômetros acima da superfície terrestre. Essa enorme distância da estação terrestre ao satélite e de seu caminho de volta à estação terrestre traz um substancial atraso de propagação de sinal de 280 milissegundos. Mesmo assim, enlaces por satélite, que podem funcionar a velocidades de centenas de $Mbits/s$, são frequentemente usados em áreas sem acesso à Internet baseada em DSL ou a cabo.

Os satélites LEO são posicionados muito mais próximos da Terra e não ficam sempre sobre um único lugar. Eles giram ao redor da Terra (exatamente como a Lua) e podem se comunicar uns com os outros e com estações terrestres. Para prover cobertura contínua em determinada área, é preciso colocar muitos satélites em órbita. Hoje, existem muitos sistemas de comunicação de baixa altitude em desenvolvimento. A tecnologia de satélites LEO poderá ser utilizada para acesso à Internet no futuro.

1.3 O NÚCLEO DA REDE

Após termos examinado a periferia da Internet, vamos agora nos aprofundar mais no núcleo da rede — a rede de nós de comutação de pacotes e enlaces que interconectam os sistemas finais da Internet. Os núcleos da rede aparecem destacados em cinza na Figura 1.10.

1.3.1 Comutação de pacotes

Em uma aplicação de rede, sistemas finais trocam **mensagens** entre si. Mensagens podem conter qualquer coisa que o projetista do protocolo queira. Podem desempenhar uma função de controle (p. ex., as mensagens “Oi” no nosso exemplo de comunicação na Figura 1.2) ou conter dados, tal como um *e-mail*, uma imagem JPEG ou um arquivo de áudio MP3. Para enviar uma mensagem de um sistema final de origem para um destino, o originador fragmenta mensagens longas em porções de dados menores, denominadas **pacotes**. Entre origem e destino, cada um deles percorre enlaces de comunicação e **nós de comutação de pacotes** (há dois tipos principais de comutadores de pacotes: **roteadores** e **switches**). Pacotes são transmitidos por cada enlace de comunicação a uma taxa igual à de transmissão *total*. Assim, se um sistema final de origem ou um nó de comutação de pacotes estiver enviando um pacote de L bits por um enlace com taxa de transmissão de R bits/s, então o tempo para transmitir o pacote é L/R segundos.

Transmissão armazena-e-reenvia

A maioria dos nós de comutação de pacotes utiliza a **transmissão armazena-e-reenvia** (*store-and-forward*) nas entradas dos enlaces. A transmissão armazena-e-reenvia significa que o nó de comutação de pacotes deve receber o pacote inteiro antes de poder começar a transmitir o primeiro *bit* para o enlace de saída. Para explorar a transmissão armazena-e-reenvia com mais detalhes, considere uma rede simples, consistindo em dois sistemas finais conectados por um único roteador, conforme mostra a Figura 1.11. Um roteador em geral terá muitos enlaces incidentes; pois sua função é transferir um pacote que chega para um enlace de saída; neste exemplo simples, o roteador tem a tarefa de transferir um pacote de um enlace (entrada) para o único outro enlace conectado. Aqui, a origem tem três pacotes, cada um consistindo em L bits, para enviar ao destino. No instante de tempo mostrado na Figura 1.11, a origem transmitiu parte do pacote 1, e a frente do pacote 1 já chegou no roteador. Como entrega a transmissão armazena-e-reenvia, nesse momento, o roteador não pode transmitir os *bits* que recebeu; em vez disso, ele precisa primeiro manter em *buffer* (i.e., “armazenar”) os *bits* do pacote. Somente depois que o roteador tiver recebido *todos* os

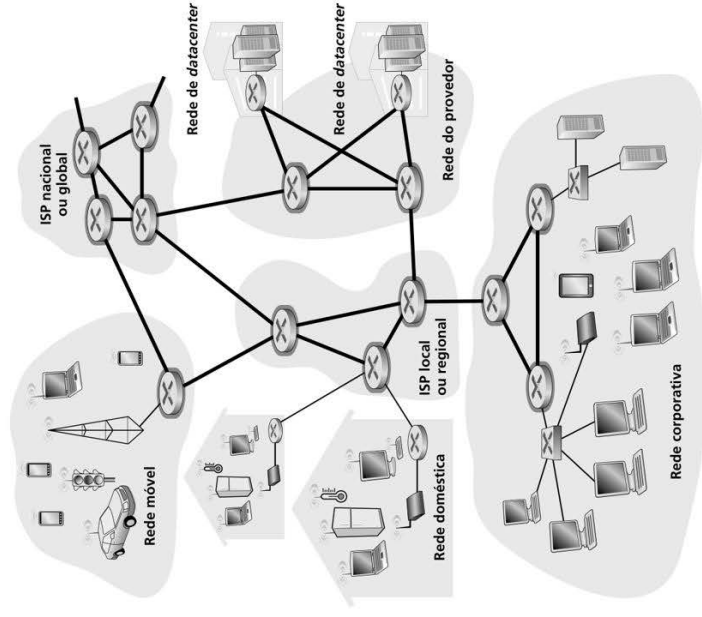


Figura 1.10 O núcleo da rede.

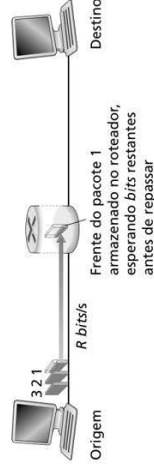


Figura 1.11 Comutação de pacotes armazena-e-reenvia.

de um pacote, poderá começar a transmitir (i.e., “reenviar”) o pacote para o enlace de saída. Para ter uma ideia da transmissão armazena-e-reenvia, vamos agora calcular a quantidade de tempo decorrido desde quando a origem começa a enviar até que o destino tenha recebido o pacote inteiro. (Aqui, ignoraremos o atraso de propagação — o tempo gasto para os *bits* atravessarem o fio em uma velocidade próxima à da luz —, o que será discutido na Seção 1.4.) A origem começa a transmitir no tempo 0; no tempo L/R segundos, a origem terá transmitido o pacote inteiro, que terá sido recebido e armazenado no roteador (pois não há atraso de

propagação). No tempo L/R segundos, como o roteador já terá recebido o pacote inteiro, ele pode começar a transmiti-lo para o enlace de saída, em direção ao destino; no tempo $2L/R$, o roteador terá transmitido o pacote inteiro, e este terá sido recebido pelo destino. Assim, o atraso total é $2L/R$. Se o roteador, em vez disso, reenviasse os *bits* assim que chegassem (sem primeiro receber o pacote inteiro), então o atraso total seria L/R , pois os *bits* não são mantidos no roteador. Mas, conforme discutiremos na Seção 1.4, os roteadores precisam receber, armazenar e *processar* o pacote inteiro antes de encaminhá-lo.

Agora vamos calcular a quantidade de tempo decorrido desde quando a origem começa a enviar o primeiro pacote até que o destino tenha recebido todos os três. Como antes, no instante L/R , o roteador começa a reenviar o primeiro pacote. Mas, também no tempo L/R , a origem começará a enviar o segundo, pois ela terá acabado de mandar o primeiro pacote inteiro. Assim, no tempo $2L/R$, o destino terá recebido o primeiro pacote e o roteador terá recebido o segundo. De modo semelhante, no instante $3L/R$, o destino terá recebido os dois primeiros pacotes e o roteador terá recebido o terceiro. Por fim, no tempo $4L/R$, o destino terá recebido todos os três pacotes!

Vamos considerar o caso geral do envio de um pacote da origem ao destino por um caminho que consiste em N enlaces, cada um com taxa R (assim, há $N - 1$ roteadores entre origem e destino). Aplicando a mesma lógica usada anteriormente, vemos que o atraso fim a fim é:

$$d_{\text{fim a fim}} = N \frac{L}{R} \quad (1.1)$$

Você poderá tentar determinar qual seria o atraso para P pacotes enviados por uma série de N enlaces.

Atrasos de fila e perda de pacote

A cada nó de comutação de pacotes estão ligados vários enlaces. Para cada um destes, o nó de comutação de pacotes tem um **buffer de saída** (também denominado **fila de saída**), que armazena pacotes prestes a serem enviados pelo roteador para aquele enlace. Os *buffers* de saída desempenham um papel fundamental na comutação de pacotes. Se um pacote que está chegando precisa ser transmitido por um enlace, mas o encontra ocupado com a transmissão de outro pacote, deve aguardar no *buffer* de saída. Desse modo, além dos atrasos de armazenagem e reenvio, os pacotes sofrem **atrasos de fila** no *buffer* de saída. Esses atrasos são variáveis e dependem do grau de congestionamento da rede. Como o espaço do *buffer* é finito, um pacote que está chegando pode encontrar o *buffer* lotado de outros que estão esperando transmissão. Nesse caso, ocorrerá uma **perda de pacote** – um pacote que está chegando ou um dos que já estão na fila é descartado.

A Figura 1.12 ilustra uma rede simples de comutação de pacotes. Como na Figura 1.11, os pacotes são representados por placas tridimensionais. A largura de uma placa representa o número de *bits* no pacote. Nessa figura, todos os pacotes têm a mesma largura, portanto, o mesmo tamanho. Suponha que os hóspedes A e B estejam enviando pacotes ao hospedeiro E. Os hóspedes A e B primeiro enviarão seus pacotes por enlaces Ethernet de 100 *Mbits/s* até o primeiro nó de comutação, que vai direcioná-los para o enlace de 15 *Mbits/s*. Se, durante um pequeno intervalo de tempo, a taxa de chegada de pacotes ao roteador (quando convertida para *bits* por segundo) for maior do que 15 *Mbits/s*, ocorrerá congestionamento no roteador, pois os pacotes formarão uma fila no *buffer* de saída do enlace antes de serem transmitidos para o enlace. Por exemplo, se cada um dos hóspedes A e B enviar uma rajada de cinco pacotes de ponta a ponta ao mesmo tempo, então a maior parte deles gastará algum tempo esperando na fila. De fato, a situação é semelhante a muitas no dia a dia – por exemplo, quando aguardamos na fila de um caixa de banco ou quando esperamos em uma cabine de pedágio. Vamos analisar esse atraso de fila mais detalhadamente na Seção 1.4.

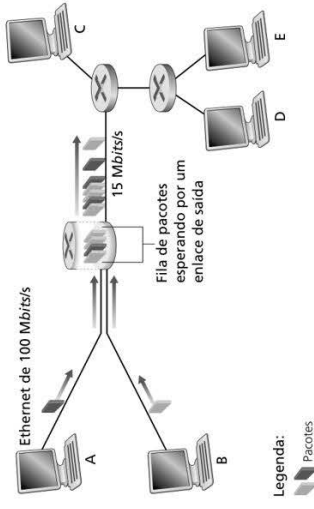


Figura 1.12 Comutação de pacotes.

Tabelas de repasse e protocolos de roteamento

Dissemos anteriormente que um roteador conduz um pacote que chega a um de seus enlaces de comunicação para outro de seus enlaces de comunicação conectados. Mas como o roteador determina o enlace ao qual deve conduzir o pacote? Na verdade, isso é feito de diversas maneiras por diferentes tipos de rede de computadores. Aqui, descreveremos de modo resumido como isso é feito pela Internet.

Na Internet, cada sistema final tem um endereço denominado endereço IP. Quando um sistema final de origem quer enviar um pacote a um destino, a origem inclui o endereço IP do destino no cabeçalho do pacote. Como os endereços postais, o IP possui uma estrutura hierárquica. Quando um pacote chega a um roteador na rede, este examina uma parte do endereço de destino e o conduz a um roteador adjacente. Mais especificamente, cada roteador possui uma **tabela de repasse** que mapeia os endereços de destino (ou partes deles) para enlaces de saída desse roteador. Quando um pacote chega a um roteador, este examina o endereço e pesquisa sua tabela de repasse, utilizando esse endereço de destino para encontrar o enlace de saída apropriado. O roteador, então, direciona o pacote a esse enlace de saída.

O processo de roteamento fim a fim é semelhante a um motorista que não quer consultar o mapa, preferindo pedir informações. Por exemplo, suponha que Joe vai dirigir da Flórida para 156 Lakeside Drive, em Orlando, Flórida. Primeiro, Joe vai ao posto de gasolina de seu bairro e pergunta como chegar a 156 Lakeside Drive, em Orlando, Flórida. O frentista do posto extrai a palavra Flórida do endereço e diz que Joe precisa pegar a interestadual I-95 Sul, cuja entrada fica ao lado do posto. Ele também diz a Joe para pedir outras informações assim que chegar à Flórida. Então, Joe pega a I-95 Sul até chegar a Jacksonville, na Flórida, onde pede mais informações a outro frentista. Este extrai a palavra Orlando do endereço a Joe para continuar na I-95 até Daytona Beach, e lá se informar de novo. Em Daytona Beach, outro frentista também extrai a palavra Orlando do endereço e pede para ele pegue a I-4 diretamente para Orlando. Joe segue suas orientações e chega a uma saída para Orlando. Ele vai até outro posto de gasolina, e dessa vez o frentista extrai a palavra Lakeside Drive do endereço e diz a ele qual estrada seguir para Lakeside Drive. Assim que Joe chega a Lakeside Drive, pergunta a uma criança andando de bicicleta como chegar a seu destino. A criança extrai o número 156 do endereço e aponta para a casa. Joe finalmente chega a seu destino. Nessa analogia, os frentistas de posto de gasolina e as crianças andando de bicicleta são semelhantes aos roteadores.

Vimos que um roteador usa um endereço de destino do pacote para indexar uma tabela de repasse e determinar o enlace de saída apropriado. Mas essa afirmação traz ainda outra questão: como as tabelas de repasse são montadas? Elas são configuradas manualmente

em cada roteador ou a Internet utiliza um procedimento mais automatizado? Essa questão será estudada com mais profundidade no Capítulo 5. Mas, para aguçar seu apetite, observe que a Internet possui uma série de **protocolos de roteamento** especiais, que são utilizados para configurar automaticamente as tabelas de repasse. Um protocolo de roteamento pode, por exemplo, determinar o caminho mais curto de cada roteador a cada destino e utilizar os resultados para configurar as tabelas de repasse nos roteadores.

1.3.2 Comutação de circuitos

Há duas abordagens fundamentais para transmissão de dados através de uma rede de enlaces e nós de comutação: **comutação de circuitos** e **comutação de pacotes**. Tendo visto estas últimas na subseção anterior, agora vamos voltar nossa atenção às redes de comutação de circuitos.

Nessas redes, os recursos necessários ao longo de um caminho (*buffer*s, taxa de transmissão de enlaces) para oferecer comunicação entre os sistemas finais são *reservados* pelo período da sessão de comunicação entre os sistemas finais. Em redes de comutação de pacotes, tais recursos *não* são reservados; as mensagens de uma sessão usam os recursos por demanda e, como consequência, poderão ter de esperar (i.e., entrar na fila) para conseguir acesso a um enlace de comunicação. Como simples analogia, considere dois restaurantes – um que exige e outro que não exige nem aceita reserva. Se quisermos ir ao restaurante que exige reserva, teremos de passar pelo aborrecimento de telefonar antes de sair de casa. Mas, quando chegarmos lá, poderemos, em princípio, ser logo atendidos e servidos. No segundo restaurante, não precisaremos nos dar ao trabalho de reservar mesa, porém, quando lá chegarmos, talvez tenhamos de esperar para sentar.

As redes de telefonia tradicionais são exemplos de redes de comutação de circuitos. Considere o que acontece quando uma pessoa quer enviar a outra uma informação (por voz ou por fax) por meio de uma rede telefônica. Antes que o remetente possa enviar a informação, a rede precisa estabelecer uma conexão entre ele e o destinatário. Essa é uma conexão forte, na qual os computadores no caminho entre o remetente e o destinatário mantêm o estado. No jargão da telefonia, essa conexão é denominada **circuito**. Quando a rede estabelece o circuito, também reserva uma taxa de transmissão constante nos enlaces da rede durante o período da conexão. Visto que foi reservada largura de banda para essa conexão remetente-destinatário, o remetente pode transferir dados ao destinatário a uma taxa constante *garantida*.

A Figura 1.13 ilustra uma rede de comutação de circuitos. Nela, os quatro nós de comutação de circuitos estão interconectados por quatro enlaces. Cada enlace tem quatro circuitos, de modo que cada um pode suportar quatro conexões simultâneas. Cada um dos hospedeiros (p. ex., PCs e estações de trabalho) está conectado diretamente a um dos circuitos.

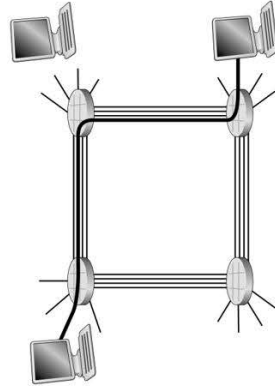


Figura 1.13 Uma rede simples de comutação de circuitos composta por quatro nós de comutação e quatro enlaces.

Quando dois sistemas finais querem se comunicar, a rede estabelece uma **conexão fim a fim** dedicada entre os dois hospedeiros. Assim, para que o sistema final A envie mensagens ao sistema final B, a rede deve primeiro reservar um circuito em cada um dos dois enlaces. Neste exemplo, a conexão fim a fim dedicada usa o segundo circuito no primeiro enlace e o quarto circuito no segundo enlace. Como cada enlace tem quatro circuitos, para cada enlace usado pela conexão fim a fim, esta fica com um quarto da capacidade de transmissão total durante o período da conexão. Assim, por exemplo, se cada enlace entre computadores adjacentes tiver uma taxa de transmissão de 1 Mbit/s, então cada conexão comutada por circuitos fim a fim obtém 250 kbit/s de taxa de transmissão dedicada.

Em contrapartida, considere o que ocorre quando um sistema final quer enviar um pacote a outro hospedeiro por uma rede de comutação de pacotes, como a Internet. Como acontece na comutação de circuitos, o pacote é transmitido por uma série de enlaces de comunicação. Mas, na comutação de pacotes, ele é enviado à rede sem reservar qualquer recurso do enlace. Se um dos enlaces estiver congestionado porque outros pacotes precisam ser transmitidos ao mesmo tempo, então nosso pacote terá de esperar em um *buffer* na extremidade de origem do enlace de transmissão e sofrerá um atraso. A Internet faz o melhor esforço para entregar os dados rapidamente, mas não dá garantia alguma.

Multiplexação em redes de comutação de circuitos

Um circuito é implementado em um enlace por **multiplexação por divisão de frequência** (FDM, do inglês *frequency-division multiplexing*) ou por **multiplexação por divisão de tempo** (TDM, do inglês *time-division multiplexing*). Com FDM, o espectro de frequência de um enlace é compartilhado entre as conexões estabelecidas através desse enlace. Ou seja, o enlace reserva uma banda de frequência para cada conexão durante o período da ligação. Em redes telefônicas, a largura dessa banda de frequência em geral é 4 kHz (i.e., 4 mil Hertz ou 4 mil ciclos por segundo). A largura da banda é denominada, claro, **largura de banda**. Estações de rádio FM também usam FDM para compartilhar o espectro de frequência entre 88 MHz e 108 MHz, sendo atribuída para cada estação uma banda de frequência específica.

Em um enlace TDM, o tempo é dividido em quadros de duração fixa, e cada quadro é dividido em um número fixo de compartimentos (*slots*). Quando estabelece uma conexão por meio de um enlace, a rede dedica à conexão um compartimento de tempo em cada quadro. Esses compartimentos são reservados para o uso exclusivo dessa conexão, e um dos compartimentos de tempo (em cada quadro) fica disponível para transmitir os dados dela.

A Figura 1.14 ilustra as técnicas FDM e TDM para um enlace de rede que suporta até quatro circuitos. Para FDM, o domínio de frequência é segmentado em quatro faixas, com largura de banda de 4 kHz cada. Para TDM, o domínio de tempo é segmentado em quadros, cada um com quatro compartimentos de tempo; a cada circuito é designado o mesmo compartimento dedicado nos quadros sucessivos. Para TDM, a taxa de transmissão de um circuito é igual à taxa do quadro multiplicada pelo número de *bits* em um compartimento. Por exemplo, se o enlace transmite 8 mil quadros por segundo e cada compartimento consiste em 8 *bits*, então a taxa de transmissão de um circuito é 64 kbit/s.

