



IPv6+智能骨干网络技术研究

刘建国¹, 渠灏¹, 黄红明², 朱智丹²

(1. 中央广播电视总台, 北京 100038; 2. 华为技术有限公司, 北京 100073)

摘要: 中央广播电视总台(以下简称总台)骨干网是连接总台、数据中心、国内外各总站/记者站和全国“百城千屏”的关键基础设施, IPv6+是下一代IP网络的主要技术路线。主要分析了总台骨干网的现状和面临的挑战, 介绍了IPv6+核心的原理、优势和技术特点, 研究和探讨了IPv6+技术如何适应广播电视业务的发展需求, 最后提出了基于IPv6+构建和部署下一代广播电视骨干网的构想和建议, 为未来广播电视台骨干网建设提供了简化协议、加速业务开通的新思路。

关键词: 骨干网; IPv6+; 网络切片; iFIT; APN6; 广联接; 超宽带

中图分类号: G220.7

文献标志码: A

doi: 10.11959/j.issn.1000-0801.2023115

Research on IPv6+ intelligent backbone network technology

LIU Jianguo¹, QU Hao¹, HUANG Hongming², ZHU Zhidan²

1. China Media Group, Beijing 100038, China

2. Huawei Technologies Co., Ltd., Beijing 100073, China

Abstract: The China Media Group backbone network is a key infrastructure that connects the headquarters, data centers, hubs/journalists at home and abroad, and “hundreds cities and thousands of screens” across the country, and IPv6+ is the main technical route of the next-generation IP network. The present situation and challenges of the backbone network of the headquarters were mainly analyzed, the principle, advantages and technical characteristics of the core of IPv6+ were introduced, and how IPv6+ technology adapts to the development of broadcast and television services was discussed. Finally, the conception and suggestion of constructing and deploying the next generation broadcast TV backbone network based on IPv6+ were put forward. It provided a new idea to simplify the protocol and speed up the service opening for the future construction of backbone network of radio and television stations.

Key words: backbone network, IPv6+, network slicing, iFIT, APN6, wide connection, ultra-broadband

0 引言

中央广播电视总台(以下简称总台)骨干网作为总台与数据中心、数据中心与各省总站、总

台与国内外总站/记者站之间安全、便捷的业务数据传输和交换网络, 其架构随着电视台业务的发展特点不断演进, 包括4K/8K百城千屏分发、新闻云采编、摄像机远程制作、视频会议、



办公 OA 等业务。现有骨干网采用 Hub-Spoke 架构, 基于网际协议 (internet protocol, IP) /多协议标签交换 (multi-protocol label switching, MPLS) 协议互联, “只能实现网络可达, 缺乏业务流量的质量保障和可视运维”, 存在网络无法根据业务实现精细保障、网络维护成本高、故障定位难、网络收敛慢等问题, 难以满足视/音频和云业务承载要求。

IPv6/IPv6+作为下一代网络的主要技术路线, 已经得到全球客户的普遍共识, 全球已有超过 100 家大型客户部署 IPv6, 涵盖运营商、金融、政府和能源等领域。当前, 我国在 IPv6+技术体系创新上处于全球领先地位。IPv6+以 SRv6、网络切片、iFIT、BIERv6、APN6 等为代表的协议创新, 在广联接、超宽带、自动化、确定性、低时延和安全 6 个维度全面提升 IP 网络能力。本文主要分析 IPv6+核心技术 (如 SRv6 Policy、TiLFA 保护、BIERv6) 的原理和优势, 研究如何适应广播电视业务的发展需求, 拟构建基于 IPv6+的下一代广播电视骨干网。

1 现状和挑战

1.1 总台骨干网现状

总台骨干网采用“北京城域+全国广域”的建设方案。在互联网线路上, 城域内租用运营商裸光纤资源, 加上自建波分设备进行互连。广域网租用多个运营商跨省点对点专线进行互连。

城域网主要承载的业务类型包括通用和专用两大类业务。其中, 通用业务包括日常办公和业务管理数据; 专用业务包括全程非线性采、编、播、存、管业务, 节目采集、上载、编目等, 基于网络的非线性高清节目编辑制作, 节目数据备份、迁移、存储等, 节目数据的播出, 节目数据的管理, 以及台内 IPTV 网络等。

城域网用于连接北京分支站址, 以 IP/MPLS VPN 为基础协议搭建统一承载网络服务平台, 支撑总台总控、播出、新闻网络制播、综合节目制

作、包装、审片、新闻和综合演播室、媒资系统、媒体数据交换、企业数据总线、综合信息门户、节目生产、办公等 90 多个业务系统。核心 P/PE 节点采用双平面、双机的方式增加可靠性保护。隧道技术采用传统的 MPLS LDP 技术, 业务承载采用 L3VPN (单播) 和 NG-MVPN (多播) 技术。选用大容量万兆以太网骨干, 适应总台节目生产所需带宽要求。

总台广域网用于连接央视总台、各省总站、新闻采集中心、区域制作中心、海外中心站、海外分台。采用星形连接拓扑, 各分支与总部之间租用运营商专线组网, 国内分支租用点到点专线和互联网专线互备, 海外分支通过二层 VPN 专线。广域网承载了融媒体演播室 (SDI)、新闻云外延采编、视频会议、总控控制、资源回传等业务。广域网通过 IP/MPLS VPN 与北京总部互连, 通过 VLAN/VPN 做隔离区分, 在设备入口做 QoS 限速。各记者站至总台专线带宽无法复用。

1.2 现网面临的挑战

(1) 网络协议多达 8 种, 协议间相互调用, 难定位故障, 运维困难

总台骨干网的单播业务, 采用 MPLS 网络运行时涉及的控制面协议多, 涉及外部边界网关协议 (EBGP)、内部边界网关协议 (IBGP)、内部网关协议 (IGP)、标签分发协议 (LDP) 4 种基于 IPv4 发展出来的控制协议信令交互, 不同协议间有时要考虑交互影响 (如需要考虑 IGP 和 LDP 之间的同步), 运维复杂, 需要考虑的问题多, 对维护人员技术要求高。

多播业务采用 MPLS 多播 VPN 方案, 涉及 BGP、NG-MVPN、PIM、IGMP 4 种协议传递多播信息, 一方面协议复杂、多播树消耗设备资源大, 每个中间节点都要维护多播状态; 另一方面可靠性弱, 网络规模扩大后, 一旦出现故障, 多播收敛时间变为分钟级, 影响业务体验。