Os defensores da comutação de pacotes sempre argumentaram que comutação de circuitos é desperdício, porque os circuitos dedicados ficam ociosos durante **períodos de silêncio**. Por exemplo, quando um dos participantes de uma conversa telefônica para de falar, os recursos ociosos da rede (bandas de frequências ou compartimentos nos enlaces ao longo da rota da conexão) não podem ser usados por outras conexões em curso. Para outro exemplo de como esses recursos podem ser subutilizados, considere um radiologista que usa uma rede de comutação de circuitos para acessar remotamente uma série de radiografias. Ele estabelece uma conexão, requisita uma imagem, examina-a e, em seguida, solicita uma nova. Recursos de rede são atribuídos à conexão, mas não são utilizados (i.e., são desperdiçados) no período em que o radiologista examina a imagem. Defensores da comutação de pacotes também gostam de destacar que estabelecer circuitos e reservar larguras de banda fim a fim

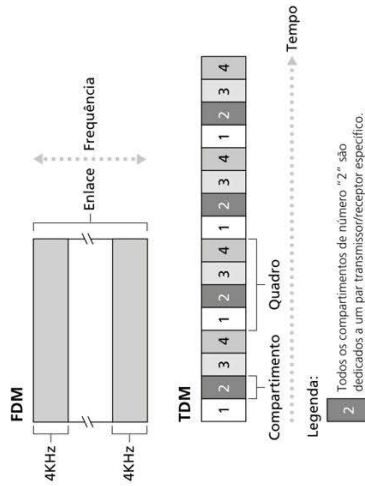


Figura 1.14 Com FDM, cada circuito dispõe continuamente de uma fração da largura de banda. Com TDM, cada circuito dispõe de toda a largura de banda periodicamente, durante breves intervalos de tempo (i.e., durante compartimentos de tempo).

são processos complicados e exigem *softwares* complexos de sinalização para coordenar a operação dos nós de comutação ao longo do caminho.

Antes de encerrarmos esta discussão sobre comutação de circuitos, examinemos um exemplo numérico que deverá esclarecer melhor o assunto. Vamos considerar o tempo que levamos para enviar um arquivo de 640 mil *bits* do hospedeiro A ao hospedeiro B por uma rede de comutação de circuitos. Suponha que todos os enlaces da rede usem TDM de 24 compartimentos e tenham uma taxa de 1,536 *Mbits/s*. Suponha também que um circuito fim a fim leva 500 milissegundos para ser ativado antes que A possa começar a transmitir o arquivo. Em quanto tempo o arquivo será enviado? Cada circuito tem uma taxa de transmissão de $(1,536 \text{ Mbits/s})/24 = 64 \text{ kbits/s}$; portanto, demorará $(640,000 \text{ bits})/(64 \text{ kbits/s}) = 10$ segundos para transmitir o arquivo. A esses 10 segundos adicionamos o tempo de ativação do circuito, resultando 10,5 segundos para o envio. Observe que o tempo de transmissão é independente do número de enlaces: o tempo de transmissão seria 10 segundos se o circuito fim a fim passasse por um ou por uma centena de enlaces. (O atraso real fim a fim também inclui um atraso de propagação; ver Seção 1.4.)

Comutação de pacotes versus comutação de circuitos

Agora que já descrevemos comutação de pacotes e comutação de circuitos, vamos comparar as duas. Opositores da comutação de pacotes costumam argumentar que ela não é adequada para serviços de tempo real (p. ex., ligações telefônicas e videoconferência) em virtude de seus atrasos fim a fim variáveis e imprevisíveis (que se devem principalmente a variáveis e imprevisíveis atrasos de fila). Defensores da comutação de pacotes argumentam que (1) ela oferece melhor compartilhamento de banda do que a comutação de circuitos e (2) sua implementação é mais simples, mais eficiente e mais barata do que a de comutação de circuitos. Uma discussão interessante sobre comutação de pacotes e comutação de circuitos pode ser encontrada em Molinero-Fernandez (2002). De modo geral, quem não gosta de perder tempo fazendo reserva de mesa em restaurantes prefere comutação de pacotes à comutação de circuitos.

Por que a comutação de pacotes é mais eficiente? Vamos examinar um exemplo simples. Suponha que usuários compartilhem um enlace de 1 *Mbit/s*. Considere também que cada usuário alterne períodos de atividade, quando gera dados a uma taxa constante de

100 *kbits/s*, e de inatividade, quando não gera dados. Imagine ainda que o usuário esteja ativo apenas 10% do tempo (e fique ocioso, tomando cafézinho, durante os restantes 90%). Com comutação de circuitos, devem ser reservados 100 *kbits/s* para cada usuário durante todo o tempo. Por exemplo, com TDM, se um quadro de 1 segundo for dividido em 10 compartimentos de tempo de 100 milissegundos cada, então seria alocado um compartimento de tempo por quadro a cada usuário.

Desse modo, o enlace de comutação de circuitos pode suportar somente 10 ($= 1 \text{ Mbit/s}/100 \text{ kbits/s}$) usuários simultaneamente. Com a comutação de pacotes, a probabilidade de haver um usuário específico ativo é 0,1 (i.e., 10%). Se houver 35 usuários, a probabilidade de haver 11 ou mais usuários ativos ao mesmo tempo é de mais ou menos 0,0004. (O Problema P8 dos Exercícios de Fixação demonstra como essa probabilidade é calculada.) Quando houver dez ou menos usuários ativos simultâneos (a probabilidade de isso acontecer é 0,9996), a taxa agregada de chegada de dados é menor ou igual a 1 *Mbit/s*, que é a taxa de saída do enlace. Assim, quando houver dez ou menos usuários ativos, pacotes de usuários fluirão pelo enlace essencialmente sem atraso, como é o caso na comutação de circuitos. Quando houver mais de dez usuários ativos ao mesmo tempo, a taxa agregada de chegada de pacotes excederá a capacidade de saída do enlace, e a fila de saída começará a crescer. (E continuará a crescer até que a velocidade agregada de entrada caia novamente para menos de 1 *Mbit/s*, ponto em que o comprimento da fila começará a diminuir.) Como a probabilidade de haver mais de dez usuários ativos é ínfima nesse exemplo, a comutação de pacotes apresenta, em essência, o mesmo desempenho da comutação de circuitos, *mas o faz para mais de três vezes o número de usuários*.

Vamos considerar agora um segundo exemplo simples. Suponha que haja dez usuários e que um deles de repente gere 1.000 pacotes de 1.000 *bits*, enquanto os outros nove permanecem inativos e não geram pacotes. Com comutação de circuitos TDM de dez compartimentos de tempo por quadro, e cada quadro consistindo em 1.000 *bits*, o usuário ativo poderá usar somente seu único compartimento por quadro para transmitir dados, enquanto os nove compartimentos restantes em cada quadro continuarão ociosos. Dez segundos se passarão antes que todo o 1 milhão de *bits* de dados do usuário ativo seja transmitido. No caso da comutação de pacotes, o usuário ativo poderá enviá-los continuamente à taxa total de 1 *Mbit/s*, visto que não haverá outros gerando pacotes que precisem ser multiplexados com os dele. Nesse caso, todos os dados do usuário ativo serão transmitidos dentro de 1 segundo.

Os exemplos citados ilustram duas maneiras pelas quais o desempenho da comutação de pacotes pode ser superior ao da comutação de circuitos. Também destacam a diferença crucial entre as duas formas de compartilhar a taxa de transmissão de um enlace entre vários fluxos de *bits*. A comutação de circuitos aloca previamente a utilização do enlace de transmissão independentemente da demanda, com desperdício de tempo de enlace desnecessário alocado e não utilizado. A comutação de pacotes, por outro lado, aloca utilização de enlace *por demanda*. A capacidade de transmissão do enlace será compartilhada pacote por pacote somente entre usuários que tenham pacotes que precisam ser transmitidos pelo enlace.

Embora tanto a comutação de pacotes quanto a de circuitos coexistam nas redes de telecomunicação de hoje, a tendência é, sem dúvida, a comutação de pacotes. Até mesmo muitas das atuais redes de telefonia de comutação de circuitos estão migrando aos poucos para a comutação de pacotes. Em especial, redes telefônicas usam comutação de pacotes na parte cara de uma chamada telefônica para o exterior.

1.3.3 Uma rede de redes

Vimos anteriormente que sistemas finais (PCs, *smartphones*, servidores Web, servidores de correio eletrônico e assim por diante) conectam-se à Internet por meio de um provedor local (ISP). Este pode fornecer conectividade tanto com ou sem fio, utilizando diversas tecnologias de acesso, que incluem DSL, cabo, FTTH, Wi-Fi e telefone celular. Observe que o provedor local não precisa ser uma operadora de telefonia ou uma empresa de TV a cabo:

pode ser, por exemplo, uma universidade (que oferece acesso à Internet para os alunos, os funcionários e o corpo docente) ou uma empresa (que oferece acesso para seus funcionários). Mas conectar usuários finais e provedores de conteúdo a um provedor de acesso (ISP) é apenas uma pequena peça do quebra-cabeça que é interligar os bilhões de sistemas finais que compõem a Internet. Para resolver esse quebra-cabeça, os próprios ISPs de acesso precisam estar interconectados. Isso é feito criando uma *rede de redes* – entender essa frase é a chave para entender a Internet.

Com o passar dos anos, a rede de redes que forma a Internet evoluiu para uma estrutura bastante complexa. Grande parte dessa evolução é controlada pela política nacional e pela economia, e não por considerações de desempenho. Para entender a estrutura de rede da Internet de hoje, vamos criar, de modo incremental, uma série de estruturas de rede, com cada nova estrutura sendo uma aproximação melhor da Internet complexa que temos. Lembre-se de que o objetivo dominante é interconectar os provedores de acesso de modo que todos os sistemas finais possam enviar pacotes entre si. Um método ingênuo seria fazer cada ISP se conectar *diretamente* a cada outro ISP. Esse projeto em malha, é evidente, seria muito caro para os ISPs, pois exigiria que cada ISP tivesse um enlace de comunicação separado para as centenas de milhares de outros ISPs do mundo inteiro.

Nossa primeira estrutura de rede, a *Estrutura de Rede 1*, interconecta todos os ISPs de acesso a um *único ISP de trânsito global*. Nosso ISP de trânsito global (imaginário) é uma rede de roteadores e enlaces de comunicação que não apenas se espalha pelo planeta, mas também tem pelo menos um roteador próximo de cada uma das centenas de milhares de ISPs de acesso. Claro, seria muito dispendioso para o ISP global montar essa rede tão extensa. Para que seja lucrativo, ele naturalmente cobraria de cada um dos ISPs de acesso pela conectividade, com o preço dependendo, mas nem sempre diretamente proporcional à quantidade de tráfego que um ISP de acesso troca com o ISP global. Como o ISP de acesso paga ao ISP de trânsito global, ele é considerado um *cliente*, e o ISP de trânsito global é considerado um *provedor*.

Agora, se alguma empresa montar e operar um ISP de trânsito global que seja lucrativo, então será natural para outras empresas montarem seus próprios ISPs de trânsito global e competirem com o original. Isso leva à *Estrutura de Rede 2*, que consiste em centenas de milhares de ISPs de acesso e múltiplos ISPs de trânsito global. Os ISPs de acesso decerto preferem a Estrutura de Rede 2 à Estrutura de Rede 1, pois agora podem escolher entre os provedores de trânsito global concorrentes comparando seus preços e serviços. Note, porém, que os próprios ISPs de trânsito global precisam se interconectar: caso contrário, os ISPs de acesso conectados a um dos provedores de trânsito global não poderiam se comunicar com os ISPs de acesso conectados aos outros provedores de trânsito global.

A Estrutura de Rede 2, que acabamos de descrever, é uma hierarquia de duas camadas com provedores de trânsito global residindo no nível superior e os ISPs de acesso no nível inferior. Com isso, considera-se que os ISPs de trânsito global não são capazes de chegar perto de todo e qualquer ISP de acesso, mas que é economicamente desejável fazer isso. Na realidade, embora alguns ISPs tenham uma cobertura global impressionante e se conectem diretamente com muitos ISPs de acesso, nenhum tem presença em toda e qualquer cidade do mundo. Em vez disso, em determinada região, pode haver um **ISP regional** ao qual os ISPs de acesso na região se conectam. Cada ISP regional, então, se conecta a **ISPs de nível 1**. Estes são semelhantes ao nosso ISP de trânsito global (imaginário); mas os ISPs de nível 1, que realmente existem, não têm uma presença em cada cidade do mundo. Existe mais ou menos uma dúzia de ISPs de nível 1, incluindo organizações como Level 3 Communications, AT&T, Sprint e NTT. É interessante que nenhum grupo sanciona oficialmente o *status* de nível 1; como diz o ditado – se você tiver que perguntar se é membro de um grupo, provavelmente não é.

Retornando a essa rede de redes, não apenas existem vários ISPs de nível 1 concorrentes, mas pode haver múltiplos ISPs regionais concorrentes em uma região. Em tal hierarquia, cada ISP de acesso paga ao regional ao qual se conecta, e cada ISP regional paga ao ISP de nível 1 ao qual se interliga. (Um ISP de acesso também pode se conectar diretamente a um

ISP de nível 1, quando pagará ao ISP de nível 1). Assim, existe uma relação cliente-provedor em cada nível da hierarquia. Observe que os ISPs de nível 1 não pagam a ninguém, pois estão no topo. Para complicar as coisas ainda mais, em algumas regiões pode haver um ISP regional maior (talvez se espalhando por um país inteiro) ao qual os ISPs regionais menores nessa região se conectam: o ISP regional maior, então, se conecta a um ISP de nível 1. Por exemplo, na China, existem ISPs de acesso em cada cidade, que se conectam a ISPs provinciais, que se ligam a ISPs nacionais, que, por fim, se interligam a ISPs de nível 1 (Tian, 2012). Chamamos a essa hierarquia multinível, que ainda é apenas uma aproximação grosseira da Internet de hoje, *Estrutura de Rede 3*.

Para montar uma rede que se assemelhe mais à Internet de hoje, temos que acrescentar pontos de presença (PoPs, do inglês *points of presence*), *multi-homing*, emparelhamento e pontos de troca da Internet (IXPs, do inglês *Internet exchange points*) à Estrutura de Rede 3. Existem PoPs em todos os níveis da hierarquia, exceto para o nível de baixo (ISP de acesso). Um **PoP** é simplesmente um grupo de um ou mais roteadores (no mesmo local) na rede do provedor, onde os ISPs clientes podem se conectar ao ISP provedor. Para que uma rede do cliente se conecte ao PoP de um provedor, ele pode alugar um enlace de alta velocidade de um provedor de telecomunicações de terceiros para conectar diretamente um de seus roteadores a um roteador no PoP. Qualquer ISP (exceto os de nível 1) pode decidir efetuar o *multi-home*, ou seja, conectar-se a dois ou mais ISPs provedores. Assim, por exemplo, um ISP de acesso pode efetuar *multi-home* com dois ISPs regionais, ou então com dois ISPs regionais e também com um ISP de nível 1. De modo semelhante, um ISP regional pode efetuar *multi-home* com vários ISPs de nível 1. Quando um ISP efetua *multi-home*, ele pode continuar a enviar e receber pacotes na Internet, mesmo que um de seus provedores apresente uma falha.

Como vimos, os ISPs clientes pagam aos seus ISPs provedores para obter interconectividade global com a Internet. O valor que um ISP cliente paga a um ISP provedor reflete a quantidade de tráfego que ele troca com o provedor. Para reduzir esses custos, um par de ISPs próximos no mesmo nível da hierarquia pode **emparelhar**, ou seja, conectar diretamente suas redes, de modo que todo o tráfego entre elas passe pela conexão direta, em vez de passar por intermediários mais à frente. Quando dois ISPs são emparelhados, isso em geral é feito em acordo, ou seja, nenhum ISP paga ao outro. Como já dissemos, os ISPs de nível 1 também são emparelhados uns com os outros, sem taxas. Para uma discussão fácil de ler sobre emparelhamento e relações cliente-provedor, consulte Van der Berg (2008). Nesses mesmos termos, uma empresa de terceiros pode criar um **IXP**, que é um ponto de encontro onde vários ISPs podem se emparelhar (quase sempre em um prédio isolado com seus próprios nós de comutação (Ager, 2012). Existem mais de 600 IXPs na Internet hoje (PeeringDB 2020). Referimo-nos a esse ecossistema – consistindo em ISPs de acesso, ISPs regionais, ISPs de nível 1, PoPs, *multi-homing*, emparelhamento e IXPs – como *Estrutura de Rede 4*.

Agora, chegamos finalmente na *Estrutura de Rede 5*, que descreve a Internet de hoje. Essa estrutura, ilustrada na Figura 1.15, se baseia no topo da Estrutura de Rede 4 acrescentando **redes de provedor de conteúdo**. A Google é um dos principais exemplos dessa rede de provedor de conteúdo. No momento, a Google tem 19 grandes *datacenters* espalhados pela América do Norte, Europa, Ásia, América do Sul e Austrália, sendo que cada um possui dezenas ou centenas de milhares de servidores. Além disso, a Google possui *datacenters* menores, acomodando apenas centenas de servidores cada; estes *datacenters* menores muitas vezes ficam localizados junto aos IXPs. Os *datacenters* da Google são todos interconectados por meio de uma rede TCP/IP privada, que se espalha pelo mundo inteiro, mas, apesar disso, é separada da Internet pública. O importante é que essa rede privada só transporta tráfego de/para servidores da Google. Como vemos na Figura 1.15, a rede privada da Google tenta “contornar” as camadas mais altas da Internet emparelhando (sem custo) com outros ISPs de nível mais baixo, seja conectando diretamente ou interligando com eles em IXPs (Labovitz, 2010). Entretanto, como muitos ISPs de acesso ainda só podem ser alcançados transitando por redes de nível 1, a rede da Google também se conecta a ISPs de nível 1 e paga a esses ISPs pelo tráfego que troca com eles. Criando sua própria rede, um provedor

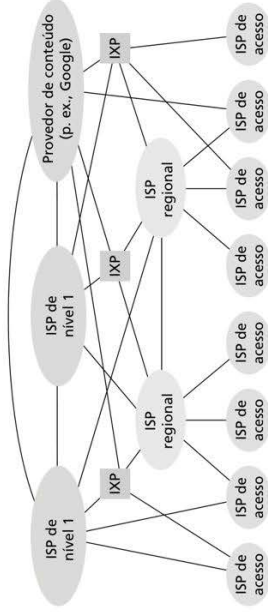


Figura 1.15 Interconexão de ISPs.

de conteúdo não apenas reduz seus pagamentos aos ISPs da camada mais alta, mas também tem maior controle de como seus serviços são entregues aos usuários finais. A infraestrutura de rede da Google é descrita com mais detalhes na Seção 2.6.

Resumindo, a topologia da Internet é complexa, consistindo em uma dúzia ou mais de ISPs de nível 1 e centenas de milhares de ISPs de níveis mais baixos. A cobertura dos ISPs é bastante diversificada; alguns abrangem vários continentes e oceanos e outros se limitam a pequenas regiões geográficas. Os ISPs de níveis mais baixos conectam-se a ISPs de níveis mais altos e estes se interconectam uns com os outros. Usuários e provedores de conteúdo são clientes de ISPs de níveis mais baixos e estes são clientes de ISPs de níveis mais altos. Nos últimos anos, os principais provedores de conteúdo também têm criado suas próprias redes e se conectam diretamente a ISPs de níveis mais baixos, quando possível.

1.4 ATRASO, PERDA E VAZÃO EM REDES DE COMUTAÇÃO DE PACOTES

Na Seção 1.1, dissemos que a Internet pode ser vista como uma infraestrutura que fornece serviços a aplicações distribuídas que são executadas nos sistemas finais. De modo ideal, gostaríamos que os serviços da Internet transferissem tantos dados quanto desejamos entre dois sistemas finais, de modo instantâneo, sem nenhuma perda. É uma pena, mas esse é um objetivo muito ambicioso, algo inalcançável. Em vez disso, as redes de computadores, necessariamente, restringem a vazão (a quantidade de dados por segundo que podem ser transferidos) entre sistemas finais, apresentam atrasos entre sistemas finais e podem perder pacotes. Por um lado, infelizmente as leis físicas da realidade introduzem atraso e perda, bem como restringem a vazão. Por outro, já que as redes de computadores têm esses problemas, existem muitas questões fascinantes sobre como lidar com eles — questões mais do que suficientes para preencher um curso de redes de computadores e motivar milhares de teses de doutorado! Nesta seção, começaremos a examinar e quantificar atraso, perda e vazão em redes de computadores.

1.4.1 Uma visão geral de atraso em redes de comutação de pacotes

Lembre-se de que um pacote começa em um sistema final (a origem), passa por uma série de roteadores e termina sua jornada em outro sistema final (o destino). Quando um pacote viaja de um nó (sistema final ou roteador) ao nó subsequente (sistema final ou roteador), sofre, ao

longo desse caminho, diversos tipos de atraso em cada nó. Os mais importantes deles são o **atraso de processamento nodal**, o **atraso de fila**, o **atraso de transmissão** e o **atraso de propagação**; juntos, eles se acumulam para formar o **atraso nodal total**. O desempenho de muitas aplicações da Internet — como busca, navegação Web, *e-mail*, mapas, mensagens instantâneas e voz sobre IP — é bastante afetado por atrasos na rede. Para entender a fundo a comutação de pacotes e redes de computadores, é preciso entender a natureza e a importância desses atrasos.

Tipos de atraso

Vamos examinar esses atrasos no contexto da Figura 1.16. Como parte de sua rota fim a fim entre origem e destino, um pacote é enviado do nó de origem por meio do roteador A até o roteador B. Nossa meta é caracterizar o atraso nodal no roteador A. Note que este tem um enlace de saída que leva ao roteador B. Esse enlace é precedido de uma fila (também conhecida como *buffer*). Quando o pacote chega ao roteador A, vindo do nó de origem, o roteador examina o cabeçalho do pacote para determinar o enlace de saída apropriado e então o direciona a esse enlace. Nesse exemplo, o enlace de saída para o pacote é o que leva ao roteador B. Um pacote pode ser transmitido por um enlace apenas se não houver nenhum outro sendo transmitido por ele e se não houver outros à sua frente na fila. Se o enlace estiver ocupado, ou com pacotes à espera, o recém-chegado entrará na fila.

Atraso de processamento

O tempo exigido para examinar o cabeçalho do pacote e determinar para onde direcioná-lo é parte do **atraso de processamento**, que pode também incluir outros fatores, como o tempo necessário para verificar os erros em *bits* existentes no pacote que ocorreram durante a transmissão dos *bits* desde o nó de origem ao roteador A. Atrasos de processamento em roteadores de alta velocidade em geral são da ordem de microssegundos, ou menos. Depois desse processamento nodal, o roteador direciona o pacote à fila que precede o enlace com o roteador B. (No Capítulo 4, estudaremos os detalhes da operação de um roteador.)

Atraso de fila

O pacote sofre um **atraso de fila** enquanto espera para ser transmitido no enlace. O tamanho desse atraso dependerá da quantidade de outros pacotes que chegarem antes e que já estiverem na fila esperando pela transmissão no enlace. Se a fila estiver vazia, e nenhum outro pacote estiver sendo transmitido naquele momento, então o tempo de fila de nosso pacote será zero. Por outro lado, se o tráfego estiver intenso e houver muitos pacotes também esperando para ser transmitidos, o atraso de fila será longo. Em breve, veremos que o número de pacotes que um determinado pacote provavelmente encontrará ao chegar é uma função da intensidade e da natureza do tráfego que está chegando à fila. Na prática, atrasos de fila podem ser da ordem de micro a milissegundos.

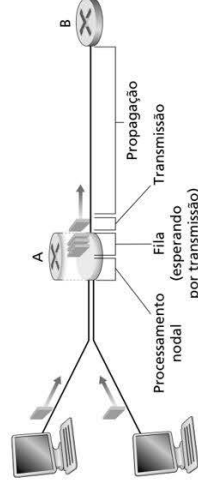


Figura 1.16 O atraso nodal no roteador A.

Atraso de transmissão

Admitindo-se que pacotes são transmitidos segundo a estratégia de “o primeiro a chegar será o primeiro a ser processado”, como é comum em redes de comutação de pacotes, o nosso sistema poderá ser transmitido depois de todos os que chegaram antes terem sido enviados. Denominemos o tamanho do pacote como L bits e a velocidade de transmissão do enlace do roteador A ao roteador B como R bits/s. Por exemplo, para um enlace Ethernet de 10 Mbits/s, a velocidade é $R = 10$ Mbits/s; para um enlace Ethernet de 100 Mbits/s, a velocidade é $R = 100$ Mbits/s. O **atraso de transmissão** é L/R . Esta é a quantidade de tempo exigida para enfileirar (i.e., transmitir) todos os bits do pacote para o enlace. Na prática, atrasos de transmissão costumam ser da ordem de micro a milissegundos.

Atraso de propagação

Assim que é lançado no enlace, um bit precisa se propagar até o roteador B. O tempo necessário para propagar o bit desde o início do enlace até o roteador B é o **atraso de propagação**. O bit se propaga à velocidade de propagação do enlace, a qual depende do meio físico (i.e., fibra ótica, par de fios de cobre trançado e assim por diante) e está na faixa de

$$2 \cdot 10^8 \text{ m/s a } 3 \cdot 10^8 \text{ m/s}$$

que é menor ou igual à velocidade da luz. O atraso de propagação é a distância entre dois roteadores dividida pela velocidade de propagação. Isto é, o atraso de propagação é d/s , sendo d a distância entre o roteador A e o roteador B, e s é a velocidade de propagação do enlace. Assim que o último bit do pacote se propagar até o nó B, ele e todos os outros bits precedentes serão armazenados no roteador B. Então, o processo inteiro continua, agora com o roteador B executando a retransmissão. Em redes WAN (do inglês *wide-area networks* — redes de área ampla), os atrasos de propagação são da ordem de milissegundos.

Comparação entre atrasos de transmissão e de propagação

Os principiantes na área de redes de computadores às vezes têm dificuldade para entender a diferença entre atrasos de transmissão e de propagação. Ela é sutil, mas importante. O atraso de transmissão é a quantidade de tempo necessária para o roteador enfileirar o pacote para fora; é uma função do comprimento do pacote e da taxa de transmissão do enlace, mas nada tem a ver com a distância entre os roteadores. O atraso de propagação, por outro lado, é o tempo que leva para um bit se propagar de um roteador até o seguinte; é uma função da distância entre os roteadores, mas nada tem a ver com o comprimento do pacote ou com a taxa de transmissão do enlace.

Podemos esclarecer melhor as noções de atrasos de transmissão e de propagação com uma analogia. Considere uma rodovia que tenha um posto de pedágio a cada 100 quilômetros, como mostrado na Figura 1.17. Imagine que os trechos da rodovia entre os postos de pedágio sejam enlaces e que os postos de pedágio sejam roteadores. Suponha que os carros trafeguem (i.e., se propaguem) pela rodovia a uma velocidade de 100 km/h (i.e., quando o carro sai de um posto de pedágio, acelera instantaneamente até 100 km/h e mantém essa velocidade entre os dois postos de pedágio). Agora, considere que dez carros viajem em comboio, um atrás do outro, em ordem fixa.

Imagine que cada carro seja um bit e que o comboio seja um pacote. Suponha ainda que cada posto de pedágio libere (i.e., transmita) um carro a cada 12 segundos, que seja tarde da noite e que os carros do comboio sejam os únicos na estrada. Por fim, imagine que, ao chegar a um posto de pedágio, o primeiro carro do comboio aguarde na entrada até que os outros nove cheguem e formem uma fila atrás dele. (Assim, o comboio inteiro deve ser “armazenado” no posto de pedágio antes de começar a ser “reenviado”). O tempo necessário para que todo o comboio passe pelo posto de pedágio e volte à estrada é de (10 carros)/

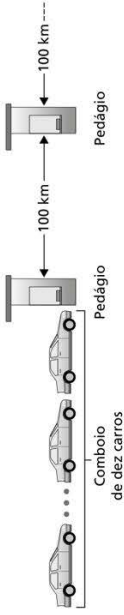


Figura 1.17 Analogia do comboio.

(5 carros/minuto) = 2 minutos, semelhante ao atraso de transmissão em um roteador. O tempo necessário para um carro trafegar da saída de um posto de pedágio até o próximo é de $(100 \text{ km})/(100 \text{ km/h}) = 1$ hora, semelhante ao atraso de propagação. Portanto, o tempo decorrido entre o instante em que o comboio é “armazenado” em frente a um posto de pedágio até o momento em que é “armazenado” em frente ao seguinte é a soma do atraso de transmissão e do atraso de propagação — nesse exemplo, 62 minutos.

Vamos explorar um pouco mais essa analogia. O que aconteceria se o tempo de liberação do comboio no posto de pedágio fosse maior do que o tempo que um carro leva para trafegar entre dois postos? Por exemplo, suponha que os carros trafeguem a uma velocidade de 1.000 km/h e que o pedágio libere um carro por minuto. Então, o atraso de trânsito entre dois postos de pedágio é de 6 minutos e o tempo de liberação do comboio no posto de pedágio é de 10 minutos. Nesse caso, os primeiros carros do comboio chegarão ao segundo posto de pedágio antes que os últimos carros saiam do primeiro posto. Essa situação também acontece em redes de comutação de pacotes — os primeiros bits de um pacote podem chegar a um roteador enquanto muitos dos remanescentes ainda estão esperando para ser transmitidos pelo roteador precedente.

Se uma imagem vale mil palavras, então uma animação vale um milhão de palavras. O site de apoio deste livro apresenta uma animação interativa que ilustra e compara o atraso de transmissão com o de propagação. Recomendamos que o leitor visite essa animação. Smith (2009) também oferece uma discussão bastante fácil de ler sobre atrasos de propagação, enfileiramento e transmissão.

Se d_{proc} , d_{fila} , d_{trans} e d_{prop} forem, respectivamente, os atrasos de processamento, de fila, de transmissão e de propagação, então o atraso nodal total é dado por:

$$d_{\text{nodal}} = d_{\text{proc}} + d_{\text{fila}} + d_{\text{trans}} + d_{\text{prop}}$$

A contribuição desses componentes do atraso pode variar significativamente. Por exemplo, d_{prop} pode ser desprezível (p. ex., 2 microssegundos) para um enlace que conecta dois roteadores no mesmo campus universitário; contudo, é de centenas de milissegundos para dois roteadores interconectados por um enlace de satélite geostacionário e pode ser o termo dominante no d_{nodal} . De maneira semelhante, d_{trans} pode variar de desprezível a significativo. Sua contribuição costuma ser desprezível para velocidades de transmissão de 10 Mbits/s e mais altas (p. ex., em LANs); contudo, pode ser de centenas de milissegundos para grandes pacotes de Internet enviados por enlaces de modems discados de baixa velocidade. O atraso de processamento, d_{proc} , é quase sempre desprezível; no entanto, tem forte influência sobre a produtividade máxima de um roteador, que é a velocidade máxima com que ele pode encaminhar pacotes.

1.4.2 Atraso de fila e perda de pacote

O mais complicado e interessante componente do atraso nodal é o atraso de fila, d_{fila} . Na verdade, o atraso de fila é tão importante e interessante em redes de computadores que milhares de artigos e numerosos livros já foram escritos sobre ele (Bertsekas, 1991; Kleinrock, 1975; Kleinrock, 1976). Neste livro, faremos apenas uma discussão intuitiva, de alto nível, sobre o

atraso de fila; o leitor mais curioso pode consultar alguns dos livros citados (ou até mesmo escrever uma tese sobre o assunto!). Diferente dos três outros atrasos (a saber, d_{proc} , d_{trans} e d_{prop}), o atraso de fila pode variar de pacote a pacote. Por exemplo, se dez pacotes chegarem a uma fila vazia ao mesmo tempo, o primeiro pacote transmitido não sofrerá nenhum atraso de fila, ao passo que o último sofrerá um atraso relativamente grande (enquanto espera que os outros nove sejam transmitidos). Por conseguinte, para se caracterizar um atraso de fila, são utilizadas em geral medições estatísticas, tais como atraso de fila médio, variância do atraso de fila e a probabilidade de que ele exceda um valor especificado.

Quando o atraso de fila é grande e quando é insignificante? A resposta a essa pergunta depende da velocidade de transmissão do enlace, da taxa com que o tráfego chega à fila e de sua natureza, isto é, se periodicamente ou de modo intermitente, em rajadas. Para entendermos melhor, vamos adotar α para representar a taxa média com que os pacotes chegam à fila (α é medida em pacotes/segundo). Lembra-se de que R é a taxa de transmissão, isto é, a taxa (em bits/segundo) com que os bits são retirados da fila. Suponha também, para simplificar, que todos os pacotes tenham L bits. Então, a taxa média com que os bits chegam à fila é La bits/s. Por fim, imagine que a fila seja muito longa, de modo que possa conter um número infinito de bits. A razão La/R , denominada **intensidade de tráfego**, costuma desempenhar um papel importante na estimativa do tamanho do atraso de fila. Se $La/R > 1$, então a velocidade média com que os bits chegam à fila excederá aquela com que eles podem ser transmitidos para fora da fila. Nessa situação desastrosa, a fila tenderá a aumentar sem limite e o atraso de fila tenderá ao infinito! Por conseguinte, uma das regras de ouro da engenharia de tráfego é: *projete seu sistema de modo que a intensidade de tráfego não seja maior do que 1*.

Agora, considere o caso em que $La/R \leq 1$. Aqui, a natureza do tráfego influencia o atraso de fila. Por exemplo, se pacotes chegarem periodicamente — isto é, se chegar um pacote a cada L/R segundos —, então todos os pacotes chegarão a uma fila vazia e não haverá atraso. Por outro lado, se chegarem em rajadas, mas periodicamente, poderá haver um significativo atraso de fila médio. Por exemplo, suponha que N pacotes cheguem ao mesmo tempo a cada $(L/R)/N$ segundos. Então, o primeiro pacote transmitido não sofrerá atraso de fila; o segundo terá um atraso de L/R segundos e, de modo mais geral, o n ésimo pacote transmitido terá um atraso de fila de $(n-1)L/R$ segundos. Deixamos como exercício para o leitor o cálculo do atraso de fila médio para esse exemplo.

Os dois exemplos de chegadas periódicas que acabamos de descrever são um tanto acadêmicos. Em geral, o processo de chegada a uma fila é *aleatório* — isto é, não segue um padrão, e os intervalos de tempo entre os pacotes são ao acaso. Nessa hipótese mais realista, a quantidade La/R quase sempre não é suficiente para caracterizar por completo a estatística do atraso. Não obstante, é útil para entender intuitivamente a extensão do atraso de fila. Em especial, se a intensidade de tráfego for próxima de zero, então as chegadas de pacotes serão poucas e bem espaçadas, e é improvável que um pacote que esteja chegando encontre outro na fila. Consequentemente, o atraso de fila médio será próximo de zero. Por outro lado, quando a intensidade de tráfego for próxima de 1, haverá intervalos de tempo em que a velocidade de chegada excederá a capacidade de transmissão (em razão das variações na taxa de chegada do pacote), e uma fila será formada durante esses períodos; quando a taxa de chegada for menor do que a capacidade de transmissão, a extensão da fila diminuirá. Todavia, à medida que a intensidade de tráfego se aproxima de 1, o comprimento médio da fila fica cada vez maior. A dependência qualitativa entre o atraso de fila médio e a intensidade de tráfego é mostrada na Figura 1.18.

Um aspecto importante a observar na Figura 1.18 é que, quando a intensidade de tráfego se aproxima de 1, o atraso de fila médio aumenta depressa. Uma pequena porcentagem de aumento na intensidade resulta em um aumento muito maior no atraso, em termos de porcentagem. Talvez você já tenha percebido esse fenômeno na estrada. Se você dirige regularmente por uma estrada que costuma ser congestionada, o fato de ela estar sempre assim significa que a intensidade de tráfego é próxima de 1. Se algum evento causar um tráfego um pouco maior do que o normal, as demoras que você sofrerá poderão ser enormes.

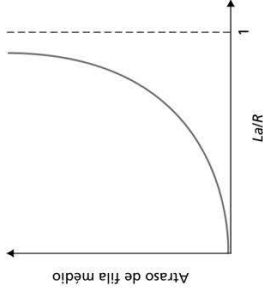


Figura 1.18 Dependência entre atraso de fila médio e intensidade de tráfego.

Para compreender um pouco mais os atrasos de fila, visite o *site* de apoio do livro, que apresenta uma animação interativa sobre uma fila. Se você aumentar a taxa de chegada do pacote o suficiente de forma que a intensidade do tráfego exceda 1, verá a fila aumentar ao longo do tempo.

Perda de pacotes

Na discussão anterior, admitimos que a fila é capaz de conter um número infinito de pacotes. Na realidade, a capacidade da fila que precede um enlace é finita, embora o seu tamanho máximo dependa bastante do projeto e do custo do nó de comutação. Como a capacidade da fila é finita, na verdade os atrasos de pacote não se aproximam do infinito quando a intensidade de tráfego se aproxima de 1. O que acontece de fato é que um pacote pode chegar e encontrar uma fila cheia. Sem espaço disponível para armazená-lo, o roteador o **descarta**; isto é, ele será **perdido**. Esse excesso em uma fila pode ser observado novamente na animação interativa quando a intensidade do tráfego é maior do que 1.

Do ponto de vista de um sistema final, uma perda de pacote é vista como um pacote que foi transmitido para o núcleo da rede, mas sem nunca ter emergido dele no destino. A fração de pacotes perdidos aumenta com o aumento da intensidade de tráfego. Por conseguinte, o desempenho em um nó costuma ser medido não apenas em termos de atraso, mas também da probabilidade de perda de pacotes. Como discutiremos nos capítulos subsequentes, um pacote perdido pode ser retransmitido fim a fim para garantir que todos os dados sejam transferidos da origem ao local de destino.

1.4.3 Atraso fim a fim

Até o momento, nossa discussão focalizou o atraso nodal, isto é, em um único roteador. Concluiremos essa discussão considerando brevemente o atraso da origem ao destino. Para entender esse conceito, suponha que haja $N-1$ roteadores entre a máquina de origem e a de destino. Imagine também que a rede não esteja congestionada (e, portanto, os atrasos de fila sejam desprezíveis), que o atraso de processamento em cada roteador e na máquina de origem seja d_{proc} , que a taxa de transmissão de saída de cada roteador e da máquina de origem seja R bits/s, e que o atraso de propagação em cada enlace seja d_{prop} . Os atrasos nodais se acumulam e resultam em um atraso fim a fim,

$$d_{\text{fim a fim}} = N(d_{\text{proc}} + d_{\text{trans}} + d_{\text{prop}}) \quad (1.2)$$

em que, mais uma vez, $d_{\text{mns}} = L/R$, e L é o tamanho do pacote. Note que a Equação 1.2 é uma generalização da Equação 1.1, na qual não levamos em conta os atrasos de processamento e propagação. Convidamos você a generalizar a Equação 1.2 para o caso de atrasos heterogêneos nos nós e para o caso de um atraso de fila médio em cada nó.

Traceroute

Para perceber o que é fato o atraso em uma rede de computadores, podemos utilizar o Traceroute, programa de diagnóstico que pode ser executado em qualquer hospedeiro da Internet. Quando o usuário especifica um nome de hospedeiro de destino, o programa no hospedeiro de origem envia vários pacotes especiais em direção àquele destino. Ao seguir seu caminho até o destino, esses pacotes passam por uma série de roteadores. Um deles recebe um desses pacotes especiais e envia à origem uma curta mensagem, contendo o nome e o endereço do roteador.