(2) 广域网星形架构建网成本高, 新技术难发挥技术价值

当前骨干网采用星形拓扑架构, 各记者站、海外中心站和分台均通过点到点线路连接回总台。此种网络架构下, 网络链路资源无法进行带宽统计复用能力, 新技术带来的技术红利无法体现 (如不降低可靠性的同时提高网络带宽利用率)。

随着全国 31 个省 (区市) 的总站相继启动建设, 网络规模将会成倍增加。参考全球运营商、金融、政府、能源等行业的广域网建设理念, 需要采用层次化网络架构的新顶层设计。

(3) IP 网络仅解决连通性, 缺少对业务的差异化保障和运维能力

当前城域网、广域网仅能保证单播和多播业务的连通性, 只做到 IP 可达, 缺少网络确定性、有差异、高质量、易运维的连接能力。虽开启了 QoS 功能, 但仍无法对某类业务或某条重保业务流进行端到端的质量保障和监控。如在发生故障的时候, 由于当前网络只能感知每台设备端口的信息, 无法准确定位到底是哪一跳设备、哪一条链路出了问题, 更无法定位是网络的问题还是应用系统的问题, 从而无法帮助系统运维人员快速定位并解决问题。

2 IPv6+技术特点

2.1 基于广电行业的 SRv6 特点

(1) SRv6 对于广电来说更容易部署

当一个传统 IP/MPLS 设备网络要部署新的业务时, 如果部署 SRv6, 可以先将一些基本设备升级到 SRv6 设备, 再用 overlay 的 SRv6 隧道部署新的业务, 另外一个设备只需要部署 IPv6。如果部署 SR-MPLS, 应该把所有的设备升级为 SR 设备, 再部署新的业务。所以 SRv6 对传统网络更友好。

(2) SRv6 支持广电大网络扩展

SRv6 采用 IPv6 地址作为网段 ID, 因此它具有天然的原生 IP 属性, 与可路由地址相连的标签

不同, SRv6 的 SID 是可路由的地址, 在 IPv6 地址可达的情况下, 可以基于 SRv6 发布业务。因此, 快速建立新业务是非常方便的, 不需要分段配置业务, 只需要升级业务创建点, 就可以拥有 VPN 业务体验。

对于 MPLS 标签, 不能做任何聚合行为, 如果想建立与某人的连接, 必须有其具体路线。这带来了巨大的复杂性, 很难对网络进行大规模扩展。而 SRv6 可以总结路由。在安全策略允许的情况下, 可以使用默认路由达到用户想要的任何点。即使用 SRv6, 可以获得几乎无限的扩展能力。

(3) SRv6 有 3 级网络编程能力

第一层: 灵活的 SID 列表来实现路径规划。这也是 SR 的 MPLS。

第二层: 每个段都有 3 个字段——定位器、函数、args, 可以实现不同的业务, 这是 SR 的 MPLS 没有的能力。

第三层: 通过选项 TLV, SRv6 的 SID 可以携带 AppID、UserID 和 meta, 实现应用程序的编程, 这也是 SR 的 MPLS 没有的能力。

2.2 基于 SRv6 Policy 的总台骨干网络

(1) SRv6 Policy 的路径可编程能力

SRv6 Policy 利用 SR 的源路由机制, 通过在 IPv6 报文中插入路由扩展头 (segment routing header, SRH), 在其中压入一个显式的 IPv6 地址栈, 中间节点转发时通过匹配目的地址和偏移地址栈的操作来完成逐跳转发。SRv6 Policy 是一种将用户意图 (SLA、业务链) 映射为网络实现的机制, 是 IPv6+ 时代实现流量工程的重要手段。SRv6 技术可以给用户带来 3 点价值: 简化网络配置、提供高可靠保护能力、提供转发路径流量编排与调优能力。

- 头节点: 标识 SRv6 Policy 的头节点, 通常为央视云 PE 节点、各省总站 PE 节点, 在广域网的两个边界侧。头节点将不同类型流量导入不同的 SRv6 Policy 中, 在头节点的 SRH 字段内, 会以每跳路由器的地址



作为标签,插入逐跳顺序。

- 颜色/优先级: 标识指定头节点和目的节点的不同 SRv6 Policy, 可与一系列业务属性相关联, 如新闻云、“百城千屏”等实时性强的业务为高优先级, 标记为 color1; 办公 OA、文件传输等业务为中优先级, 标记为 color2。color 不同的流量进入不同的 SRv6 Policy 隧道, 即可实现业务和网络的关联。
- 目的节点: 标识 SRv6 Policy 的目的地址。总台广域网连接拓扑如图 1 所示。结合全台骨干网未来承载业务和拓扑规划, 以北京总台云中心节点的云出口 PE 为起点, 深圳总站 PE 为终点, 设计不同业务的转发路径。如对于新闻云采编业务, 选择从北京到深圳最短路径转发 (如图 1 浅灰色 SRv6 Policy 隧道标识), 保证业务质量。而对于其他办公 OA 或下载类业务 (如图 1 深灰色 SRv6 Policy 隧道标识), 可以选择非最短但流量轻载路径转发。这种能力即 SRv6 Policy 的路径可编程能力, 可以保证全网的平均带宽利用率提升, 降低网络带宽建设成本, 对网络运维和运营带来巨大价值。

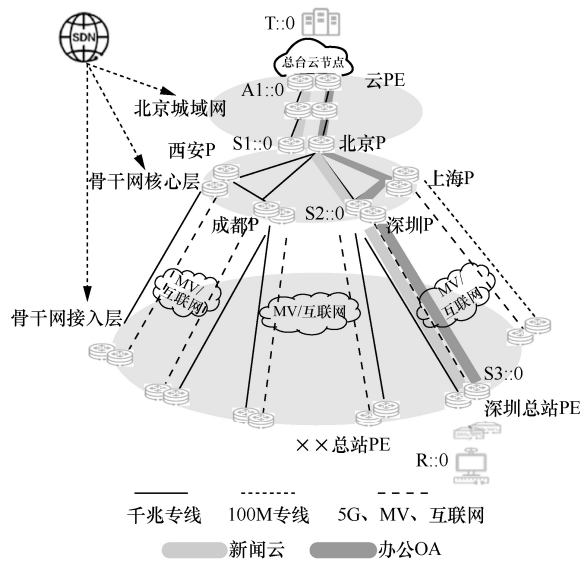


图1 总台广域网连接拓扑

(2) 与传统网络、多类型专线网络融合互通

网络改造升级是逐步的, 当前总台广域网络均具备 IPv6 能力, 但有部分设备不支持 SRv6 能力, 此时如果希望部署 SRv6, 需要能穿越 IPv6 native 区域。虚拟链路就是用于该场景, 虚拟链路是指在两个 SRv6 节点间建立虚拟的链路跨越 IPv6 native 网络。链路连接示意图 (一) 如图 2 所示, 该场景共分为 3 类。

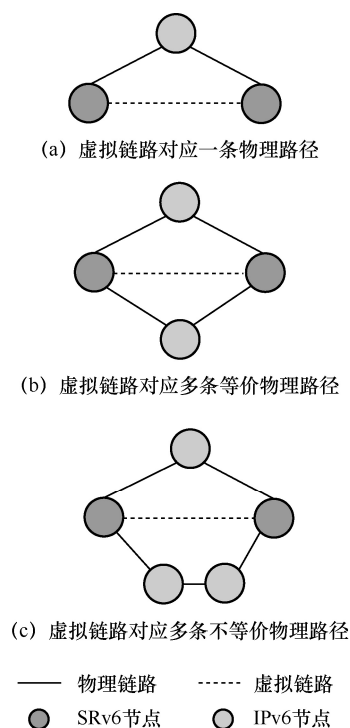


图2 链路连接示意图 (一)

上述场景是租用运营商点到点专线将网络设备进行连接, 由于广域网线路成本高昂, 有时还要租用其他类型专线, 如 MPLS VPN 专线或互联网专线等, 这些专线是运营商网络上的 3 层专线, 结合 SRv6 场景, 实际上是需要虚链路穿越运营商的 IP 网络。链路连接示意图 (二) 如图 3 所示, 共有 3 类场景。

对于以上与传统网络、多类型专线网络融合互通, 通过虚拟链路的创建, 连接支持 SRv6 的节点, 实现 SRv6 Policy 穿越场景的业务创建。同时, 还可以在虚拟链路上配置静态 IPv6 时延、SRv6

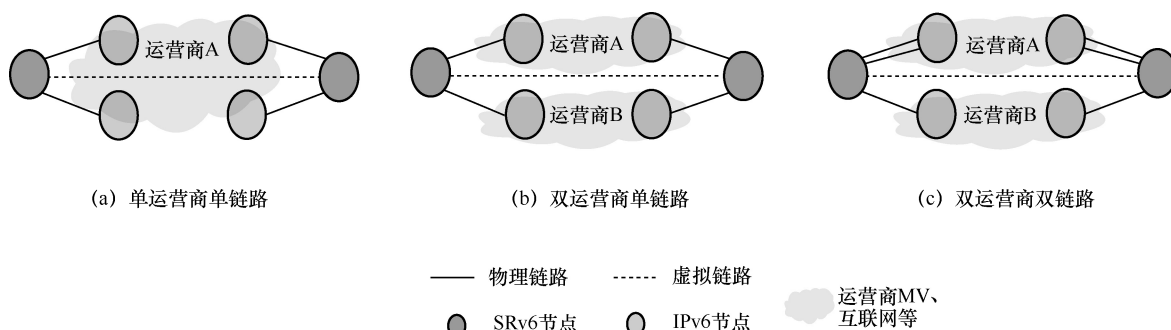


图3 链路连接示意图(二)

Policy Metric 度量值、亲和属性等相关的属性值，完成 SRv6 Policy 路径计算和流量转发。

(3) SRv6 Policy 路径优化能力

SRv6 Policy 通过编排 SRH 中的显式 IPv6 地址栈，在网络中形成多条业务隧道，实现了对业务路径编排和控制，实现最优性价比、关键业务保障、流量负载分担等众多流量调优效果。

目前的 SRv6 路径调优参数约束与功能见表 1，约束参数包括 5 类。

表 1 SRv6 路径调优参数约束与功能

约束参数	功能
优先级	指定不同隧道之间的优先顺序，高优先级隧道可以抢占低优先级隧道的带宽资源
带宽	带宽包括链路带宽与采集的实时流量带宽，使能实时流量带宽后，SDN 控制器会使用实时带宽进行路径计算，忽略隧道配置带宽
跳数限制	根据跳数约束算路
时延门限	算路选择路径时延在门限范围内的路径
显式路径	支持严格和松散模式，严格模式是指每一跳路径都指定，松散模式是指指定其中几个关键节点，其余路径按照网络配置自选

调优的策略支持时延最小、开销(cost)最小、链路时间可用度最优。实际部署中，建议如下。

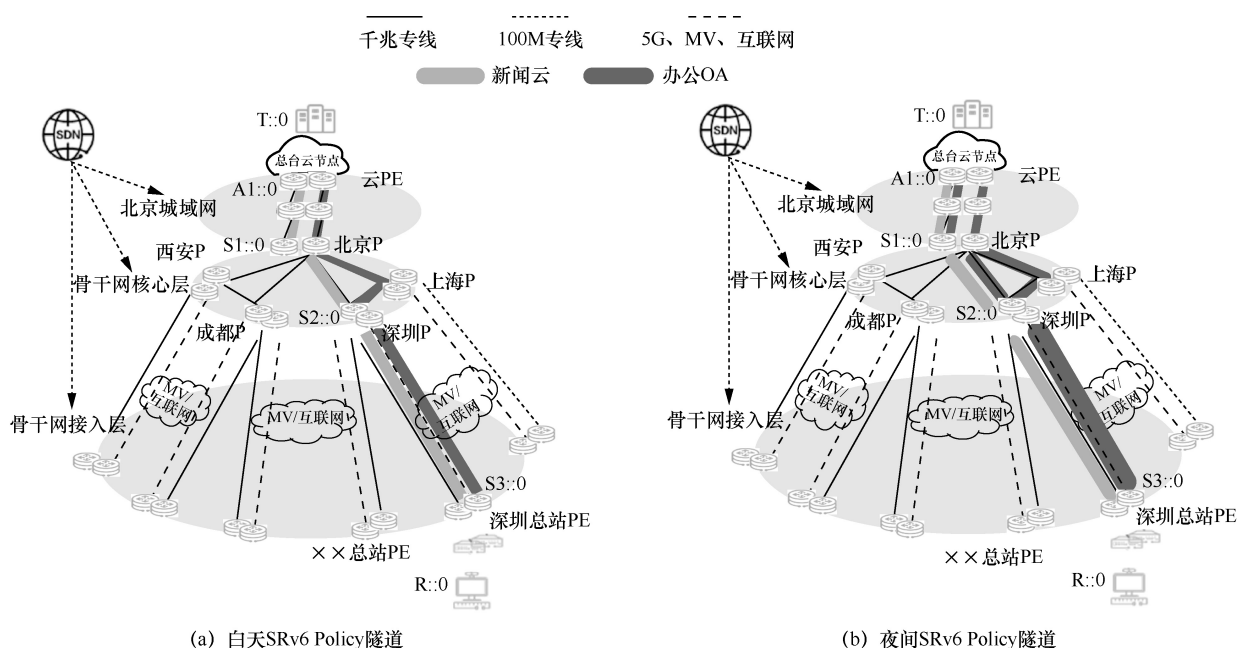
- 对时延比较敏感的实时性业务（如视频直播、视频会议等）可以选择时延最小算法。
- 普通业务建议使用开销最小算法。
- 大带宽、实时性要求不高的业务（如数据同步、文件同步等）可以采用链路时间可用度最优算法。链路带宽依据业务