Mais especificamente, suponha que haja $N - 1$ roteadores entre a origem e o destino. Então, a fonte enviará N pacotes especiais à rede, e cada um deles estará endereçado ao destino final. Esses N pacotes especiais serão marcados de 1 a N , sendo a marca do primeiro pacote 1 e a do último, N . Assim que o enésimo roteador recebe o enésimo pacote com a marca n , não envia o pacote a seu destino, mas uma mensagem à origem. Quando o hospedeiro de destino recebe o pacote N , também envia uma mensagem à origem, que registra o tempo transcorrido entre o envio de um pacote e o recebimento da mensagem de retorno correspondente. A origem registra também o nome e o endereço do roteador (ou do hospedeiro de destino) que retorna a mensagem. Dessa maneira, a origem pode reconstruir a rota tomada pelos pacotes que vão da origem ao destino e pode determinar os atrasos de ida e volta para todos os roteadores intermediários. Na realidade, o programa Traceroute repete o processo que acabamos de descrever três vezes, de modo que a fonte envia, na verdade, $3 \cdot N$ pacotes ao destino. O RFC 1393 descreve detalhadamente o Traceroute.

Eis um exemplo de resultado do programa Traceroute, no qual a rota traçada ia do hospedeiro de origem gaia.cs.umass.edu (na Universidade de Massachusetts) até um host pederio no departamento de ciências da computação na Universidade Sorbonne, em Paris (antiga UPMC). O resultado tem seis colunas: a primeira é o valor n descrito, isto é, o número do roteador ao longo da rota; a segunda é o nome do roteador; a terceira é o endereço do roteador (na forma xxx.xxx.xxx.xxx); as últimas três são os atrasos de ida e volta para três tentativas. Se a fonte receber menos do que três mensagens de qualquer roteador determinado (em virtude da perda de pacotes na rede), o Traceroute coloca um asterisco logo após o número do roteador e registra menos do que três tempos de duração de ida e volta para aquele roteador.

```

1 gw-vlan-2451.cs.umass.edu (128.119.245.1) 1,899 ms 3,266 ms 3,280 ms
2 j-cs-gw-int-10-240.cs.umass.edu (10.119.240.254) 1,296 ms 1,276 ms 1,245 ms
3 n5-rt-1-l-xe-2-l-0-gw.umass.edu (128.119.3.33) 2,237 ms 2,217 ms 2,187 ms
4 core1-rt-et-5-2-0-gw.umass.edu (128.119.0.9) 0,351 ms 0,392 ms 0,380 ms
5 border1-rt-et-5-0-0-gw.umass.edu (192.80.83.102) 0,345 ms 0,345 ms 0,344 ms
6 nox300gw1-umass-re.nox.org (192.5.89.101) 3,260 ms 0,416 ms 3,127 ms
7 nox300gw1-umass-re.nox.org (192.5.89.101) 3,165 ms 7,326 ms 7,311 ms
8 198.71.45.237 (198.71.45.237) 77,826 ms 77,246 ms 77,744 ms
9 renater-lbl-gw.mxl.par.fr.geant.net (62.40.124.70) 79,357 ms 77,729 ms 79,152 ms
10 193.51.180.109 (193.51.180.109) 78,379 ms 79,936 ms 80,042 ms
11 * 193.51.180.109 (193.51.180.109) 80,640 ms *
12 * 195.221.127.182 (195.221.127.182) 78,408 ms *
13 195.221.127.182 (195.221.127.182) 80,686 ms 80,796 ms 78,434 ms
14 r-upmcl.reseau.jussieu.fr (134.157.254.10) 78,399 ms * 81,353 ms

```

No exemplo anterior, há 14 roteadores entre a origem e o destino. Quase todos eles têm um nome e todos têm endereços. Por exemplo, o nome do roteador 4 é `core1-rt-et-5-2-0-gw.umass.edu`, e seu endereço é 128.119.0.9. Examinando os dados apresentados para ele, verificamos que, na primeira das três tentativas, o atraso de ida e volta entre a origem e o roteador foi de 0,351 ms. Os atrasos de ida e volta para as duas tentativas subsequentes foram 0,392 e 0,380 ms, e incluem todos os atrasos que acabamos de discutir, ou seja, de transmissão, de propagação, de processamento do roteador e de fila.

Como o atraso de fila varia com o tempo, o atraso de ida e volta do pacote n enviado a um roteador n pode, às vezes, ser maior do que o do pacote $n+1$ enviado ao roteador $n+1$. Realmente, observamos esse fenômeno no exemplo anterior: o atraso do roteador 12 é menor que o do roteador 11! Observe também o grande aumento no atraso de ida e volta do roteador 7 para o roteador 8. Isso se deve ao enlace de fibra ótica transatlântica entre os dois, o que dá origem a um atraso de propagação relativamente grande. Diversos programas de *software* gratuitos oferecem uma interface gráfica para o Traceroute; um dos nossos favoritos é o PingPlotter (PingPlotter 2020).

Sistema final, aplicativo e outros atrasos

Além dos atrasos de processamento, transmissão e de propagação, os sistemas finais podem adicionar outros atrasos significativos. Por exemplo, um sistema final que quer transmitir um pacote para uma mídia compartilhada (p. ex., como em um cenário WiFi ou *modem a cabo*) pode, *intencionalmente*, atrasar sua transmissão como parte de seu protocolo para compartilhar a mídia com outros sistemas finais; vamos analisar tais protocolos em detalhes no Capítulo 6. Outro importante atraso é o atraso de empacotamento de mídia, o qual está presente nos aplicativos VoIP (voz sobre IP). No VoIP, o remetente deve primeiro carregar um pacote com voz digitalizada e codificada antes de transmitir o pacote para a Internet. Esse tempo para carregar um pacote — chamado de atraso de empacotamento — pode ser significativo e ter impacto sobre a qualidade visível pelo usuário de uma chamada VoIP. Esse assunto será explorado mais adiante nos exercícios de fixação no final deste capítulo.

1.4.4 Vazão nas redes de computadores

Além do atraso e da perda de pacotes, outra medida de desempenho importante em redes de computadores é a vazão fim a fim. Para definir vazão, considere a transferência de um arquivo grande do hospedeiro A para o hospedeiro B por uma rede de computadores. Essa transferência pode ser, por exemplo, um arquivo de vídeo extenso de um computador para outro. A **vazão instantânea** em um dado momento é a taxa (em *bits/s*) à qual o hospedeiro B está recebendo o arquivo. (Muitos aplicativos exibem a vazão instantânea durante os *downloads* na interface do usuário — talvez você já tenha observado isso! Você poderia medir o atraso fim a fim e a vazão de *download* entre o seu sistema e servidores na Internet usando a aplicação de teste de velocidade [Speedtest, 2020].) Se o arquivo consistir em F *bits* e a transferência levar T segundos para o hospedeiro B receber todos os F *bits*, então a **vazão média** da transferência do arquivo é F/T *bits/s*. Para algumas aplicações, como a telefonia via Internet, é desejável ter um atraso baixo e uma vazão instantânea acima de algum limiar (p. ex., superior a 24 *kbits/s* para aplicações de telefonia via Internet, e superior a 256 *kbits/s* para algumas aplicações de vídeo em tempo real). Para outras aplicações, incluindo as de transferência de arquivo, o atraso não é importante, mas é recomendado ter a vazão mais alta possível.

Para obter uma visão mais detalhada do importante conceito de vazão, vamos analisar alguns exemplos. A Figura 1.19(a) mostra dois sistemas finais, um servidor e um cliente, conectados por dois enlaces de comunicação e um roteador. Considere a vazão para uma transferência de arquivo do servidor para o cliente. Suponha que R_s seja a taxa do enlace entre o servidor e o roteador; e R_c seja a taxa do enlace entre o roteador e o cliente. Imagine

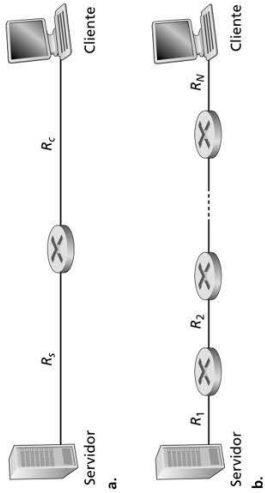


Figura 1.19 Vazão para uma transferência de arquivo do servidor ao cliente.

que os únicos *bits* enviados na rede inteira sejam os do servidor para o cliente. Agora vem a pergunta: nesse cenário ideal, qual é a vazão servidor-para-cliente? Para responder, pense nos *bits* como um *fluido* e nos enlaces de comunicação como *tubos*. Claro, o servidor não pode enviar os *bits* através de seu enlace a uma taxa mais rápida do que R_s *bits/s*; e o roteador não pode encaminhar os *bits* a uma taxa mais rápida do que R_i *bits/s*. Se $R_s < R_i$, então os *bits* enviados pelo servidor “fluirão” diretamente pelo roteador e chegarão ao cliente a uma taxa de R_s *bits/s*, gerando uma vazão de R_s *bits/s*. Se, por outro lado, $R_s > R_i$, então o roteador não poderá encaminhar os *bits* tão rápido quanto ele os recebe. Neste caso, os *bits* somente deixarão o roteador a uma taxa R_i , dando uma vazão fim a fim de R_i . (Observe também que, se os *bits* continuarem a chegar no roteador a uma taxa R_s , e a deixá-lo a uma taxa R_i , o acúmulo de *bits* esperando para transmissão ao cliente só aumentará – uma situação extremamente indesejável!) Assim, para essa rede simples de dois enlaces, a vazão é $\min\{R_s, R_i\}$, ou seja, é a taxa de transmissão do **enlace de gargalo**. Após determinar a vazão, agora podemos aproximar o tempo que leva para transferir um arquivo grande de F *bits* do servidor ao cliente como $F/\min\{R_s, R_i\}$. Para um exemplo específico, suponha que você está fazendo o *download* de um arquivo MP3 de $F = 32$ milhões de *bits*, o servidor tem uma taxa de transmissão de $R_s = 2$ *Mbits/s*, e você tem um enlace de acesso de $R_i = 1$ *Mbit/s*. O tempo necessário para transferir o arquivo é, então, 32 segundos. Claro que essas expressões para tempo de vazão e de transferência são apenas aproximações, já que elas não consideram os atrasos para armazenar-e-reenviar e de processamento, bem como questões relativas a protocolos.

A Figura 1.19(b) agora mostra uma rede com N enlaces entre o servidor e o cliente, com as taxas de transmissão dos N enlaces sendo $R_1, R_2, \dots, R_N$. Aplicando a mesma análise da rede de dois enlaces, descobrimos que a vazão para uma transferência de arquivo do servidor ao cliente é $\min\{R_1, R_2, \dots, R_N\}$, a qual é novamente a taxa de transmissão do enlace de gargalo ao longo do caminho entre o servidor e o cliente.

Agora considere outro exemplo motivado pela Internet de hoje. A Figura 1.20(a) mostra dois sistemas finais, um servidor e um cliente, conectados a uma rede de computadores. Considere a vazão para uma transferência de arquivo do servidor ao cliente. O servidor está conectado à rede com um enlace de acesso de taxa R_s , e o cliente está conectado à rede com um enlace de acesso de R_c . Agora suponha que todos os enlaces no núcleo da rede de comunicação tenham taxas de transmissão muito altas, muito maiores do que R_s e R_c . De fato, hoje, o núcleo da Internet está superdimensionado com enlaces de alta velocidade que sofrem pouco congestionamento. Suponha, também, que os únicos *bits* que estão sendo enviados em toda a rede sejam os do servidor para o cliente. Já que o núcleo da rede de computadores é como um tubo largo neste exemplo, a taxa em que os *bits* correm da origem ao destino é novamente o mínimo de R_s e R_c , ou seja, vazão = $\min\{R_s, R_c\}$. Portanto, o fator restritivo para vazão na Internet de hoje é, em geral, a rede de acesso.

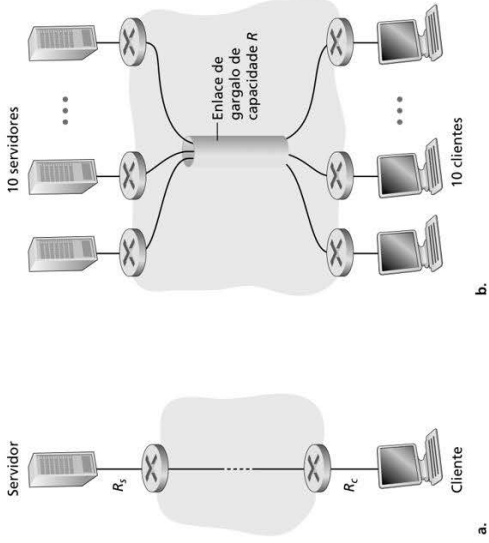


Figura 1.20 Vazão fim a fim: (a) O cliente baixa um arquivo do servidor; (b) 10 clientes fazem o download com 10 servidores.

Para um exemplo final, considere a Figura 1.20(b), na qual existem dez servidores e dez clientes conectados ao núcleo da rede de computadores. Nesse exemplo, dez *downloads* simultâneos estão sendo realizados, envolvendo dez pares cliente-servidor. Suponha que esses *downloads* sejam o único tráfego na rede no momento. Como mostrado na figura, há um enlace no núcleo que é atravessado por todos os dez *downloads*. Considere R a taxa de transmissão desse enlace R . Imagine que todos os enlaces de acesso do servidor possuem a mesma taxa R_s , todos os enlaces de acesso do cliente possuem a mesma taxa R_c , e a taxa de transmissão de todos os enlaces no núcleo – com exceção de um enlace comum de taxa R – sejam muito maiores do que R_s , R_c e R . Agora perguntamos: quais são as vazões de *download*? É claro que se a taxa do enlace comum, R , é grande – digamos, 100 vezes maior do que R_s e R_c –, então a vazão para cada *download* será novamente $\min\{R_s, R_c\}$. Mas e se essa taxa for da mesma ordem que R_s e R_c ? Qual será a vazão nesse caso? Vamos observar um exemplo específico. Suponha que $R_s = 2$ *Mbits/s*, $R_c = 1$ *Mbit/s*, $R = 5$ *Mbits/s*, e o enlace comum divide sua taxa de transmissão por igual entre 10 *downloads*. Então, o gargalo para cada *download* não se encontra mais na rede de acesso, mas é o enlace compartilhado no núcleo, que somente fornece para cada *download* 500 *kbits/s* de vazão. Desse modo, a vazão fim a fim é agora reduzida a 500 *kbits/s* por *download*.

Os exemplos nas Figuras 1.19 e 1.20(a) mostram que a vazão depende das taxas de transmissão dos enlaces sobre os quais os dados fluem. Vimos que quando não há tráfego interveniente, a vazão pode apenas ser aproximada como a taxa de transmissão mínima ao longo do caminho entre a origem e o local de destino. O exemplo na Figura 1.20(b) mostra que, de modo geral, a vazão depende não somente das taxas de transmissão dos enlaces ao longo do caminho, mas também do tráfego interveniente. Em especial, um enlace com uma alta taxa de transmissão pode, apesar disso, ser o enlace de gargalo para uma transferência de arquivo, caso muitos outros fluxos de dados estejam também passando por aquele enlace. Analisaremos em mais detalhes a vazão em redes de computadores nos exercícios de fixação e nos capítulos subsequentes.

1.5 CAMADAS DE PROTOCOLO E SEUS MODELOS DE SERVIÇO

Até aqui, nossa discussão demonstrou que a Internet é um sistema *extremamente* complexo e que possui muitos componentes: inúmeras aplicações e protocolos, vários tipos de sistemas finais e conexões entre eles, nós de comutação de pacotes, além de vários tipos de mídia em nível de enlace. Dada essa enorme complexidade, há alguma esperança de organizar a arquitetura de rede ou, ao menos, nossa discussão sobre ela? Felizmente, a resposta a ambas as perguntas é sim.

1.5.1 Arquitetura de camadas

Antes de tentarmos organizar nosso raciocínio sobre a arquitetura da Internet, vamos procurar uma analogia humana. Na verdade, lidamos com sistemas complexos o tempo todo em nosso dia a dia. Imagine se alguém pedisse que você descrevesse, por exemplo, o sistema de uma companhia aérea. Como você encontraria a estrutura para descrever esse sistema complexo que tem agências de emissão de passagens, pessoal para embarcar a bagagem, pessoal para ficar no portão de embarque, pilotos, aviões, controle de tráfego aéreo e um sistema mundial de roteamento de aeronaves? Um modo poderia ser apresentar a relação de uma série de ações que você realiza (ou que outros executam para você) quando voa por uma empresa aérea. Você compra a passagem, despacha suas malas, dirige-se ao portão de embarque e, por fim, entra no avião, que decola e segue uma rota até seu destino. Após a aterrissagem, você desembarca no portão designado e recupera suas malas. Se a viagem foi ruim, você reclama na agência que lhe vendeu a passagem (esforço em vão). Esse cenário é ilustrado na Figura 1.21.

Já podemos notar aqui algumas analogias com redes de computadores: você está sendo despachado da origem ao destino pela companhia aérea; um pacote é despachado da máquina de origem à máquina de destino na Internet. Mas essa não é exatamente a analogia que buscamos. Estamos tentando encontrar alguma *estrutura* na Figura 1.21. Observando-a, notamos que há uma função referente à passagem em cada porta: há também uma função de bagagem para passageiros que já apresentaram o bilhete e uma de portão de embarque para os que já apresentaram o tiquete e despacharam as malas. Para passageiros que já passaram pelo portão de embarque (i.e., aqueles que já apresentaram a passagem, despacharam a bagagem e passaram pelo portão), há uma função de decolagem e de aterrissagem e, durante o voo, uma função de roteamento do avião. Isso sugere que podemos examinar a funcionalidade na Figura 1.21 na *horizontal*, como mostra a Figura 1.22.

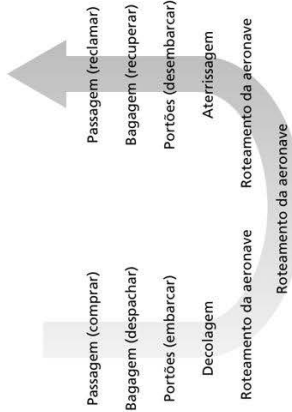


Figura 1.21 Uma viagem de avião: ações.

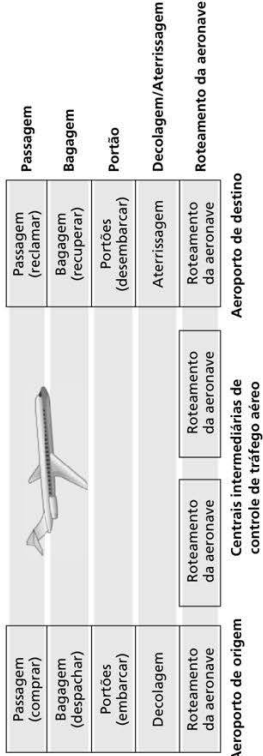


Figura 1.22 Camadas horizontais da funcionalidade de linha aérea.

A Figura 1.22 dividiu a funcionalidade da linha aérea em camadas, provendo uma estrutura com a qual podemos discutir a viagem aérea. Note que cada camada, combinada com as que estão abaixo dela, implementa alguma funcionalidade, algum *serviço*. Na camada da passagem aérea e abaixo dela, é realizada a transferência “balcão-de-linha-aérea-balcão-de-linha-aérea” de um passageiro. Na camada de bagagem e abaixo dela, é realizada a transferência “despacho-de-bagagem-recuperação-de-bagagem” de um passageiro e de suas malas. Note que a camada da bagagem provê esse serviço apenas para a pessoa que já apresentou o bilhete. Na camada do portão, é realizada a transferência “portão-de-embarque-portão-de-desembarque” do viajante e de suas malas. Na camada de decolagem/aterrissagem, é realizada a transferência “pista-a-pista” de passageiros e de suas bagagens. Cada camada provê seu serviço (1) realizando certas ações dentro dela (p. ex., na camada do portão, embarcar e desembarcar pessoas de um avião) e (2) utilizando os serviços da camada imediatamente inferior (p. ex., na do portão, aproveitando o serviço de transferência “pista-a-pista” de passageiros da camada de decolagem/aterrissagem).