特点灵活调整如图 4 所示，夜间没有“百城千屏”业务时，由 SDN 控制器将浅灰色隧道带宽变小，深灰色隧道也分流走到浅灰色隧道所在的物理链路上，将流量在不同物理链路上实现 1:1 的等价负载分担，或者 N:1 的非等价负载分担，节省扩容链路的成本。当白天“百城千屏”业务启动时，重新调整回隧道策略，保证视频直播业务带宽。

2.3 面向视音频业务，提供高可靠网络保护能力

交互式视音频业务（如 VoIP）对网络丢包非常敏感，通常只能容忍数十毫秒的网络丢包，而网络中链路或设备发生故障时，网络自动恢复/网络收敛时间通常为数百毫秒甚至达到秒级，造成视频卡顿、花屏、重新连接等情况。

为最大限度地减少流量损失，网络设备预先配置一条备份路径，当故障发生的时候，由邻近故障点的路由器（本地修复节点（point of local repair, PLR））快速切换到备份路径，从而最大限度地减少网络故障丢包，提升收敛性能，这种机制称为快速重路由（fast reroute, FRR），它是当前 IP 网络中收敛性能最佳的解决方案。FRR 存在 3 种算法：LFA（loop-free alternate）、RLFA（remote LFA）以及 TiLFA（topology-independent LFA）。由于 3 种算法的计算方式不同，LFA 和 RLFA 在某些特定的拓扑场景无法计算得到备份路径（LFA 和 RLFA 分别能解决

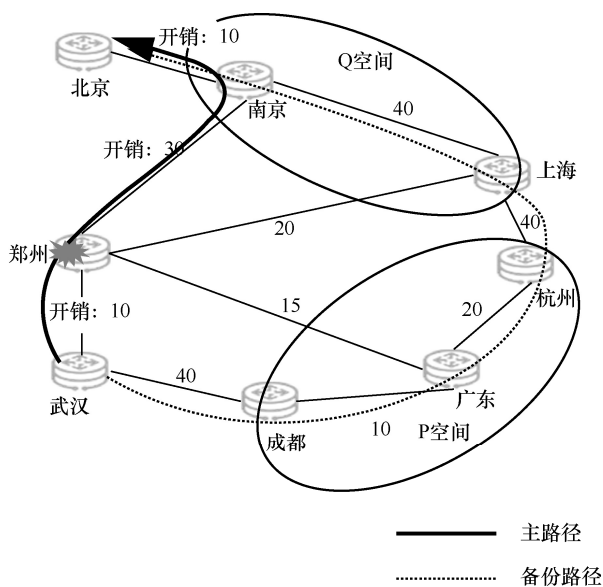


75%、96%场景), 而 TiLFA 采用拓扑无关的无环替换路径快速重路由机制, 可以实现 100%的拓扑和海量业务场景下任意故障 50 ms 快速倒换, 保障用户业务体验。

TiLFA 在计算备份路由时, 首先视保护链路(保护节点)不存在, 计算得到一条最短路径, 然后在最短路径上寻找扩展 P 空间(以保护链路源端为根节点建立 SPF 树, 所有从根节点不经过保护链路可达的节点集合称为 P 空间)和 Q 空间(以保护链路末端为根节点建立反向 SPF 树, 所有从根节点不经过保护链路可达的节点集合称为 Q 空间)。武汉至北京 TiLFA 保护路径如图 5 所示, 从武汉到北京的流量, 根据开销(cost 值)计算得到最优路径(见图 5 中加粗线型), 同时通过 TiLFA 计算得到郑州节点故障时的保护路径(见图 5 中虚线线型)。

具体过程如下。

步骤 1 从武汉转发到北京的流量, 正常的转发路径为“武汉→郑州→南京→北京”。假设故障点为郑州节点, 现在要计算郑州节点的保护路径。



步骤 2 计算 P 空间。以源节点(武汉)为根, 计算到其他节点的最短路径转发(SPF)树, 删掉被保护路径及其下挂的分支和节点, 剩下可达的节点(不包含根节点)满足表达式: “cost(邻居节点→扩展 P) < cost(邻居节点→保护节点) + cost(保护节点→扩展 P)” 的集合就是 P 空间。即: 从源节点到 P 节点不经过被保护链路(被保护节点),

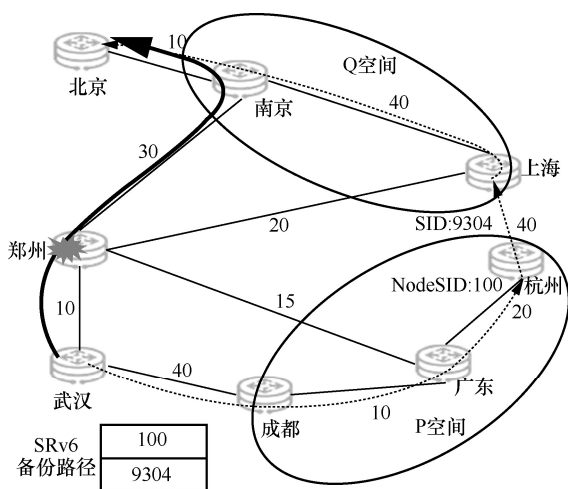
这样的节点集合 P 空间为{成都, 广东, 杭州}。

步骤 3 计算 Q 空间: 以目的节点 (北京) 为根, 其他节点到根的 R-SPF 树, 删掉被保护路径及其下挂的分支和节点, 剩下可达的节点 (不包含根节点) 满足表达式 “ $\text{Cost}(Q \rightarrow \text{目的节点}) < \text{Cost}(Q \rightarrow \text{保护节点}) + \text{Cost}(\text{保护节点} \rightarrow \text{目的节点})$ ” 的集合就是 Q 空间。即: 从 Q 节点到目的节点不经过被保护链路 (被保护节点), 这样的节点集合 Q 空间为{南京、上海}。

步骤 4 计算备份出接口和 Repair List。Repair List 是一个约束路径, 用来指示如何到达 Q 节点, Repair List 由 “P 节点标签+P 到 Q 路径上的邻接标签” 组成。

处理方法是建一条到最远的 P 节点的 SRv6 SID, 而远端 P 节点到最近的 Q 节点的隧道段采用压入邻接标签的方式, 进入 Q 节点之后, 后面就可以使用基于最短路径到达目的节点的方式了, 因为后面都是 Q 节点, 到达目的节点的最短路径都不回头。

注意, P 空间和 Q 空间不存在交集的情况下, TiLFA 能够计算出备份路径, 但是 LFA 和 RLFA 算法失效。所以 TiLFA 能够达到 100% 的场景覆盖。武汉至北京 TiLFA 保护 SRv6 转发路径如图 6 所示。



NodeSID:100→要求从武汉节点转发到杭州节点
SID:9304→要求从杭州特定链路转发到上海

图 6 武汉至北京 TiLFA 保护 SRv6 转发路径

步骤 5 如图 6 所示, 先找到最远的 P 节点杭州, 压入杭州的 NodeSID:100 到标签栈, 再找到最近的 Q 节点上海, 这时要压入上海的 link SID, 保证报文可以转发到 Q 节点上海, 最后按照 SPF 算法进行最短路径转发。备份转发路径为: 武汉→成都→广州→杭州→上海→南京→北京。

2.4 基于 BIERv6 的极简多播设计

位索引显式复制 (bit index explicit replication, BIER) 技术, 是由 RFC 8279 定义的一种新型的多播报文转发架构。此架构对多播转发域中的接收者节点编码, 并将所有接收者节点以比特串的形式表达。在多播源给接收端发送数据时, 源节点将标识了接收端节点的比特串封装在多播报文中, 多播报文根据比特串所代表的接收端节点复制多播报文, 达到多播转发的目的。BIERv6 是 BIER 的子集, 是指在 IPv6 网络中采用 BIER 技术架构完成多播复制转发的方案。

相比与传统多播协议, BIERv6 协议与单播转发机制相融合, 采用基础的 IGP 和单播路由转发机制, 结合代表多播接收端的比特串完成多播报文复制转发, 具有如下优点。

- 简化协议: 只需单播路由协议 IGP/BGP 扩展, 跟随单播路由转发; 无须创建多播分发树 (如现网的 PIM 多播树、NG-MVPN P2MP 隧道), 无须额外的协议 (如 mLDP/TE P2MP), 无须共享树和源树切换等复杂处理。
- 简化运维: 网络的中间节点只需一份转发表, 无须持续维护基于流的多播表项, 极大地节约了资源, 多播组应用规模不受限, 多播节目部署发生变化时中间节点不感知, 无须重建和撤销等复杂运维, 简化了大网部署。
- 可靠性高: 由于不需要维护大量基于流的多播树表项, 设备存储表项少; 当网络出



现一个故障点，设备只刷新一条表项，故障收敛快（百毫秒级），同时结合双源与双流选收等保护手段，可实现多播业务 E2E 的高可靠性。

- 业务体验好：用户观看视频加入多播，通过 BGP 直接向头端 PE 经一跳直接加入，不需要经过中间节点逐跳加入，时延小，用户体验好。

传统 PIM 多播与 BIER 多播的能力对比见表 2。

BIERv6 转发主要涉及几个关键角色：位转发路由器（bit-forwarding router, BFR）、位转发入口路由器（bit-forwarding ingress router, BFIR）、位转发出口路由器（bit forwarding egress router, BFER）。

- BFR：支持 BIERv6 转发的多播报文的转发设备。
- BFIR：BFIR 是一种特殊的 BFR，处于多播流进入多播转发域的边缘。
- BFER：BFER 也是一种特殊的 BFR，处于多播流离开多播转发域的边缘。
- BIER 域：所有 BFR 组成的网络形成 BIER 域（BIER domain）。

- BFR-ID：每台 BIERv6 域内的转发设备必须配置一个唯一的编号 BFR-ID，BFR-ID 是取值为 1~65 535 的整数。

每个 BFIR/BFER 需要分配唯一 bit 标识 BFR-ID 与对应的 BFR-Prefix 地址（设备特定的 IPv6 地址，标识为 BEIR 节点）。每个叶子节点通过 IGP 在 BIERv6 域中洪泛 BFR-ID 标识和 BFR-Prefix 映射关系。所有的 BFR 设备基于路由计算到 BFR-Prefix 的最佳路径，同时基于 bit 位置生成到每一个邻居节点（BFR-NBR）的转发掩码。

举例如下，边缘节点分别分配唯一的 BFR-ID，比如 A 节点 BFR-ID 为 0:1 000。所有 BIER 域中路由器完成 IGP 收敛后，BFR 设备完成邻居等信息学习，并在每台设备上形成一个包含 BFER 的 BFR-ID、BFR-Prefix 和 BFR-NBR 邻居关系的比特索引转发表（BIFT），如图 7 所示。

BIERv6 可承载公网和 MVPN 多播业务，可同时承载 IPv4 与 IPv6 业务。BIERv6 转发区域仅在承载网络层运行，需要运行 IPv6 协议。多播源与接受者仍采用传统的 IGMP/PIM 协议进行多播加入，每个频道对应一个多播组 G。

表 2 传统 PIM 多播与 BIER 多播的能力对比

对比项	PIM	BIERv6	比较分析
控制机制	需要为每条流建立多播分发树，需要选择 RP、RPF 处理，共享树和源树切换	不需要建立多播分发树，不需要建立专门的 P2MP 隧道	传统多播技术为每条流建立多播分发树，当多播业务发展增多时，会导致树的数量急剧增大，资源占用多，运维困难
转发机制	PIM 多播树复制转发	根据单播路由转发	与单播转发一致便于业务统一运维管理
可扩展性	需要中间节点保留多播流和树的状态，不利于扩展	中间节点只需要一份转发表即可	传统多播协议需每个转发点保留大量多播状态信息，导致：网络变化时，收敛慢；新用户加入时，时延大；用户体验差
可靠性	PIM MoFRR	双根 1+1；单播 FRR	BIERv6 保护机制更简单，可靠性更高
用户体验	新用户加入需逐跳加入多播树，响应慢	新用户加入只需在 BFIR 和 BFER 间发请求即可	传统多播协议新用户加入请求要逐跳加入多播转发树中，导致用户观看初始延时大、体验差
设备能力要求	所有节点需保留多播树和多播流状态	报文携带 BitString，BitString 占用部分业务带宽，芯片支持复制转发	传统多播协议对设备表项资源要求高；BIER 转发对设备芯片复制能力要求高
跨域网络能力	分域分段部署 PIM，建多播分发树，E2E 设备要求支持 PIM	BIERv6 可通过 IP 可达来跨越不支持节点	BIERv6 由于其依据 IP 路由可达性，可以跨越不支持节点，便于穿越不支持设备或网络