Uma arquitetura de camadas nos permite discutir uma parcela específica e bem definida de um sistema grande e complexo. Essa simplificação tem considerável valor intrínseco, pois provê modularidade, tornando muito mais fácil modificar a execução do serviço prestado pela camada. Contudo que a camada forneça o mesmo serviço para a que está acima e use os mesmos serviços da que vem abaixo dela, o restante do sistema permanece inalterado quando a sua realização é modificada. (Note que modificar a implementação de um serviço é muito diferente de mudar o serviço em si!) Por exemplo, se as funções de portão fossem modificadas (digamos que passassem a embarcar e desembarcar passageiros por ordem de altura), o restante do sistema da linha aérea permaneceria inalterado, já que a camada do portão continuaria a prover a mesma função (embarcar e desembarcar passageiros); ela apenas executaria aquela função de maneira diferente após a alteração. Para sistemas grandes e complexos que são atualizados constantemente, a capacidade de modificar a realização de um serviço sem afetar outros componentes do sistema é outra vantagem importante da divisão em camadas.

Camadas de protocolo

Mas chega de linhas aéreas! Vamos agora voltar nossa atenção a protocolos de rede. Para prover uma estrutura para o projeto, projetistas de rede organizam protocolos – e o *hardware* e o *software* de rede que os executam – em **camadas**. Cada protocolo pertence a uma das camadas, assim como cada função na arquitetura de linha aérea da Figura 1.22 pertence a uma camada. Mais uma vez, estamos interessados no conjunto de **serviços** que uma camada oferece à camada acima dela – denominado **modelo de serviço**. Assim como em nosso exemplo da linha aérea, cada camada provê seu serviço (1) executando certas ações dentro dela e (2) utilizando os serviços da camada diretamente abaixo dela. Por exemplo, os

serviços providos pela camada n podem incluir entrega confiável de mensagens de uma extremidade da rede à outra, que pode ser implementada utilizando um serviço não confiável de entrega de mensagem fim a fim da camada $n-1$ e adicionando funcionalidade da camada n para detectar e retransmitir mensagens perdidas.

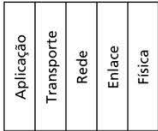
Uma camada de protocolo pode ser implementada em *software*, em *hardware*, ou em uma combinação dos dois. Protocolos de camada de aplicação – como HTTP e SMTP – quase sempre são realizados em *software* nos sistemas finais; o mesmo acontece com protocolos de camada de transporte. Como a camada física e as de enlace de dados são responsáveis pelo manuseio da comunicação por um enlace específico, em geral são executadas em uma placa de interface de rede (p. ex., placas de interface Ethernet ou WiFi) associadas a determinado enlace. A camada de rede muitas vezes é uma execução mista de *hardware* e *software*. Note também que, tal como as funções na arquitetura em camadas da linha aérea eram distribuídas entre os vários aeroportos e centrais de controle de tráfego aéreo que compunham o sistema, um protocolo de camada n é *distribuído* entre sistemas finais, nós de comutação de pacote e outros componentes que formam a rede. Isto é, há sempre uma parte de um protocolo de camada n em cada componente de rede.

O sistema de camadas de protocolos tem vantagens conceituais e estruturais (RFC 3439). Como vimos, a divisão em camadas proporciona um modo estruturado de discutir componentes de sistemas. A modularidade facilita a atualização de componentes de sistema. Devemos mencionar, no entanto, que alguns pesquisadores e engenheiros de rede se opõem veementemente ao sistema de camadas (Wakeman, 1992). Uma desvantagem potencial é que uma camada pode duplicar a funcionalidade de uma camada inferior. Por exemplo, muitas pilhas de protocolos oferecem serviço de recuperação de erros para cada enlace e também de modo fim a fim. Uma segunda desvantagem é que a funcionalidade em uma camada pode necessitar de informações (p. ex., um valor de marca de tempo) que estão presentes somente em outra, o que infringe o objetivo de separação de camadas.

Quando tomados em conjunto, os protocolos das várias camadas são denominados **pilha de protocolos**. A pilha de protocolos da Internet é formada por cinco camadas: física, de enlace, de rede, de transporte e de aplicação, como mostra a Figura 1.23. Se você verificar o sumário, verá que organizamos este livro utilizando as camadas da pilha de protocolos da Internet. Fazemos uma **abordagem top-down** (de cima para baixo), primeiro abordando a camada de aplicação e prosseguindo para baixo.

Camada de aplicação

A camada de aplicação é onde residem aplicações de rede e seus protocolos. A camada de aplicação da Internet inclui muitos protocolos, tais como o HTTP (que provê requisição e transferência de documentos pela Web), o SMTP (que provê transferência de mensagens de correio eletrônico) e o FTP (que provê a transferência de arquivos entre dois sistemas finais). Veremos que certas funções de rede, como a tradução de nomes mnemônicos, que são dados



Pilha de protocolos da Internet de cinco camadas

Figura 1.23 A pilha de protocolos da Internet.

a sistemas finais da Internet (p. ex., de `<www.ietf.org>` para um endereço de rede de 32 bits), também são executadas com a ajuda de um protocolo de camada de aplicação, no caso, o sistema de nomes de domínio (DNS, do inglês *domain name system*). Veremos no Capítulo 2 que é muito fácil criar nossos próprios novos protocolos de camada de aplicação.

Um protocolo de camada de aplicação é distribuído por diversos sistemas finais, e a aplicação em um sistema final utiliza o protocolo para trocar pacotes de informação com a aplicação em outro sistema final. Chamaremos de **mensagem** esse pacote de informação na camada de aplicação.

Camada de transporte

A camada de transporte da Internet carrega mensagens da camada de aplicação entre os lados do cliente e servidor de uma aplicação. Há dois protocolos de transporte na Internet: TCP e UDP* (do inglês *User Datagram Protocol* – Protocolo de Datagrama de Usuário), e qualquer um pode levar mensagens da camada de aplicação. O TCP provê serviços orientados à conexão para suas aplicações. Alguns desses serviços são a entrega garantida de mensagens da camada de aplicação ao destino e controle de fluxo (i.e., adequação das velocidades do remetente e do receptor). O TCP também fragmenta mensagens longas em segmentos mais curtos e provê mecanismo de controle de congestionamento, de modo que uma origem reduza sua velocidade de transmissão quando a rede está congestionada. O protocolo UDP provê serviço não orientado à conexão para suas aplicações. É um serviço econômico que não oferece confiabilidade, nem controle de fluxo ou de congestionamento. Neste livro, chamaremos de **segmento** um pacote da camada de transporte.

Camada de rede

A camada de rede da Internet é responsável pela movimentação, de um hospedeiro para outro, de pacotes da camada de rede, conhecidos como **datagramas**. O protocolo de camada de transporte da Internet (TCP ou UDP) em um hospedeiro de origem passa um segmento da camada de transporte e um endereço de destino à camada de rede, exatamente como você passaria ao serviço de correios uma carta com um endereço de destinatário. A camada de rede então provê o serviço de entrega do segmento à camada de transporte no hospedeiro de destino.

Essa camada inclui o famoso protocolo IP, que define os campos no datagrama e o modo como os sistemas finais e os roteadores agem nesses campos. Existe apenas um único protocolo IP, e todos os componentes da Internet que têm uma camada de rede devem executá-lo. A camada de rede da Internet também contém protocolos de roteamento que determinam as rotas que os datagramas seguem entre origens e destinos. A Internet tem muitos protocolos de roteamento. Como vimos na Seção 1.3, a Internet é uma rede de redes e, dentro de uma delas, o administrador pode executar qualquer protocolo de roteamento. Embora a camada de rede contenha o protocolo IP e também numerosos outros de roteamento, ela quase sempre é denominada apenas camada IP, refletindo o fato de que ele é o elemento fundamental que mantém a integridade da Internet.

Camada de enlace

A camada de rede roteia um datagrama por meio de uma série de roteadores entre a origem e o destino. Para levar um pacote de um nó (hospedeiro ou roteador) ao nó seguinte na rota, a camada de rede depende dos serviços da camada de enlace. Em especial, em cada nó, a camada de rede passa o datagrama para a camada de enlace, que o entrega, ao longo da rota, ao nó seguinte, no qual o datagrama é passado da camada de enlace para a de rede.

*N. de T.: Na verdade existem muitos outros protocolos de transporte, mas esses são os mais comuns.

Os serviços prestados pela camada de enlace dependem do protocolo específico empregado no enlace. Por exemplo, alguns desses protocolos proveem entrega garantida entre enlaces, isto é, desde o nó transmissor, passando por um único enlace, até o nó receptor. Note que esse serviço confiável de entrega é diferente do de entrega garantida do TCP, que provê serviço de entrega garantida de um sistema final a outro. Exemplos de protocolos de camadas de enlace são Ethernet, WiFi e o protocolo DOCSIS da rede de acesso por cabo. Como datagramas normalmente precisam transitar por diversos enlaces para trem da origem ao destino, serão manuseados por diferentes protocolos de camada de enlace em diversos enlaces ao longo de sua rota, podendo ser manuseados por Ethernet em um e por PPP no seguinte. A camada de rede receberá um serviço diferente de cada um dos variados protocolos de camada de enlace. Neste livro, pacotes de camada de enlace serão denominados **quadros**.

Camada física

Enquanto a tarefa da camada de enlace é movimentar quadros inteiros de um elemento da rede até um elemento adjacente, a da camada física é movimentar os *bits individuais* que estão dentro do quadro de um nó para o seguinte. Os protocolos nessa camada de novo dependem do enlace e, além disso, do próprio meio de transmissão do enlace (p. ex., fios de cobre trançado ou fibra ótica monomodo). Por exemplo, a Ethernet tem muitos protocolos de camada física: um para par de fios de cobre trançado, outro para cabo coaxial, mais um para fibra e assim por diante. Em cada caso, o *bit* atravessa o enlace de um modo diferente.

1.5.2 Encapsulamento

A Figura 1.24 apresenta o caminho físico que os dados percorrem: para baixo na pilha de protocolos de um sistema final emissor, para cima e para baixo nas pilhas de protocolos de um *switch* e roteador intermediários, e depois para cima na pilha de protocolos do sistema final receptor. Como discutiremos mais adiante neste livro, ambos, roteadores e *switches*, são computadores de pacotes. De modo semelhante a sistemas finais, ambos organizam seu

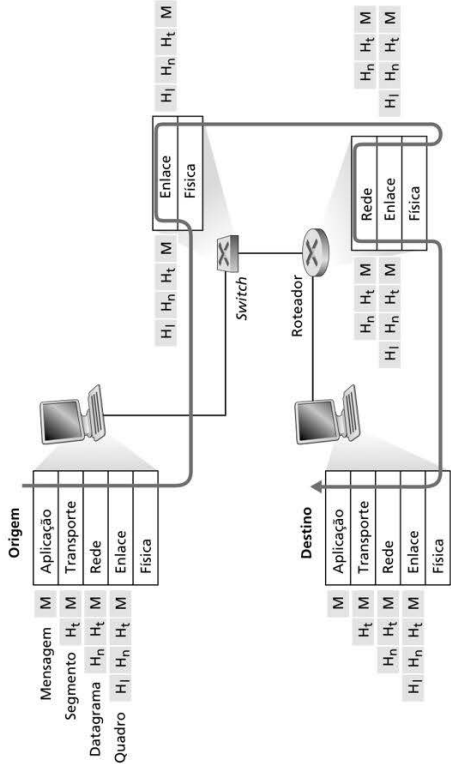


Figura 1.24 Hospedeiros, roteadores e switches; cada um contém um conjunto diferente de camadas, refletindo suas diferenças em funcionalidade.

hardware e *software* de rede em camadas. Mas não implementam *todas* as camadas da pilha de protocolos; em geral, executam apenas as camadas de baixo. Como ilustra a Figura 1.24, *switches* realizam as camadas 1 e 2; roteadores executam as camadas 1, 2 e 3. Isso significa, por exemplo, que roteadores da Internet são capazes de executar o protocolo IP (da camada 3), mas *switches* não. Veremos mais adiante que, embora não reconheçam endereços IP, *switches* são capazes de reconhecer endereços de camada 2, os da Ethernet. Note que os hospedeiros implementam todas as cinco camadas, o que é consistente com a noção de que a arquitetura da Internet concentra sua complexidade na periferia da rede.

A Figura 1.24 também ilustra o importante conceito de **encapsulamento**. Uma **mensagem da camada de aplicação** na máquina emissora (M na Figura 1.24) é passada para a camada de transporte. No caso mais simples, esta pega a mensagem e anexa informações adicionais (denominadas informações de cabeçalho de camada de transporte, H_t na Figura 1.24) que serão usadas pela camada de transporte do lado receptor. A mensagem da camada de aplicação e as informações de cabeçalho da camada de transporte, juntas, constituem o **segmento da camada de transporte**, que encapsula a mensagem da camada de aplicação. As informações adicionadas podem incluir dados que habilitem a camada de transporte do lado do receptor a entregar a mensagem à aplicação apropriada, além de *bits* de detecção de erro que permitem que o receptor determine se os *bits* da mensagem foram modificados em trânsito. A camada de transporte então passa o segmento à camada de rede, que adiciona informações de cabeçalho de camada de rede (H_r na Figura 1.24), como endereços de sistemas finais de origem e de destino, criando um **datagrama de camada de rede**. Este é então passado para a camada de enlace, que (é claro!) adicionará suas próprias informações de cabeçalho e criará um **quadro de camada de enlace**. Assim, vemos que, em cada camada, um pacote possui dois tipos de campos: campos de cabeçalho e um **campo de carga útil**. A carga útil é em geral um pacote da camada acima.

Uma analogia útil que podemos usar aqui é o envio de um memorando entre escritórios de uma empresa pelo correio de uma filial a outra. Suponha que Alice, que está em uma filial, queira enviar um memorando a Bob, que está na outra filial. O *memorando* representa a *mensagem da camada de aplicação*. Alice coloca o memorando em um envelope de correspondência interna em cuja face são escritos o nome e o departamento de Bob. O *envelope de correspondência interna* representa o *segmento da camada de transporte* – contém as informações de cabeçalho (o nome de Bob e seu departamento) e encapsula a mensagem de camada de aplicação (o memorando). Quando a central de correspondência do escritório emissor recebe o envelope, ele é colocado dentro de outro, adequado para envio pelo correio. A central de correspondência emissora também escreve o endereço postal do remetente e do destinatário no envelope postal. Nesse ponto, o *envelope postal* é análogo ao *datagrama* – encapsula o segmento de camada de transporte (o envelope de correspondência interna), que, por sua vez, encapsula a mensagem original (o memorando). O correio entrega o envelope postal à central de correspondência do escritório destinatário. Nesse local, o processo de desencapsulamento se inicia. A central de correspondência retira o memorando e o encaminhava a Bob. Este, por fim, abre o envelope e retira o memorando.

O processo de encapsulamento pode ser mais complexo do que o descrito. Por exemplo, uma mensagem grande pode ser dividida em vários segmentos de camada de transporte (que também podem ser divididos em vários datagramas de camada de rede). Na extremidade receptora, cada segmento deve ser reconstruído a partir dos datagramas que o compõem.

1.6 REDES SOB AMEAÇA

A Internet se tornou essencial para muitas instituições, incluindo empresas grandes e pequenas, universidades e órgãos do governo. Muitas pessoas também contam com a Internet para suas atividades profissionais, sociais e pessoais. Bilhões de “coisas”, incluindo

eletrodomésticos e dispositivos pessoais, estão sendo conectados à Internet. Mas atrás de toda essa utilidade e entusiasmo, existe o lado escuro, um lado no qual “vilões” tentam causar problemas em nosso cotidiano, danificando nossos computadores conectados à Internet, violando nossa privacidade e tornando inoperantes os serviços da rede dos quais dependemos.

A área de segurança trata de como esses vilões podem ameaçar as redes de computadores e como nós, futuros especialistas no assunto, podemos defender a rede contra essas ameaças ou, melhor ainda, criar novas arquiteturas intrinsecamente imunes a tais riscos. Dadas a frequência e a variedade das ameaças existentes, bem como o perigo de novos e mais destrutivos futuros ataques, a segurança se tornou um assunto principal na área de redes de computadores. Um dos objetivos deste livro é trazer as questões de segurança de rede para o primeiro plano.

Visto que ainda não temos o *know-how* em rede de computadores e em protocolos da Internet, comecemos com uma análise de alguns dos atuais problemas predominantes relacionados à segurança. Isto irá aguçar nosso apetite para discussões mais importantes nos capítulos futuros. Começamos com a pergunta: o que pode dar errado? Como as redes de computadores são vulneráveis? Quais são alguns dos tipos de ameaças predominantes hoje?

Os vilões podem colocar “malware” em seu hospedeiro por meio da Internet

Conectamos aparelhos à Internet porque queremos receber/enviar dados de/para a rede. Isso inclui todos os tipos de recursos vantajosos, como postagens no Instagram, resultados de buscas, *streaming* de música, chamadas de videoconferência, *streaming* de filmes e assim por diante. Infelizmente, no entanto, junto com esses recursos vantajosos, aparecem, os maliciosos – chamados coletivamente de **malware** – que podem entrar e infectar nossos aparelhos. Uma vez que o **malware** infecta nosso aparelho, ele é capaz de fazer coisas perversas, como apagar nossos arquivos; instalar *spyware* que coleta informações particulares, como nosso número de cartão de crédito, senhas e combinação de telas, e as envia (pela Internet, é claro!) de volta aos vilões. Nosso hospedeiro comprometido pode estar, também, envolvido em uma rede de milhares de aparelhos comprometidos, conhecida como **botnet**, a qual é controlada e utilizada pelos vilões para distribuição de *spams* ou ataques de recusa de serviço distribuídos (que serão discutidos mais adiante) contra hospedeiros direcionados.

Boa parte do **malware** existente hoje é **autorreprodutivo**: depois que infecta um hospedeiro, ele busca outros hospedeiros na Internet para infectar; destes hospedeiros recém-infectados, ele busca então invadir ainda mais sistemas. Dessa forma, o **malware** autorreprodutivo pode se disseminar em velocidade exponencial.

Os vilões podem atacar servidores e infraestrutura de redes

Um amplo grupo de ameaças à segurança pode ser classificado como **ataques de recusa de serviços** (DoS, do inglês **Denial-of-Service**). Como o nome sugere, um ataque DoS torna uma rede, hospedeiro ou outra parte da infraestrutura inutilizável por usuários verdadeiros. Servidores da Web, de *e-mail* e DNS (discutidos no Capítulo 2) e redes institucionais podem estar sujeitos aos ataques DoS. O *site* Digital Attack Map oferece uma visão dos principais ataques de DoS diários em todo o mundo (DAM, 2020). A maioria dos ataques DoS na Internet pode ser dividida em três categorias:

- **Ataque de vulnerabilidade.** Envolve o envio de algumas mensagens bem elaboradas a uma aplicação vulnerável ou a um sistema operacional sendo executado em um hospedeiro direcionado. Se a sequência correta de pacotes é enviada a uma aplicação ou sistema operacional vulnerável, o serviço pode parar ou, pior, o hospedeiro pode estragar.

- **Inundação na largura de banda.** O atacante envia um grande número de pacotes ao hospedeiro direcionado – tantos pacotes que o enlace de acesso do alvo congestionaria, impedindo os pacotes legítimos de alcançarem o servidor.
- **Inundação na conexão.** O atacante estabelece um grande número de conexões TCP semiabertas ou abertas (as conexões TCP são discutidas no Capítulo 3) no hospedeiro-alvo. O hospedeiro pode ficar tão atolado com essas conexões falsas que deixa de aceitar conexões legítimas.