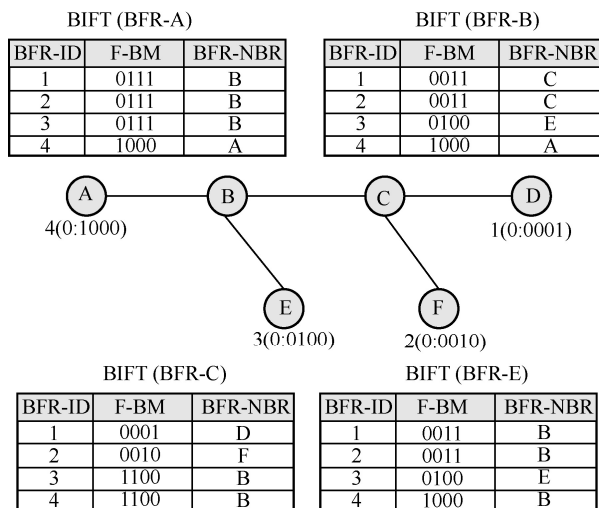


图7 比特索引转发表

BIERv6 MVPN 如图8所示, 视频内容源输出两个频道, 每个频道对应一个多播组。其中频道1发往重庆、湖南与福建, 频道2发往福建与广西。因为福建的两个大屏需要接收不同的频道内容, 因此福建对应的 BFR 节点需要同时申请加入与接收频道1与频道2的内容, 收到视频内容后, 再根据

出接口选择向不同的大屏转发对应的频道内容。

具体的 BIERv6 业务转发流程如下。

步骤1 BIERv6 区域内路由器设备使能 BIER 功能, 配置 BFR-ID、BFR-Prefix 等, IGP 计算生成 BIFT 转发表。

步骤2 多播接收端 (以重庆大屏为例) 发起 IGMP 请求, 申请加入多播组 G1。接入交换机收到申请后向重庆 BFR 节点发起 PIM 请求, 申请加入多播组(S,G1)。

步骤3 重庆 BFR 节点收到 PIM 请求后, 通过 BGP MVPN 信令向北京多播源传递加入消息, 携带 BFR-ID 00001, BFIR 北京节点依此建立 (vrf,S,G1)对应的 BitString 信息。因为重庆节点、湖南节点、福建节点均申请加入 G1 多播组, 因此生成 bitstring 00011, 同理, 针对 G2 多播组生成 bitstring 01100。

步骤4 BFIR 节点通过用户报文中(S,G), 查找(vrf,S,G)获取 bitstring, 根据 BFIR 设备的 BFIR

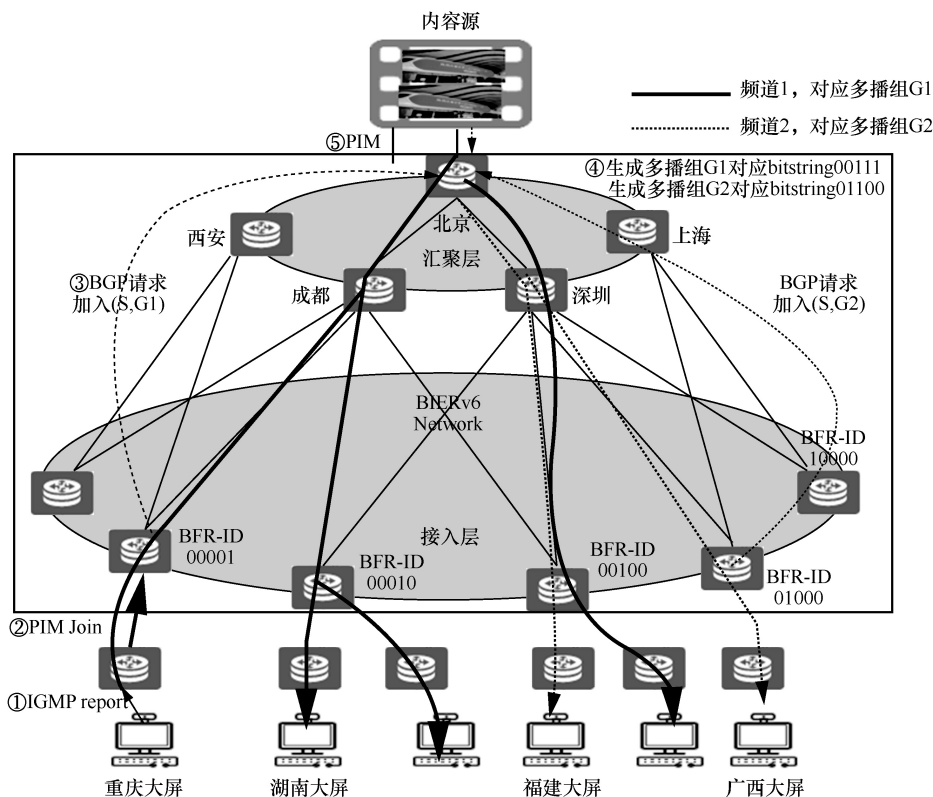


图8 BIERv6 MVPN



表复制多播报文, 进入 BIER 域。北京 BFIR 节点收到多播报文目的地址是 G1 时, 封装 bitstring 00011, 开始进行 BIER 转发。

步骤 5 中间 BFR 节点不关注多播报文内容与多播组地址, 仅根据 BIFT 转发表参考 bitstring 内容进行查表转发。因为中间设备无多播组状态, 所以可以实现大规格与快速手链。

步骤 6 BFEF 节点(重庆)根据报文携带的信息, 查找(vrf,S,G), 完成多播在出口路由器的转发。

3 IPV6+在总台骨干网中的部署

3.1 基于IPv6+打造总台IT基础设施底座, 扩大能力, 提升服务

总台 IPv6+骨干网 IT 基础设施底座如图 9 所示。

在媒体融合向纵深发展的背景下, 新一代基础网络平台将实现总台本部与国内外总站/记者站之间安全、便捷的业务数据传输和交换系统, 实现总台国内外各办公区域网络高速又安全的互联互通。采用 IPv6/IPv6+、5G、边缘计算等新一代网络通信技术, 大幅提升总台基础网络在数据处理、智能感知、安全防护等方面的能力, 同时降低运维成本, 实现总台基础信息网络跨越式发展。

- 百城千屏: 在广域骨干网总台和全国总站之间开展 4K/8K 超高清多播分发时, 通过 BIERv6 多播承载技术高效复制多播直播视频流, 简化网络配置并提升网络可靠性。
- 内容回传: 在总台/总站/记者站之间进行内容回传时, 采用 SRv6 Policy 实现这些节点之间的快速连接、负载分担、业务隔离以及端到端服务质量保证。尤其是基于网络切片严格保障视/音频关键业务质量。
- 总台北京各办公区确定性体验: 在北京城域骨干网, 通过网络切片保障, 支持切片规划、部署、业务灵活映射切片、业务 SLA 可视、引入机器学习等 AI 技术实现对资源的最优分配, 提升网络利用率。
- 智能运维: 采用统一 SDN 控制器统管总台城域网和广域网, 呈现网络可视化, 提供物理拓扑、逻辑拓扑、设备性能、业务等直观展示视图, 为网络管理提供依据。随流检测技术实时感知业务质量和呈现。

3.2 基于IPv6+技术体系, 构筑总台“5+X+Y”智能广域网络

总台“5+X+Y”智能广域网络架构如图 10 所示。

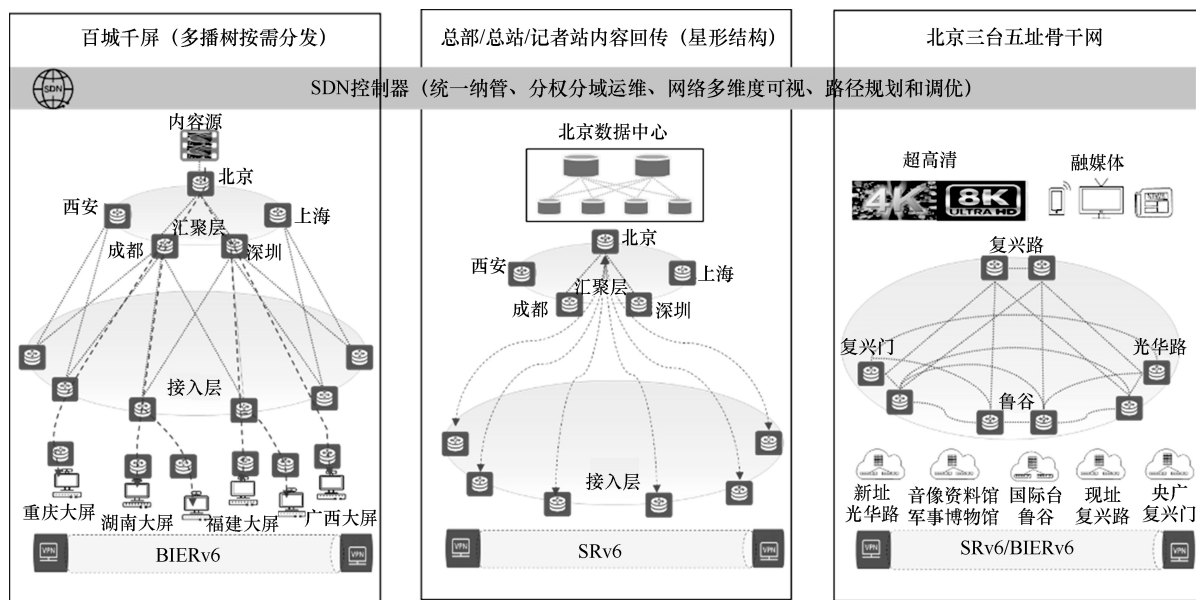


图 9 总台 IPv6+骨干网 IT 基础设施底座

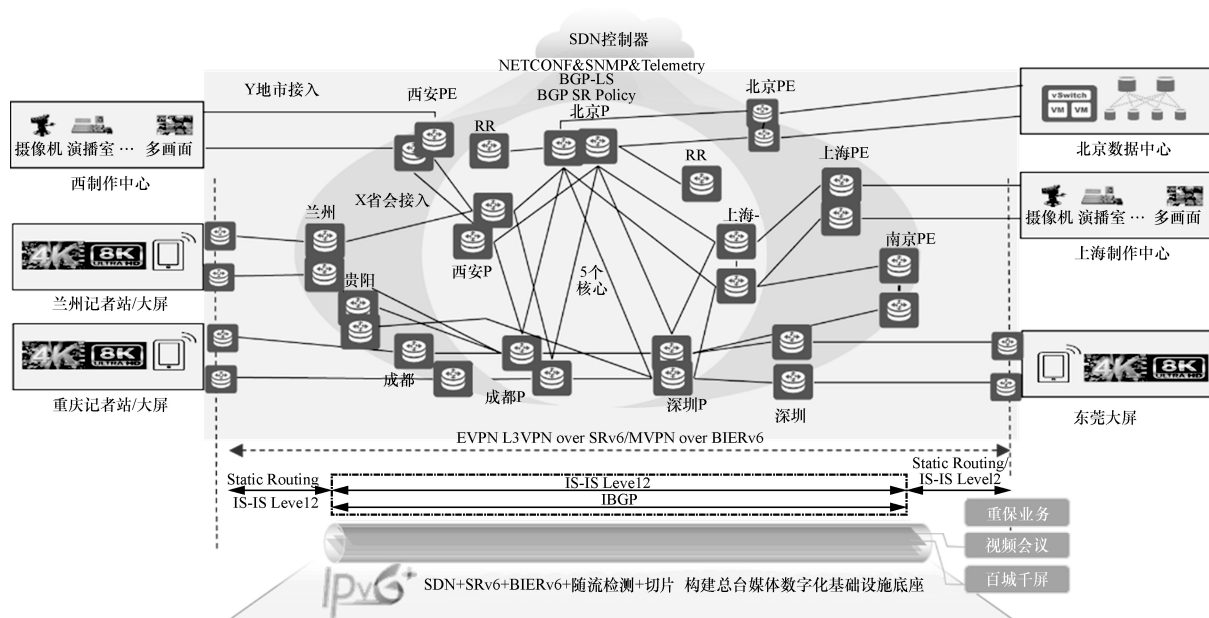


图10 总台“5+X+Y”智能广域网络架构

(1) 网络架构

全国设置5个核心节点，每个节点部署两台核心IP转发设备，每个省设置一个汇聚节点，部署两台汇聚IP转发设备，省内各记者站/大屏作为接入节点，部署1~2台接入IP转发设备。