Vamos agora explorar mais detalhadamente o ataque de inundação na largura de banda. Lembrando de nossa análise sobre atraso e perda na Seção 1.4.2, é evidente que se o servidor possui uma taxa de acesso de R bits/s, o atacante precisará enviar tráfego a uma taxa de, mais ou menos, R bits/s para causar dano. Se R for muito grande, uma fonte de ataque única pode não ser capaz de gerar tráfego suficiente para prejudicar o servidor. Além disso, se todo o tráfego emanar de uma fonte única, um roteador mais adiante pode conseguir detectar o ataque e bloquear todo o tráfego da fonte antes que ele se aproxime do servidor. Em um **ataque DoS distribuído** (DDoS, do inglês **Distributed DoS**), ilustrado na Figura 1.25, o atacante controla múltiplas fontes que sobrecarregam o alvo. Com essa tática, a taxa de tráfego agregada por todas as fontes controladas precisa ser, aproximadamente, R para incapacitar o serviço. Os ataques DDoS que potencializam *botnets* com centenas de hospedeiros comprometidos são uma ocorrência comum hoje em dia (DAM, 2020). Os ataques DDoS são muito mais difíceis de detectar e de prevenir do que um ataque DoS de um único hospedeiro.

Encorajamos o leitor a considerar a seguinte questão à medida que trabalhar com este livro: o que os projetistas de redes de computadores podem fazer para se protegerem contra ataques DoS? Veremos que são necessárias diferentes defesas para os três tipos de ataques DoS.

Os vilões podem analisar pacotes

Muitos usuários hoje acessam a Internet por meio de aparelhos sem fio, como *notebooks* conectados à tecnologia WiFi ou aparelhos portáteis com conexões à Internet via telefone celular (abordado no Capítulo 7). Embora o acesso onipresente à Internet seja de extrema conveniência e disponibilize novas aplicações sensacionais aos usuários móveis, ele também cria uma grande vulnerabilidade de segurança – posicionando um receptor passivo nas proximidades do transmissor sem fio, o receptor pode obter uma cópia de cada pacote transmitido! Esses pacotes podem conter todo tipo de informações confidenciais, incluindo

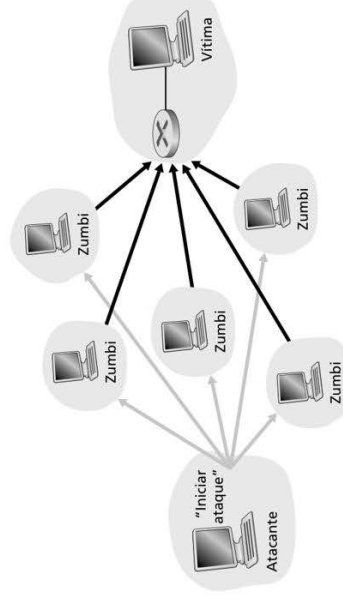


Figura 1.25 Um ataque de recusa de serviço distribuído (DDoS).

senhas, número de identificação, segredos comerciais e mensagens pessoais. Um receptor passivo que grava uma cópia de cada pacote que passa é denominado **analisador de pacotes** (*packet sniffer*).

Os analisadores também podem estar distribuídos em ambientes de conexão com fio. Nesses ambientes, como em muitas LANs Ethernet, um analisador de pacotes pode obter cópias de todos os pacotes enviados pela LAN. Como descrito na Seção 1.2, as tecnologias de acesso a cabo transmitem pacotes e são, dessa forma, vulneráveis à análise. Além disso, um vilão que quer ganhar acesso ao roteador de acesso de uma instituição ou enlace de acesso para a Internet pode instalar um analisador que faça uma cópia de cada pacote trocado com a empresa. Os pacotes capturados podem, então, ser analisados *offline* em busca de informações confidenciais.

O *software* para analisar pacotes está disponível gratuitamente em diversos *sites* da Internet e em produtos comerciais. Professores que ministraram um curso de redes passam exercícios que envolvem a escrita de um programa de reconstrução de dados da camada de aplicação e um programa analisador de pacotes. De fato, os Wireshark labs (Wireshark, 2020) associados a este texto (veja o Wireshark lab introdutório ao final deste capítulo) utilizam exatamente tal analisador de pacotes.

Como os analisadores de pacote são passivos – ou seja, não introduzem pacotes no canal – eles são difíceis de detectar. Portanto, quando enviamos pacotes para um canal sem fio, devemos aceitar a possibilidade de que alguém possa estar copiando nossos pacotes. Como você deve ter imaginado, uma das melhores defesas contra a análise de pacote envolve a criptografia, que será explicada no Capítulo 8, já que se aplica à segurança de rede.

Os vilões podem se passar por alguém de sua confiança

Por incrível que pareça, é fácil (você saberá como fazer isso à medida que ler este livro!) criar um pacote com qualquer endereço de origem, conteúdo de pacote e endereço de destino e, depois, transmiti-lo para a Internet, que, obediência, o encaminhará ao destino. Imagine que um receptor inocente (digamos, um roteador da Internet) que recebe tal pacote acredita que o endereço de origem (falso) seja confiável e então executa um comando integrado ao conteúdo do pacote (digamos, que modifica sua base de encaminhamento). A capacidade de introduzir pacotes na Internet com um endereço de origem falso é conhecida como **IP spoofing**, e é uma das muitas maneiras pelas quais um usuário pode se passar por outro.

Para resolver esse problema, precisamos de uma *autenticação do ponto final*, ou seja, um mecanismo que nos permita determinar com certeza se uma mensagem se origina de onde pensamos. Mais uma vez, sugerimos que pense em como isso pode ser feito em aplicações de rede e protocolos à medida que avança sua leitura pelos capítulos deste livro. Exploraremos mais mecanismos para autenticação da fonte no Capítulo 8.

Para encerrar esta seção, vale considerar como a Internet se tornou um local inseguro, antes de tudo. A resposta breve é que a Internet foi, a princípio, criada dessa maneira, baseada no modelo de “um grupo de usuários de confiança mútua ligados a uma rede transparente” (Blumenthal, 2001) – um modelo no qual (por definição) não há necessidade de segurança. Muitos aspectos da arquitetura inicial da Internet refletem profundamente essa noção de confiança mútua. Por exemplo, a capacidade de um usuário enviar um pacote a qualquer outro é o padrão, não um recurso solicitado/concedido, e acredita-se piamente na identidade do usuário, em vez de ser autenticada como padrão.

Mas a Internet de hoje decerto não envolve “usuários de confiança mútua”. Contudo, os usuários atuais ainda precisam se comunicar mesmo quando não confiam um no outro, podem querer se comunicar de modo anônimo, podem se comunicar indiretamente por terceiros (p. ex., *Web caches*, que serão estudados no Capítulo 2, ou agentes móveis para assistência, que serão estudados no Capítulo 7), e podem desconfiar do *hardware*, *software* e até mesmo do ar pelo qual eles se comunicam. Temos agora muitos desafios relacionados à segurança perante nós à medida que prosseguimos com o livro: devemos buscar proteção contra a análise,

disfarce da origem, ataques *man-in-the-middle*, ataques DDoS, *malware* e outros. Precisamos manter em mente que a comunicação entre usuários de confiança mútua é mais uma exceção do que uma regra. Seja bem-vindo ao mundo da moderna rede de computadores!

1.7 HISTÓRIA DAS REDES DE COMPUTADORES E DA INTERNET

Da Seção 1.1 à 1.6, apresentamos um panorama da tecnologia de redes de computadores e da Internet. Agora, você já deve saber o suficiente para impressionar sua família e seus amigos! Contudo, se quiser ser mesmo o maior sucesso na próxima festa, você deve recheiar seu discurso com pérolas da fascinante história da Internet (Segaller, 1998).

1.7.1 Desenvolvimento da comutação de pacotes: 1961–1972

Os primeiros passos da disciplina de redes de computadores e da Internet atual podem ser traçados desde o início da década de 1960, quando a rede telefônica era a rede de comunicação dominante no mundo inteiro. Lembre-se de que na Seção 1.3 dissemos que a rede de telefonia usa comutação de circuitos para transmitir informações de uma origem a um destino – uma escolha acertada, já que a voz é transmitida a uma taxa constante entre os pontos. Dada a importância cada vez maior dos computadores no início da década de 1960 e o advento de computadores com tempo compartilhado, nada seria mais natural do que considerar a questão de como interligar computadores para que pudessem ser compartilhados entre usuários geograficamente dispersos. O tráfego gerado por esses usuários provavelmente era feito por *raijadas* – períodos de atividade, como o envio de um comando a um computador remoto, seguidos de períodos de inatividade, como a espera por uma resposta ou o exame de uma resposta recebida.

Três grupos de pesquisa ao redor do mundo, sem que nenhum tivesse conhecimento do trabalho do outro (Leiner, 1998), começaram a inventar a comutação de pacotes como uma alternativa poderosa e eficiente à comutação de circuitos. O primeiro trabalho publicado sobre técnicas de comutação de pacotes foi o de Leonard Kleinrock (Kleinrock, 1961; 1964), que, naquela época, era um aluno de graduação no MIT. Usando a teoria de filas, seu trabalho demonstrou, com elegância, a eficácia da abordagem da comutação de pacotes para fontes de tráfego intermitentes (em rajadas). Em 1964, Paul Baran (Baran, 1964), do Rand Institute, começou a investigar a utilização de comutação de pacotes na transmissão segura de voz pelas redes militares, ao mesmo tempo que Donald Davies e Roger Scantlebury desenvolviam suas ideias sobre esse assunto no National Physical Laboratory (NPL), na Inglaterra.

Os trabalhos desenvolvidos no MIT, no Rand Institute e no NPL foram os alicerces do que hoje é a Internet. Mas a Internet tem uma longa história de atitudes do tipo “construir e demonstrar”, que também data do início da década de 1960. J. C. R. Licklider (DEC, 1990) e Lawrence Roberts, ambos colegas de Kleinrock no MIT, posteriormente lideraram o programa de ciência da computação na ARPA (Advanced Research Projects Agency – Agência de Projetos de Pesquisa Avançada), nos Estados Unidos. Roberts publicou um plano geral para a ARPANet (Roberts, 1967), a primeira rede de computadores por comutação de pacotes e uma ancestral direta da Internet pública de hoje. Em 1969, no Dia do Trabalho nos Estados Unidos, foi instalado o primeiro roteador de pacotes na UCLA (Universidade da Califórnia em Los Angeles) sob a supervisão de Kleinrock. Pouco tempo depois, foram instalados três roteadores de pacotes adicionais no Stanford Research Institute (SRI), na Universidade da Califórnia em Santa Bárbara e na Universidade de Utah (Figura 1.26). O incipiente precursor da Internet tinha quatro nós no final de 1969. Kleinrock recorda que a primeiríssima utilização da rede foi fazer um *login* remoto entre a UCLA e o SRI, derrubando o sistema (Kleinrock, 2004).



Mark J. Terrell/AP Photo

Figura 1.26 Um dos primeiros computadores de pacotes.

Em 1972, a ARPAnet tinha cerca de 15 nós, e foi apresentada publicamente pela primeira vez por Robert Kahn. O primeiro protocolo fim a fim entre sistemas finais da ARPAnet, conhecido como protocolo de controle de rede (NCP, do inglês *network-control protocol*), estava concluído (RFC 001). Com um protocolo fim a fim à disposição, a escrita de aplicações tornou-se possível. Em 1972, Ray Tomlinson, da BBN,* escreveu o primeiro programa de *e-mail*.

1.7.2 Redes proprietárias e interligação de redes: 1972 a 1980

A ARPAnet inicial era uma rede isolada, fechada. Para se comunicar com uma máquina da ARPAnet, era preciso estar ligado a um outro processador de interface de mensagens (IMP, do inglês *interface message processor*) dessa rede. Do início a meados de 1970, surgiram novas redes independentes de comutação de pacotes: ALOHAnet, uma rede de micro-ondas ligando universidades das ilhas do Havaí (Abramson, 1970), bem como as redes de pacotes por satélite (RFC 829) e por rádio (Kahn, 1978) da DARPA; Telenet, uma rede comercial de comutação de pacotes da BBN baseada na tecnologia ARPAnet; Cyclades, uma rede de comutação de pacotes pioneira na França, montada por Louis Pouzin (Think, 2002); redes de tempo compartilhado como a Tymnet e a rede GE Information Services, entre outras que

*N. de T.: Acrônimo para Bolt, Berenick and Newman, nome da empresa que desenvolveu os primeiros IMPs da ARPAnet.

surgiram no final da década de 1960 e início da década de 1970 (Schwartz, 1977): rede SNA da IBM (1969–1974), cujo trabalho comparava-se ao da ARPAnet (Schwartz, 1977).

O número de redes estava crescendo. Hoje, com perfeita visão do passado, podemos perceber que aquela era a hora certa para desenvolver uma arquitetura abrangente para conectar redes. O trabalho pioneiro de interconexão de redes, sob o patrocínio da DARPA (Defense Advanced Research Projects Agency – Agência de Projetos de Pesquisa Avançada de Defesa), criou basicamente uma *rede de redes*, e foi realizado por Vinton Cerf e Robert Kahn (Cerf, 1974); o termo *internetting* foi cunhado para descrever esse trabalho.

Esses princípios de arquitetura foram incorporados ao TCP. As primeiras versões desse protocolo, contudo, eram muito diferentes do TCP de hoje. Elas combinavam uma entrega sequencial confiável de dados via retransmissão por sistema final (que ainda faz parte do TCP de hoje) com funções de envio (que hoje são desempenhadas pelo IP). As primeiras experiências com o TCP, combinadas com o reconhecimento da importância de um serviço de transporte fim a fim não confiável, sem controle de fluxo, para aplicações como voz em pacotes, levaram à separação entre IP e TCP e ao desenvolvimento do protocolo UDP. Os três protocolos fundamentais da Internet que temos hoje – TCP, UDP e IP – estavam conceitualmente disponíveis no final da década de 1970.

Além das pesquisas sobre a Internet realizadas pela DARPA, muitas outras atividades importantes relacionadas ao trabalho em rede estavam em andamento. No Havaí, Norman Abramson estava desenvolvendo a ALOHAnet, uma rede de pacotes por rádio que permitia que vários lugares remotos das ilhas havaianas se comunicassem entre si. O ALOHA (Abramson, 1970) foi o primeiro protocolo de acesso múltiplo que permitiu que usuários geograficamente dispersos compartilhassem um único meio de comunicação *broadcast* (uma frequência de rádio). Metcalfe e Boggs se basearam no trabalho de Abramson sobre protocolo de múltiplo acesso quando desenvolveram o protocolo Ethernet (Metcalfe, 1976) para redes compartilhadas de transmissão *broadcast* por fio. O interessante é que o protocolo Ethernet de Metcalfe e Boggs foi motivado pela necessidade de conectar vários PCs, impressoras e discos compartilhados (Perkins, 1994). Há 25 anos, bem antes da revolução do PC e da explosão das redes, Metcalfe e Boggs estavam lançando as bases para as LANs de PCs de hoje.

1.7.3 Proliferação de redes: 1980 a 1990

Ao final da década de 1970, cerca de 200 máquinas estavam conectadas à ARPAnet. Ao final da década de 1980, o número de máquinas ligadas à Internet pública, uma confederação de redes muito parecida com a Internet de hoje, alcançou 100 mil. A década de 1980 foi uma época de formidável crescimento.

Grande parte daquele crescimento foi consequência de vários esforços distintos para criar redes de computadores para interligar universidades. A BITNET processava *e-mails* e fazia transferência de arquivos entre diversas universidades do nordeste dos Estados Unidos. A CSNET (Computer Science Network – rede da ciência de computadores) foi formada para interligar pesquisadores de universidades que não tinham acesso à ARPAnet. Em 1986, foi criada a NSFNET para prover acesso a centros de supercomputação patrocinados pela National Science Foundation (NSF). Partindo de uma velocidade inicial de 56 *kbits/s*, ao final da década, o *backbone* da NSFNET estaria funcionando a 1,5 *Mbits/s* e servindo como *backbone* primário para a interligação de redes regionais.

Na comunidade da ARPAnet, já estavam sendo encaixados muitos dos componentes finais da arquitetura da Internet de hoje. No dia 1º de janeiro de 1983, o TCP/IP foi adotado oficialmente como o novo padrão de protocolo de máquinas para a ARPAnet (em substituição ao protocolo NCP). Pela importância do evento, o dia da transição do NCP para o TCP/IP (RFC 801) foi marcado com antecedência – a partir daquele dia, todas as máquinas tiveram de adotar o TCP/IP. No final da década de 1980, foram agregadas importantes extensões ao TCP para implementação do controle de congestionamento baseado em hospedeiros (Jacobson, 1988). Também foi desenvolvido o sistema de nomes de domínios

(DNS, do inglês *domain name system*) utilizado para mapear nomes da Internet mnemônicos (p. ex., *gaia.cs.umass.edu*) para seus endereços IP de 32 *bits* (RFC 1034).

Em paralelo ao desenvolvimento da ARPAnet (que em sua maior parte deve-se aos Estados Unidos), no início da década de 1980, os franceses lançaram o projeto Minitel, um plano ambicioso para levar as redes de dados para todos os lares. Patrocinado pelo governo francês, o sistema consistia em uma rede pública de comutação de pacotes (baseada no conjunto de protocolos X.25, que usava circuitos virtuais), servidores Minitel e terminais baratos com *modems* de baixa velocidade embutidos. O Minitel transformou-se em um enorme sucesso em 1984, quando o governo francês forneceu, gratuitamente, um terminal para toda residência francesa que quisesse. O sistema incluía *sites* de livre acesso – como o da lista telefônica – e também particulares, que cobravam uma taxa de cada usuário baseada no tempo de utilização. No seu auge, em meados de 1990, o Minitel oferecia mais de 20 mil serviços, que iam desde *home banking* até bancos de dados especializados para pesquisa. Estava presente em grande parte dos lares franceses dez anos antes sequer de a maioria dos norte-americanos ouvir falar de Internet.

1.7.4 A explosão da Internet: a década de 1990

A década de 1990 estreou com vários eventos que simbolizaram a evolução contínua e a comercialização iminente da Internet. A ARPAnet, a progenitora da Internet, deixou de existir. Em 1991, a NSFNET extinguiu as restrições que impunha à sua utilização com finalidades comerciais, mas, em 1995, perdeu seu mandato quando o tráfego de *backbone* da Internet passou a ser carregado por provedores de serviços.

O principal evento da década de 1990, no entanto, foi o surgimento da World Wide Web, que levou a Internet para os lares e as empresas de milhões de pessoas no mundo inteiro. A Web serviu também como plataforma para a habilitação e o oferecimento de centenas de novas aplicações, inclusive busca (p. ex., Google e Bing), comércio pela Internet (p. ex., Amazon e eBay) e redes sociais (p. ex., Facebook).

A Web foi inventada no CERN (European Center for Nuclear Physics – Centro Europeu para Física Nuclear) por Tim Berners-Lee entre 1989 e 1991 (Berners-Lee, 1989), com base em ideias originadas de trabalhos anteriores sobre hipertexto realizados por Vannevar Bush (Bush, 1945), na década de 1940, e por Ted Nelson (Xanadu, 2012), na década de 1960. Berners-Lee e seus companheiros desenvolveram versões iniciais de HTML, HTTP, um servidor Web e um navegador (*browser*) – os quatro componentes fundamentais da Web. Por volta de 1993, havia cerca de 200 servidores Web em operação, e esse conjunto era apenas um prenúncio do que estava por vir. Nessa época, vários pesquisadores estavam desenvolvendo navegadores Web com interfaces gráficas de usuário (GUI, do inglês *graphical user interface*), entre eles Marc Andreessen, que liderou o desenvolvimento do popular navegador Mosaic, junto com Jim Clark, que formaram a Mosaic Communications, que mais tarde se transformou na Netscape Communications Corporation (Cusumano, 1998; Quittner, 1998). Em 1995, estudantes universitários estavam usando navegadores Mosaic e Netscape para navegar na Web diariamente. Na época, empresas – grandes e pequenas – começaram a operar servidores e a realizar transações comerciais pela Web. Em 1996, a Microsoft começou a desenvolver navegadores, dando início à guerra entre Netscape e Microsoft, vencida pela última alguns anos mais tarde (Cusumano, 1998).

A segunda metade da década de 1990 foi um período de tremendo crescimento e inovação, com grandes corporações e milhares de novas empresas criando produtos e serviços para a Internet. No final do milênio, a Internet dava suporte a centenas de aplicações populares, entre elas quatro de enorme sucesso:

- *e-mail*, incluindo anexos e correio eletrônico com acesso pela Web;
- a Web, incluindo navegação pela Web e comércio pela Internet;
- serviço de mensagem instantânea, com listas de contato;
- compartilhamento *peer-to-peer* (P2P) de arquivos MP3, cujo pioneiro foi o Napster.

O interessante é que as duas primeiras dessas aplicações de sucesso arrasador vieram da comunidade de pesquisas, ao passo que as duas últimas foram criadas por alguns jovens empreendedores.

No período de 1995 a 2001, a Internet realizou uma viagem vertiginosa nos mercados financeiros. Antes mesmo de se mostrarem lucrativas, centenas de novas empresas faziam suas ofertas públicas iniciais de ações e começavam a ser negociadas em bolsas de valores. Muitas empresas eram avaliadas em bilhões de dólares sem ter nenhum fluxo significativo de receita. As ações da Internet sofreram uma queda também vertiginosa em 2000-2001, e muitas novas empresas fecharam. Não obstante, várias outras surgiram como grandes vencedoras no mundo da Internet, entre elas Microsoft, Cisco, Yahoo, eBay, Google e Amazon.

1.7.5 O novo milênio

Nas duas primeiras décadas do século XXI, é possível que nenhuma tecnologia tenha transformado a sociedade mais do que a Internet combinada aos *smartphones* conectados à Internet. E a inovação na área de redes de computadores continua a passos largos. Há progressos em todas as frentes, incluindo distribuição de roteadores mais velozes e velocidades de transmissão mais altas nas redes de acesso e nos *backbones* da rede. Mas os seguintes desenvolvimentos merecem atenção especial:

- Desde o início do milênio, vimos a implementação agressiva do acesso residencial à Internet por banda larga – não apenas *modems* a cabo e DSL, mas também por fibra ótica, e agora também por acesso sem fio fixo 5G, conforme discutimos na Seção 1.2. Esse acesso à Internet de alta velocidade preparou a cena para uma série de aplicações de vídeo, incluindo a distribuição de vídeo gerado pelo usuário (p. ex., YouTube), *streaming* sob demanda de filmes e seriados de televisão (p. ex., Netflix) e videoconferência entre várias pessoas (p. ex., Skype, Facetime e Google Hangouts).
- A onipresença cada vez maior do acesso à Internet sem fio de alta velocidade não apenas está possibilitando permanecer constantemente conectado enquanto se desloca, mas também permite novas aplicações específicas à localização, como Yelp, Tinder e Waze. O número de dispositivos sem fio conectados ultrapassou o número de dispositivos com fio em 2011. Esse acesso sem fio em alta velocidade preparou a cena para o rápido surgimento de computadores portáteis (iPhones, Androids, iPads, etc.) que possuem acesso constante e livre à Internet.
- Redes sociais *online*, como Facebook, Instagram, Twitter e WeChat (incrivelmente popular na China), criaram grandes redes pessoais em cima da Internet. Muitas dessas redes sociais são amplamente utilizadas para troca de mensagens e compartilhamento de fotos. Muitos usuários hoje “vivem” principalmente dentro de uma ou mais redes sociais. Através de suas interfaces de programação de aplicações (APIs, do inglês *application programming interfaces*), as redes sociais *online* criam plataformas para novas aplicações em rede, incluindo pagamentos móveis e jogos distribuídos.
- Conforme discutimos na Seção 1.3.3, os provedores de serviços *online*, como Google e Microsoft, implementaram suas próprias amplas redes privadas, que não apenas conectam seus *datacenters* distribuídos em todo o planeta, mas são usadas para evitar a Internet ao máximo possível, emparelhando diretamente com ISPs de nível mais baixo. Como resultado, Google oferece resultados de busca e acesso a *e-mail* quase instantaneamente, como se seus *datacenters* estivessem rodando dentro do computador de cada usuário.
- Muitas empresas de comércio na Internet agora estão rodando suas aplicações na “nuvem” – como na EC2 da Amazon, na Azure da Microsoft ou na Alibaba Cloud. Diversas empresas e universidades também migraram suas aplicações da Internet (p. ex., *e-mail* e hospedagem de páginas Web) para a nuvem. Empresas de nuvem não apenas oferecem ambientes de computação e armazenamento escaláveis às aplicações, mas também lhes oferecem acesso implícito às suas redes privadas de alto desempenho.

1.8 RESUMO

Neste capítulo, abordamos uma quantidade imensa de assuntos. Examinamos as várias peças de *hardware* e *software* que compõem a Internet, em especial, e redes de computadores, em geral. Começamos pela periferia da rede, examinando sistemas finais e aplicações, além do serviço de transporte fornecido às aplicações que executam nos sistemas finais. Também vimos as tecnologias de camada de enlace e meio físico encontradas na rede de acesso. Em seguida, mergulhamos no interior da rede e chegamos ao seu núcleo, identificando comutação de pacotes e comutação de circuitos como as duas abordagens básicas do transporte de dados por uma rede de telecomunicações, expondo os pontos fortes e fracos de cada uma delas. Examinamos, então, as partes inferiores (do ponto de vista da arquitetura) da rede — as tecnologias de camada de enlace e os meios físicos comumente encontrados na rede de acesso. Estudamos também a estrutura da Internet global e aprendemos que ela é uma rede de redes. Vimos que a estrutura hierárquica da Internet, composta por ISPs de níveis mais altos e mais baixos, permitiu que ela se expandisse e incluisse milhares de redes.

Na segunda parte deste capítulo introdutório, abordamos diversos tópicos fundamentais da área de redes de computadores. Primeiro, examinamos as causas de atrasos e perdas de pacotes em uma rede de comutação de pacotes. Desenvolvemos modelos quantitativos simples de atrasos de transmissão, de propagação e de fila, bem como modelos de vazão; esses modelos de atrasos serão muito usados nos problemas propostos em todo o livro. Em seguida, examinamos camadas de protocolo e modelos de serviço, princípios fundamentais de arquitetura de redes aos quais voltaremos a nos referir neste livro. Analisamos também alguns dos ataques mais comuns na Internet. Terminamos nossa introdução sobre redes com uma breve história das redes de computadores. O primeiro capítulo constitui um minicurso sobre redes de computadores.

Portanto, percorremos de fato um extraordinário caminho neste primeiro capítulo! Se você estiver um pouco assustado, não se preocupe. Abordaremos todas essas ideias em detalhes nos capítulos seguintes (é uma promessa, e não uma ameaça!). Por enquanto, esperamos que, ao encerrar este capítulo, você tenha adquirido uma noção, ainda que incipiente, das partes que formam uma rede, um domínio ainda em desenvolvimento do vocabulário de redes (não se acanhe de voltar aqui para consulta) e um desejo cada vez maior de aprender mais sobre elas. Essa é a tarefa que nos espera no restante deste livro.

O guia deste livro

Antes de iniciarmos qualquer viagem, sempre é bom consultar um guia para nos familiarizarmos com as estradas principais e os desvios que encontraremos pela frente. O destino final da viagem que estamos prestes a empreender é um entendimento profundo do como, do quê e do porquê das redes de computadores. Nosso guia é a sequência de capítulos:

1. Redes de computadores e a Internet
2. Camada de aplicação
3. Camada de transporte
4. A camada de rede: plano de dados
5. A camada de rede: plano de controle
6. A camada de enlace e as LANs
7. Redes sem fio e móveis
8. Segurança em redes de computadores

Os Capítulos 2 a 6 são os cinco capítulos centrais deste livro. Note que eles estão organizados segundo as quatro camadas superiores da pilha de cinco camadas de protocolos da Internet. Note também que nossa jornada começará no topo da pilha, a saber, a camada de aplicação, e prosseguirá daí para baixo. O princípio racional que orienta essa jornada de cima para baixo é que, entendidas as aplicações, podemos compreender os serviços de rede

necessários para dar-lhes suporte. Então, poderemos examinar, um por um, os vários modos como esses serviços poderiam ser executados por uma arquitetura de rede. Assim, o estudo das aplicações logo no início dá motivação para o restante do livro.

A segunda metade — Capítulos 7 e 8 — aborda dois tópicos de extrema importância (e de certa maneira independentes) para as redes modernas. No Capítulo 7, examinamos as redes sem fio e móveis, incluindo LANs sem fio (incluindo WiFi e Bluetooth), redes celulares (incluindo 4G e 5G) e mobilidade. No Capítulo 8, sobre segurança em redes de computadores, analisamos, primeiro, os fundamentos da criptografia e da segurança de redes e, em seguida, de que modo a teoria básica está sendo aplicada a um amplo leque de contextos da Internet.

Exercícios de fixação e perguntas

Questões de revisão do Capítulo 1

SEÇÃO 1.1

- R1. Qual é a diferença entre um hospedeiro e um sistema final? Cite os tipos de sistemas finais. Um servidor Web é um sistema final?
- R2. A palavra *protocolo* é muito usada para descrever relações diplomáticas. Como a Wikipédia descreve um protocolo diplomático?
- R3. Por que os padrões são importantes para os protocolos?

SEÇÃO 1.2

- R4. Cite quatro tecnologias de acesso. Classifique cada uma delas nas categorias acesso residencial, acesso corporativo ou acesso móvel.
- R5. A taxa de transmissão HFC é dedicada ou é compartilhada entre usuários? É possível haver colisões na direção provedor-usuário de um canal HFC? Por quê?
- R6. Cite as tecnologias de acesso residencial disponíveis em sua cidade. Para cada tipo de acesso, apresente a taxa *downstream*, a taxa *upstream* e o preço mensal anunciados.
- R7. Qual é a taxa de transmissão de LANs Ethernet?
- R8. Cite alguns meios físicos utilizados para instalar a Ethernet.
- R9. HFC, DSL e FTTH são usados para acesso residencial. Para cada uma dessas tecnologias de acesso, cite uma faixa de taxas de transmissão e comente se a taxa de transmissão é compartilhada ou dedicada.
- R10. Descreva as tecnologias de acesso sem fio mais populares atualmente. Faça uma comparação entre elas.

SEÇÃO 1.3

- R11. Suponha que exista exatamente um nó de comutação de pacotes entre um computador de origem e um de destino. As taxas de transmissão entre a máquina de origem e o computador e entre este e a máquina de destino são R_1 e R_2 , respectivamente. Admitindo que um roteador use comutação de pacotes do tipo *armazena-e-reenvia*, qual é o atraso total fim a fim para enviar um pacote de comprimento L ? (Desconsidere formação de fila, atraso de propagação e atraso de processamento.)

R12. Qual é a vantagem de uma rede de comutação de circuitos em relação a uma de comutação de pacotes? Quais são as vantagens da TDM sobre a FDM em uma rede de comutação de circuitos?

R13. Suponha que usuários compartilhem um enlace de 2 Mbits/s e que cada usuário transmita continuamente a 1 Mbit/s, mas cada um deles transmite apenas 20% do tempo. (Veja a discussão sobre multiplexação estatística na Seção 1.3.)

a. Quando a comutação de circuitos é utilizada, quantos usuários podem ser admitidos?

b. Para o restante deste problema, suponha que seja utilizada a comutação de pacotes. Por que não haverá atraso de fila antes de um enlace se dois ou menos usuários transmitirem ao mesmo tempo? Por que haverá atraso de fila se três usuários transmitirem ao mesmo tempo?

c. Determine a probabilidade de um dado usuário estar transmitindo.

d. Suponha agora que haja três usuários. Determine a probabilidade de, a qualquer momento, os três usuários transmitirem simultaneamente. Determine a fração de tempo durante o qual a fila cresce.

R14. Por que dois ISPs no mesmo nível de hierarquia farão emparelhamento? Como um IXP consegue ter lucro?

R15. Alguns provedores de conteúdo criaram suas próprias redes. Descreva a rede da Google. O que motiva os provedores de conteúdo a criar essas redes?

SEÇÃO 1.4

R16. Considere o envio de um pacote de uma máquina de origem a uma de destino por uma rota fixa. Relacione os componentes do atraso que formam o atraso fim a fim. Quais deles são constantes e quais são variáveis?

R17. Visite a animação interativa “Transmission versus Propagation Delay” no site de apoio do livro. Entre as taxas, o atraso de propagação e os tamanhos de pacote disponíveis, determine uma combinação para a qual o emissor termine de transmitir antes que o primeiro *bit* do pacote chegue ao receptor. Ache outra combinação para a qual o primeiro *bit* do pacote alcança o receptor antes que o emissor termine de transmitir.

R18. Quanto tempo um pacote de 1.000 bytes leva para se propagar através de um enlace de 2.500 km de distância, com uma velocidade de propagação de $2,5 \times 10^8$ m/s e uma taxa de transmissão de 2 Mbits/s? Em geral, quanto tempo um pacote de comprimento L leva para se propagar através de um enlace de distância d , velocidade de propagação s , e taxa de transmissão de R bits/s? Esse atraso depende do comprimento do pacote? Depende da taxa de transmissão?

R19. Suponha que o provedor A queira enviar um arquivo grande para o provedor B. O percurso de A para B possui três enlaces, de taxas $R_1 = 500$ kbits/s, $R_2 = 2$ Mbits/s e $R_3 = 1$ Mbit/s.

a. Considerando que não haja nenhum outro tráfego na rede, qual é a vazão para a transferência de arquivo?

b. Suponha que o arquivo tenha 4 milhões de bytes. Dividindo o tamanho do arquivo pela vazão, quanto tempo levará a transferência para o provedor B?

c. Repita os itens “a” e “b”, mas agora com R_2 reduzido a 100 kbits/s.

R20. Suponha que o sistema final A queira enviar um arquivo grande para o sistema final B. Em um nível muito alto, descreva como o sistema A cria pacotes a partir do arquivo. Quando um desses arquivos chega ao roteador, quais informações no pacote

o roteador utiliza para determinar o enlace através do qual o pacote é encaminhado? Por que a comutação de pacotes na Internet é semelhante a dirigir de uma cidade para outra pedindo informações ao longo do caminho?

R21. Visite a animação interativa “Queuing and Loss” no site de apoio do livro. Qual é a taxa de transmissão máxima e a taxa de transmissão mínima? Com essas taxas, qual é a intensidade do tráfego? Execute a animação interativa com essas taxas e determine o tempo que leva a ocorrência de uma perda de pacote. Repita o procedimento mais uma vez e determine de novo o tempo de ocorrência para a perda de pacote. Os resultados são diferentes? Por que? Por que não?

SEÇÃO 1.5

R22. Cite cinco tarefas que uma camada pode executar. É possível que uma (ou mais) dessas tarefas seja(m) realizada(s) por duas (ou mais) camadas?

R23. Quais são as cinco camadas da pilha de protocolos da Internet? Quais as principais responsabilidades de cada uma dessas camadas?

R24. O que é uma mensagem de camada de aplicação? Um segmento de camada de transporte? Um datagrama de camada de rede? Um quadro de camada de enlace?

R25. Quais camadas da pilha do protocolo da Internet um roteador processa? Quais camadas um *switch* processa? Quais camadas um sistema final processa?

SEÇÃO 1.6

R26. O que é um *malware* autorreprodutivo?

R27. Descreva como pode ser criada uma *botnet* e como ela pode ser utilizada no ataque DDoS.

R28. Suponha que Alice e Bob estejam enviando pacotes um para o outro por uma rede de computadores e que Trudy se posicione na rede para poder capturar todos os pacotes enviados por Alice e enviar o que quiser para Bob; ela também consegue capturar todos os pacotes enviados por Bob e enviar o que quiser para Alice. Cite algumas atitudes maliciosas que Trudy pode fazer a partir de sua posição.

Problemas

P1. Projete e descreva um protocolo de nível de aplicação para ser usado entre um caixa eletrônico e o computador central de um banco. Esse protocolo deve permitir verificação do cartão e da senha de um usuário, consulta do saldo de sua conta (que é mantido no computador central) e saque de dinheiro (i.e., entrega de dinheiro ao usuário). As entidades do protocolo devem estar preparadas para resolver o caso comum em que não há dinheiro suficiente na conta para cobrir o saque. Especifique seu protocolo relacionando as mensagens trocadas e as ações realizadas pelo caixa automático ou pelo computador central do banco na transmissão e recepção de mensagens. Esquematize a operação de seu protocolo para o caso de um saque simples sem erros, usando um diagrama semelhante ao da Figura 1.2. Descreva explicitamente o que seu protocolo espera do serviço de transporte fim a fim.

P2. A Equação 1.1 contém uma fórmula para o atraso fim a fim do envio de um pacote de comprimento L por N enlaces com taxa de transmissão R . Generalize essa fórmula para enviar P desses pacotes de ponta a ponta pelos N enlaces.

P3. Considere uma aplicação que transmite dados a uma taxa constante (p. ex., a origem gera uma unidade de dados de N bits a cada k unidades de tempo, em que k é pequeno e fixo). Considere também que, quando essa aplicação começa, continuará em funcionamento por um período relativamente longo. Responda às seguintes perguntas, dando uma breve justificativa para suas respostas:

- O que seria mais apropriado para essa aplicação: uma rede de comutação de circuitos ou uma rede de comutação de pacotes? Por quê?
- Suponha que seja usada uma rede de comutação de pacotes e que o único tráfego venha de aplicações de dados a descrita anteriormente. Além disso, imagine que a soma das velocidades de dados da aplicação seja menor do que a capacidade de cada enlace. Será necessário algum tipo de controle de congestionamento? Por quê?

P4. Considere a rede de comutação de circuitos da Figura 1.13. Lembre-se de que há quatro circuitos em cada enlace. Rotule os quatro switches A, B, C e D, seguindo no sentido horário.

- Qual é o número máximo de conexões simultâneas que podem estar em curso a qualquer instante nessa rede?
- Suponha que todas as conexões sejam entre os switches A e C. Qual é o número máximo de conexões simultâneas que podem estar em curso?
- Suponha que queiramos fazer quatro conexões entre os switches A e C, e outras quatro conexões entre os switches B e D. Podemos rotear essas chamadas pelos quatro enlaces para acomodar todas as oito conexões?

P5. Considere novamente a analogia do comboio de carros da Seção 1.4. Admita uma velocidade de propagação de 100 km/h.

- Suponha que o comboio viaje 175 km, começando em frente ao primeiro dos postos de pedágio, passando por um segundo e terminando após um terceiro. Qual é o atraso so fim a fim?

b. Repita o item “a” admitindo agora que haja oito carros no comboio em vez de dez.

P6. Este problema elementar começa a explorar atrasos de propagação e de transmissão, dois conceitos centrais em redes de computadores. Considere dois hospedeiros, A e B, conectados por um único enlace de taxa R bits/s. Suponha que eles estejam separados por m metros e que a velocidade de propagação ao longo do enlace seja de s metros/segundo. O hospedeiro A tem de enviar um pacote de L bits ao hospedeiro B.

- Expresse o atraso de propagação, d_{prop} , em termos de m e s .
- Determine o tempo de transmissão do pacote, d_{trans} , em termos de L e R .
- Ignorando os atrasos de processamento e de fila, obtenha uma expressão para o atraso fim a fim.
- Suponha que o hospedeiro A comece a transmitir o pacote no instante $t = 0$. No instante $t = d_{trans}$, onde estará o último bit do pacote?

e. Imagine que d_{prop} seja maior do que d_{trans} . Onde estará o primeiro bit do pacote no instante $t = d_{trans}$?

f. Considere que d_{prop} seja menor do que d_{trans} . Onde estará o primeiro bit do pacote no instante $t = d_{trans}$?

g. Suponha que $s = 2,5 \times 10^8$ m/s, $L = 1.500$ bytes e $R = 10$ Mbits/s. Encontre a distância m de modo que d_{prop} seja igual a d_{trans} .