北京市区内采用裸光纤连接，核心层之间、核心层与省会节点之间通过运营商专线链路互联，省会节点与省内记者站/大屏之间通过Internet专线连接。

北京市区内、核心层之间、核心层与省会节点之间部署IS-ISv6 (IS-ISv6版本，即intermediate system-to-intermediate system，中间系统到中间系统)，省会节点与省内记者站/大屏之间通过IPv6静态路由传递IPv6基础路由信息。

记者站与省会节点部署IBGP，省会节点与核心层节点部署IBGP，省会节点作为二级RR，核心层选择2台设备作为一级核心RR，全网通过IBGP传递IPv4/IPv6业务路由信息。

(2) IPv6+单播业务网络

全网部署SRv6隧道，IPv4、IPv6业务全部承载在SRv6隧道上。总台、省会节点、各记者站/大屏之间基于SRv6隧道实现业务全联通。不同业

务系统（如生产、办公、财务等）可划分独立VPN进行业务隔离。

关键业务采用网络切片、SRv6 Policy和iFIT保障，保障业务质量，全网采用NCE-IP管理系统实现统一管理，便于业务部署与调优。支持基于五元组通过SRv6 Policy将不同业务负载分担到网络出接口链路。

(3) IPv6+多播业务网络

全网部署BIERv6隧道，IPv4、IPv6多播业务承载在BIERv6隧道上。总台、5个核心、其他省会节点以及各记者站分级进行多播复制。

BIERv6控制及转发面报文穿越点对点专线、MPLS VPN专线/专网、互联网专线进行交互，打造全球首个超高清视音频IPv6+多播示范应用，承载“百城千屏”业务。

3.3 部署SRv6之后的优势

(1) 简化网络协议

SRv6是对网络架构的一次革命性简化，面对传统架构中复杂的MPLS协议，SRv6实现了网络协议的去MPLS化，所有现网业务无须进行MPLS配置，仅保留了IGP与BGP两种基础网络协议。



(2) 业务开通快速

基于 IPv6 路由的可达性, SRv6 简化了跨域的难度。一方面, 对于非关键业务保持 SRv6 BE 承载, 业务仅两端部署, 中间节点仅需支持 IPv6, 从而最大限度地提升网络运维效率。另一方面, 通过 SRv6 TE Policy 承载 B2B 专线业务, 实现了专线业务自动发放, 将开通时间缩短到 1 天。

(3) 可持续演进

基于 SRH 的可编程性, SRv6 实现了网络与业务的解耦, 在未来演进中也无须新增其他协议, 而且支持网络可持续演进。

4 结束语

面向总台全面启动“十四五”转型, 以“超清化、移动化、智能化”高质量发展为目标, 以“5G+4K/8K+AI”为战略方向, 需要建设新一代基础网络平台, 实现总台本部与国内外总站/记者站之间安全、便捷的业务数据传输和交换, 实现总台国内外各办公区域网络高速又安全的互联互通。需要结合 IPv6、5G 网络切片、边缘计算等新一代网络通信技术, 大幅提升总台基础网络在数据处理、智能感知、安全防护等方面的能力, 同时降低运维成本, 保证全国“百城千屏”4K/8K 视频广域网多播传送、“新闻云远程编辑”、“摄像机远程制作”等特色视音频、媒体控制类业务, 实现总台基础信息网络作为广电行业 IPv6+骨干网络标杆项目跨越式发展。

参考文献:

- [1] 徐进. 中央广播电视总台“十四五”技术发展规划[J]. 现代电视技术, 2021(6): 22-27.
XU J. Technical development plan of the “14th Five-Year Plan” of China Media Group[J]. Advanced Television Engineering, 2021(6): 22-27.
- [2] OMDIA. Advancing the digital economy with IPv6 and IPv6+[R]. 2021.
- [3] 推进 IPv6 规模部署专家委员会. “IPv6+”技术创新愿景与展望白皮书[R]. 2021.
Expert Committee on Promoting IPv6 Scale Deployment. “IPv6+” technical innovation vision and outlook white paper[R]. 2021.
- [4] 李振斌. SRv6 网络编程: 开启 IP 网络新时代[M]. 北京: 人民邮电出版社, 2020.
LI Z B. SRv6 network programming: ushering in a new era of IP networks[M]. Beijing: Posts & Telecom Press, 2020.
- [5] 克拉伦斯·菲尔斯菲尔斯. Segment Routing 详解[M]. 北京: 人民邮电出版社, 2020.
FILSFILS C. Segment routing details[M]. Beijing: Posts & Telecom Press, 2020.
- [6] IETF. Multicast using bit index explicit replication (BIER): RFC 8279[R]. 2017.
- [7] IETF. Multicast VPN using bit index explicit replication (BIER): RFC 8556[R]. 2018.
- [8] IETF. Segment routing architecture: RFC 8402[R]. 2018.
- [9] IETF. Internet protocol, version 6 (IPv6) specification: RFC 8200[R]. 2017.
- [10] IETF. IPv6 segment routing header (SRH): draft-ietf-6man-segment-routing-header[R]. 2019.
- [11] IETF. SRv6 network programming: draft-filsfils-spring-SRv6-network-programming[R]. 2019.
- [12] IETF. Segment routing policy architecture: draft-ietf-spring-segment-routing-policy[R]. 2019.
- [13] IETF. BGP MPLS-based ethernet VPN: RFC 7432[R]. 2015.
- [14] IETF. Multicast using bit index explicit replication (BIER): RFC 8279[R]. 2017.
- [15] IETF. Encapsulation for bit index explicit replication (BIER) in MPLS and non-MPLS networks: RFC 8279[R]. 2018.

[作者简介]



刘建国(1974—), 男, 中央广播电视总台高级工程师, 主要研究方向为 IP 网络新技术和演进方向在广播电视行业的应用和实践。

渠灏(1982—), 男, 中央广播电视总台工程师, 主要研究方向为广播电视网络系统建设及自动化运维。

黄红明(1977—), 男, 现就职于华为技术有限公司, 主要研究方向为通信行业 IP 网络新技术和演进方向。

朱智丹(1981—), 男, 现就职于华为技术有限公司, 主要研究方向为数据通信产品 SRv6、多播等路由协议的创新规划。