P7. Neste problema, consideramos o envio de voz em tempo real do hospedeiro A para o hospedeiro B por meio de uma rede de comutação de pacotes (VoIP). O hospedeiro A converte voz analógica para uma cadeia digital de bits de 64 kbits/s e, em seguida, agrupa os bits em pacotes de 56 bytes. Há apenas um enlace entre os hospedeiros A e B; sua taxa de transmissão é de 10 Mbits/s e seu atraso de propagação, de 10 ms.

Assim que o hospedeiro A monta um pacote, ele o envia ao hospedeiro B. Quando recebe um pacote completo, o hospedeiro B converte os bits do pacote em um sinal analógico. Quanto tempo decorre entre o momento em que um bit é criado (a partir do sinal analógico no hospedeiro A) e o momento em que ele é decodificado (como parte do sinal analógico no hospedeiro B)?

P8. Suponha que usuários compartilhem um enlace de 10 Mbits/s e que cada usuário precise de 200 kbits/s para transmitir, mas que transmita apenas durante 10% do tempo. (Veja a discussão sobre comutação de pacotes versus comutação de circuitos na Seção 1.3.)

- Quando a comutação de circuitos é utilizada, quantos usuários podem ser admitidos?
- Para o restante deste problema, suponha que seja usada a comutação de pacotes. Determine a probabilidade de que determinado usuário esteja transmitindo. Determine a probabilidade de que haja 120 usuários. Determine a probabilidade de que, a um tempo dado, exatamente n usuários estejam transmitindo simultaneamente. (Dica: Use a distribuição binomial.)
- Suponha que haja 120 usuários. Determine a probabilidade de que, a um tempo dado, exatamente n usuários estejam transmitindo simultaneamente. (Dica: Use a distribuição binomial.)
- Determine a probabilidade de haver 51 ou mais usuários transmitindo simultaneamente.

P9. Considere a discussão na Seção 1.3 sobre comutação de pacotes versus comutação de circuitos, na qual é dado um exemplo com um enlace de 1 Mbit/s. Quando em atividade, os usuários estão gerando dados a uma taxa de 100 kbits/s; mas a probabilidade de estarem em atividade, gerando dados, é de $p = 0,1$. Suponha que o enlace de 1 Mbit/s seja substituído por um de 1 Gbit/s.

- Qual é o número máximo de usuários, N , que pode ser suportado simultaneamente por comutação de pacotes?
- Agora considere comutação de circuitos e um número M de usuários. Elabore uma fórmula (em termos de p , M , N) para a probabilidade de que mais de N usuários estejam enviando dados.

P10. Considere um pacote de comprimento L que se inicia no sistema final A e percorre três enlaces até um sistema final de destino. Eles estão conectados por dois nós de comutação de pacotes. Suponha que d_i , s_i e R_i representem o comprimento, a velocidade de propagação e a taxa de transmissão do enlace i , sendo $i = 1, 2, 3$. O nó de comutação de pacote atrasa cada pacote por d_{proc} . Considerando que não haja nenhum atraso de fila, em relação a d_i , s_i e R_i ($i = 1, 2, 3$) e L , qual é o atraso fim a fim total para o pacote? Suponha agora que o pacote tenha 1.500 bytes, a velocidade de propagação de ambos os enlaces seja $2,5 \times 10^8$ m/s, as taxas de transmissão dos três enlaces sejam 2,5 Mbits/s, o atraso de processamento do comutador de pacotes seja de 3 ms, o comprimento do primeiro enlace seja 5.000 km, o do segundo seja 4.000 km e, do último, seja 1.000 km. Dados esses valores, qual é o atraso fim a fim?

P11. No problema anterior, suponha que $R_i = R$, $s_i = s$ e $d_{proc} = 0$. Suponha que o comutador de pacote não armazene e recieve pacotes, mas transmita imediatamente cada bit recebido antes de esperar o pacote chegar. Qual é o atraso fim a fim?

P12. Um comutador de pacotes recebe um pacote e determina o enlace de saída pelo qual deve ser enviado. Quando o pacote chega, outro já está sendo transmitido nesse enlace de saída e outros quatro já estão esperando para serem transmitidos. Os pacotes são transmitidos em ordem de chegada. Suponha que todos os pacotes tenham 1.500 bytes e que a taxa do enlace seja 2,5 Mbits/s. Qual é o atraso de fila para o pacote? De modo geral, qual é o atraso de fila quando todos os pacotes possuem comprimento L , a taxa de transmissão é R , x bits do pacote sendo transmitido já foram transmitidos e n pacotes já estão na fila?

P13. (a) Suponha que N pacotes cheguem simultaneamente ao enlace no qual não há pacotes sendo transmitidos e nem pacotes enfileirados. Cada pacote tem L de comprimento e é transmitido à taxa R . Qual é o atraso médio para os N pacotes?

- (b) Agora considere que N desses pacotes cheguem ao enlace a cada LN/R segundos. Qual é o atraso de fila médio de um pacote?

P14. Considere o atraso de fila em um *buffer* de roteador, sendo I a intensidade de tráfego; isto é, $I = \lambda a/R$. Suponha que o atraso de fila tome a forma de $IL/R(1-I)$ para $I < 1$.

- a. Deduza uma fórmula para o atraso total, isto é, para o atraso de fila mais o atraso de transmissão.

b. Faça um gráfico do atraso total como uma função de L/R .

P15. Sendo a a taxa de pacotes que chegam a um enlace em pacotes/s, e μ a taxa de transmissão de enlaces em pacotes/s, com base na fórmula do atraso total (i.e., o atraso de fila mais o atraso de transmissão) do problema anterior, deduza uma fórmula para o atraso total em relação a a e μ .

P16. Considere um *buffer* de roteador anterior a um enlace de saída. Neste problema, você usará a fórmula de Little, uma famosa fórmula da teoria das filas. Considere N o número médio de pacotes no *buffer* mais o pacote sendo transmitido, a a taxa de pacotes que chegam no enlace e d o atraso total médio (i.e., o atraso de fila mais o atraso de transmissão) sofrido pelo pacote. Dada a fórmula de Little $N = a \cdot d$, suponha que, na média, o *buffer* contenha 100 pacotes, o atraso de fila de pacote médio seja 20 ms, e a taxa de transmissão do enlace seja 100 pacotes/s. Utilizando tal fórmula, qual é a taxa média de chegada, considerando que não há perda de pacotes?

P17. (a) Generalize a Equação 1.2 na Seção 1.4.3 para taxas de processamento heterogêneas, taxas de transmissão e atrasos de propagação.

- (b) Repita o item (a), mas suponha também que haja um atraso de fila médio d_{fila} em cada nó.

P18. Execute o programa Traceroute para verificar a rota entre uma origem e um destino, no mesmo continente, para três horários diferentes do dia.

- a. Determine a média e o desvio-padrão dos atrasos de ida e volta para cada um dos três horários.

b. Determine o número de roteadores no caminho para cada um dos três. Os caminhos mudam em algum dos horários?

c. Tente identificar o número de redes de ISP pelas quais o pacote do Traceroute passa entre origem e destino. Roteadores com nomes semelhantes e/ou endereços IP semelhantes devem ser considerados parte do mesmo ISP. Em suas respostas, os maiores atrasos ocorrem nas interfaces de emparelhamento entre ISPs adjacentes?

- d. Faça o mesmo para uma origem e um destino em continentes diferentes. Compare os resultados dentro do mesmo continente com os resultados entre continentes diferentes.

P19. A Lei de Metcalfe afirma que o valor de uma rede de computadores é proporcional ao quadrado do número de usuários conectados do sistema. Considere n o número de usuários em uma rede de computadores. Supondo que cada usuário envie uma mensagem para cada um dos outros usuários, quantas mensagens serão enviadas? A sua resposta apoia a Lei de Metcalfe?

P20. Considere o exemplo de vazão correspondente à Figura 1.20(b). Agora imagine que haja M pares de cliente-servidor em vez de 10. R_s , R_c e R representam as taxas do enlace do servidor, enlaces do cliente e enlace da rede. Suponha que os outros enlaces possam capacidade abundante e que não haja outro tráfego na rede além daquele gerado pelos M pares cliente-servidor. Deduza uma expressão geral para a vazão em relação a R_s , R_c , R e M .

P21. Considere a Figura 1.19(b). Agora suponha que haja M percursos entre o servidor e o cliente. Dois percursos nunca compartilham qualquer enlace. O percurso k ($k = 1, \dots, M$) consiste em N enlaces com taxas de transmissão $R_1^k, R_2^k, \dots, R_N^k$. Se o servidor pode

usar somente um percurso para enviar dados ao cliente, qual é a vazão máxima que ele pode atingir? Se o servidor pode usar todos os M percursos para enviar dados, qual é a vazão máxima que ele pode atingir?

P22. Considere a Figura 1.19(b). Suponha que cada enlace entre o servidor e o cliente possua uma probabilidade de perda de pacote p , e que as probabilidades de perda de pacote para esses enlaces sejam independentes. Qual é a probabilidade de um pacote (enviado pelo servidor) ser recebido com sucesso pelo receptor? Se o pacote se perder no percurso do servidor para o cliente, então o servidor retransmitirá o pacote. Na média, quantas vezes o servidor retransmitirá o pacote para que o cliente o receba com sucesso?

P23. Considere a Figura 1.19(a). Suponha que o enlace de gargalo ao longo do percurso do servidor para o cliente seja o primeiro com a taxa R_1 bits/s. Imagine que enviemos um par de pacotes um após o outro do servidor para o cliente, e que não haja outro tráfego nesse percurso. Suponha também que cada pacote de tamanho L bits e os dois enlaces tenham o mesmo atraso de propagação d_{prop} .

- a. Qual é o tempo entre chegadas do pacote ao destino? Isto é, quanto tempo transcorre desde quando o último *bit* do primeiro pacote chega até quando o último *bit* do segundo pacote chega?

b. Agora suponha que o segundo enlace seja o de gargalo (i.e., $R_2 < R_1$). É possível que o segundo pacote entre na fila de entrada do segundo enlace? Explique. Agora imagine que o servidor envie o segundo pacote T segundos após enviar o primeiro. Qual deverá ser o tamanho de T para garantir que não haja uma fila antes do segundo enlace? Explique.

P24. Imagine que você queira enviar, com urgência, 50 *terabytes* de dados de Boston para Los Angeles. Você tem disponível um enlace dedicado de 100 Mbit/s para transferência de dados. Escolheria transmitir os dados por meio desse enlace ou usar um serviço de entrega em 24 horas? Explique.

P25. Suponha que dois hospedeiros, A e B, estejam separados por uma distância de 20 mil quilômetros e conectados por um enlace direto de $R = 5$ Mbit/s. Suponha que a velocidade de propagação pelo enlace seja de $2,5 \times 10^8$ m/s.

- a. Calcule o produto largura de banda-atraso $R \cdot d_{\text{prop}}$.
- b. Considere o envio de um arquivo de 800 mil *bits* do hospedeiro A para o hospedeiro B. Suponha que o arquivo seja enviado continuamente, como se fosse uma única grande mensagem. Qual é o número máximo de *bits* que estará no enlace a qualquer dado instante?

c. Interprete o produto largura de banda $\times$ atraso.

- d. Qual é o comprimento (em metros) de um *bit* no enlace? É maior do que o de um campo de futebol?

e. Derive uma expressão geral para o comprimento de um *bit* em termos da velocidade de propagação s , da velocidade de transmissão R e do comprimento do enlace m .

P26. Com referência ao Problema P24, suponha que possamos modificar R . Para qual valor de R o comprimento de um *bit* será o mesmo que o comprimento do enlace?

P27. Considere o Problema P24, mas agora com um enlace de $R = 500$ Mbit/s.

- a. Calcule o produto largura de banda-atraso $R \cdot d_{\text{prop}}$.

b. Considere o envio de um arquivo de 800 mil *bits* do hospedeiro A para o hospedeiro B. Suponha que o arquivo seja enviado continuamente, como se fosse uma única grande mensagem. Qual será o número máximo de *bits* que estará no enlace a qualquer dado instante?

- c. Qual é o comprimento (em metros) de um *bit* no enlace?

P28. Considere novamente o Problema P24.

- Quanto tempo demora para mandar o arquivo, admitindo que ele seja enviado continuamente?
- Suponha agora que o arquivo seja fragmentado em 20 pacotes, e que cada um contenha 40 mil *bits*. Imagine que cada pacote seja verificado pelo receptor e que o tempo de transmissão de uma verificação de pacote seja desprezível. Por fim, admita que o emissor não possa enviar um pacote até que o anterior tenha sido reconhecido. Quanto tempo demorará para enviar o arquivo?
- Compare os resultados de “a” e “b”.

P29. Suponha que haja um enlace de micro-ondas de 10 *Mbits/s* entre um satélite geostacionário e sua estação-base na Terra. A cada minuto, o satélite tira uma foto digital e a envia à estação-base. Considere uma velocidade de propagação de $2,4 \times 10^8$ m/s.

- Qual é o atraso de propagação do enlace?
- Qual é o produto largura de banda-atraso, $R \cdot d_{\text{prop}}$, das ondas transmitidas continuamente?
- Seja x o tamanho da foto. Qual é o valor mínimo de x para que o enlace de micro-ondas transmita continuamente?

P30. Considere a analogia da viagem aérea que utilizamos em nossa discussão sobre camadas na Seção 1.5, e o acréscimo de cabeçalhos a unidades de dados de protocolo enquanto passam pela pilha. Existe uma noção equivalente de acréscimo de informações de cabeçalho à movimentação de passageiros e suas malas pela pilha de protocolos da linha aérea?

P31. Em redes modernas de comutação de pacotes, inclusive a Internet, o hospedeiro de origem segmenta mensagens longas de camada de aplicação (p. ex., uma imagem ou um arquivo de música) em pacotes menores e os envia pela rede. O destinatário, então, monta novamente os pacotes restaurando a mensagem original. Denominamos esse processo *segmentação de mensagem*. A Figura 1.27 ilustra o transporte fim a fim de uma mensagem com e sem segmentação. Considere que uma mensagem de 10^6 *bits* de comprimento tenha de ser enviada da origem ao destino na Figura 1.27. Suponha que a velocidade de cada enlace da figura seja 5 *Mbits/s*. Ignore atrasos de propagação, de fila e de processamento.

- Considere o envio da mensagem da origem ao destino *sem* segmentação. Quanto tempo essa mensagem levará para ir do hospedeiro de origem até o primeiro comutador de pacotes? Tendo em mente que cada nó de comutação usa comutação de pacotes do tipo armazenar-e-reenvia, qual é o tempo total para levar a mensagem do hospedeiro de origem ao hospedeiro de destino?

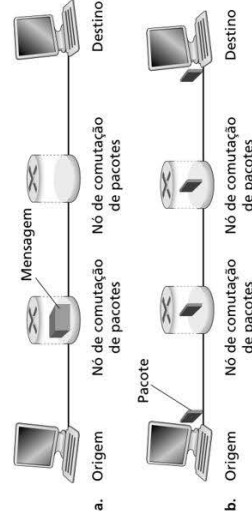


Figura 1.27 Transporte fim a fim de mensagem: (a) sem segmentação de mensagem; (b) com segmentação de mensagem.

- Agora suponha que a mensagem seja segmentada em 100 pacotes, cada um com 10.000 *bits* de comprimento. Quanto tempo demorará para o primeiro pacote ir do hospedeiro de origem até o primeiro nó de comutação? Quando o primeiro pacote está sendo enviado do primeiro ao segundo nó de comutação, o segundo pacote está sendo enviado da máquina de origem ao primeiro nó de comutação. Em que instante o segundo pacote terá sido completamente recebido no primeiro nó de comutação?
- Quanto tempo demorará para movimentar o arquivo do hospedeiro de origem até o hospedeiro de destino quando é usada segmentação de mensagem? Compare este resultado com sua resposta no item “a” e comente.
- Além de reduzir o atraso, quais são as razões para usar a segmentação de mensagem?
- Discuta as desvantagens da segmentação de mensagem.

P32. Experimente a animação interativa “Message Segmentation”, apresentada no *site* deste livro. Os atrasos na animação interativa correspondem aos atrasos obtidos no problema anterior? Como os atrasos de propagação no enlace afetam o atraso total fim a fim na comutação de pacotes (com segmentação de mensagem) e na comutação de mensagens?

P33. Considere o envio de um arquivo grande de F *bits* do hospedeiro A para o hospedeiro B. Há três enlaces (e dois nós de comutação) entre A e B, e os enlaces não estão congestionados (i.e., não há atrasos de fila). O hospedeiro A fragmenta o arquivo em segmentos de S *bits* cada e adiciona 80 *bits* de cabeçalho a cada segmento, formando pacotes de $L = 80 + S$ *bits*. Cada enlace tem uma taxa de transmissão de R *bits/s*. Qual é o valor de S que minimiza o atraso para levar o arquivo de A para B? Desconsidere o atraso de propagação.

P34. O Skype oferece um serviço que lhe permite fazer uma ligação telefônica de um PC para um telefone comum. Isso significa que a chamada de voz precisa passar pela Internet e por uma rede telefônica. Discuta como isso poderia ser feito.

Wireshark Lab

“Conte-me e eu esquecerei. Mostre-me e eu lembrarei. Envolve-me e eu entenderei.”
Provérbio chinês

A compreensão de protocolos de rede pode ser muito mais profunda se os vimos em ação e interagirmos com eles — observando a sequência de mensagens trocadas entre duas entidades de protocolo, pesquisando detalhes de sua operação, fazendo que eles executem determinadas ações e observando tais ações e suas consequências. Isso pode ser feito em cenários simulados ou em um ambiente real de rede, tal como a Internet. As animações interativas apresentadas (em inglês) no *site* deste livro adotam a primeira abordagem. Nos Wireshark labs, adotaremos a última. Você executará aplicações de rede em vários cenários utilizando seu computador no escritório, em casa ou em um laboratório e observará também os protocolos de rede interagindo e trocando mensagens com entidades de protocolo que estão executando em outros lugares da Internet. Assim, você e seu computador serão partes integrantes desses laboratórios ao vivo. Você observará — e aprenderá — fazendo.

A ferramenta básica para observar as mensagens trocadas entre entidades de protocolos em execução é denominada **analisador de pacotes** (*packet sniffer*). Como o nome sugere, um analisador de pacotes copia passivamente mensagens enviadas e recebidas por seu computador, também exibe o conteúdo dos vários campos de protocolo das mensagens que captura.

Uma tela do analisador de pacotes Wireshark é mostrada na Figura 1.28. O Wireshark é um analisador de pacotes gratuito que funciona em computadores com sistemas operacionais Windows, Linux/Unix e Mac. Por todo o livro, você encontrará Wireshark labs que o habilitarão a explorar vários dos protocolos estudados em cada capítulo. Neste primeiro Wireshark lab, você obterá e instalará uma cópia do programa, acessará um *site* e examinará as mensagens de protocolo trocadas entre seu navegador e o servidor Web.

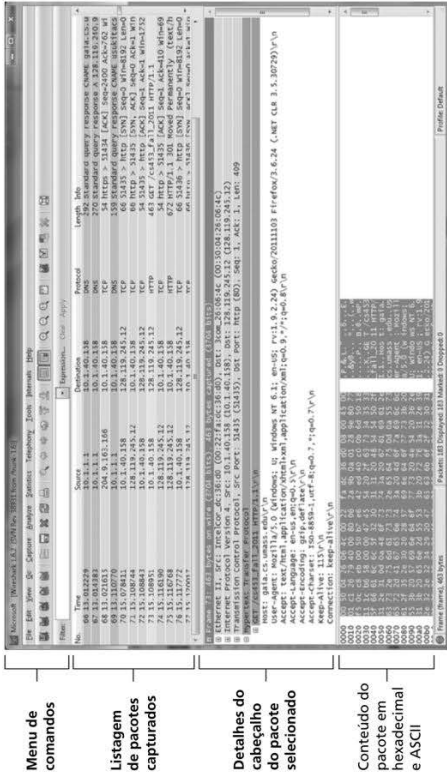


Figura 1.28 Uma amostra de tela do programa Wireshark (amostra de tela do Wireshark reimpressa com permissão da Wireshark Foundation).