

Глава 1

Интернет-протокол IP и глобальная система адресации

Следует определить различия между именами, адресами и маршрутами. Имя определяет то, что мы пытаемся найти. Адрес указывает, где это находится. Маршрут показывает, как туда попасть.

RFC 760¹, первая спецификация интернет-протокола IPv4, 1980 г.

Три дня рождения Интернета

Ранним утром 29 октября 1969 года произошло историческое событие — рождение Интернета. В тот момент мало кто осознавал значимость этого события. Чарли Кляйн (Charley Kline) на своем терминале в Калифорнийском университете в Лос-Анджелесе (UCLA) набрал слово LOGIN, чтобы отправить эту команду компьютеру в Стэнфордском исследовательском институте (SRI), за которым ожидал коллега Чарли, Билл Дювал (Bill Duvall). Первый символ 'L' проделал путь в 500 км, был принят компьютером Билла и послан обратно, появившись на терминале Чарли. За ним последовал символ 'O'. На символе 'G' система сломалась, но была полностью восстановлена часом позже. Так был рожден Интернет.

Чарли и Билл были молодыми программистами, сотрудниками двух крупнейших американских научных центров. Рождению Интернета предшествовало десятилетие научных исследований, а свой вклад внесли десятки, если не сотни людей, разработавших базовые концепции архитектуры Интернета.

¹ RFC 760: DoD Standard Internet Protocol,
URL: <https://www.rfc-editor.org/rfc/rfc760>

Еще в начале 1960-х гг. ряд исследователей, многие из которых в дальнейшем участвовали в проекте ARPANET, увидели огромные перспективы в способности компьютеров обмениваться друг с другом данными. В 1965 году было установлено тестовое соединение между компьютерами Массачусетского института технологии и Университета Южной Калифорнии — использовалась традиционная телефонная технология синхронной коммутации каналов. Стало очевидно, что такая коммутация не позволяет эффективно использовать канал связи, но именно в ходе этого эксперимента начал обретать очертания «эмбрион» будущего Интернета.

Слово «Интернет» вошло в обиход в середине 1970-х, а до того Сеть называлась ARPANET. По сравнению с телефонными сетями, основанными на коммутации каналов, в ARPANET было решено использовать технологию коммутации пакетов, или дейтаграмм — данных ограниченного объема, заключенных в «конверты» с указанием источника и получателя. Поскольку каждый пакет обрабатывался независимо, сети не требовалось хранить информацию о соединениях между оконечными компьютерами и потоках данных между ними. Этот подход позволил существенно упростить архитектуру сети и повысить ее надежность. Узел сети мог выйти из строя — но его функцию немедленно брал на себя другой, рабочий узел. Кроме того, асинхронная пакетная передача больше соответствовала характеру работы многозадачных операционных систем. Так, ОС Unix позволяла разделять ресурсы между несколькими задачами одновременно — процессор занимался и обработкой команд с многочисленных терминалов, и вычислением крупных массивов данных. Telnet (удаленный доступ в режиме терминала) и электронная почта (e-mail) появились в ARPANET в 1972 г., а ftp (обмен файлами) — годом позже. Первое время для обмена данными между компьютерами, или хостами, использовался протокол NCP (Network Control Protocol), предтеча сегодняшнего TCP/IP.

Функциональность протокола NCP ограничивалась тем, что это, по существу, был транспортный протокол. Он не был хорошо приспособлен для работы с разнообразными технологиями — например, цифровой радио- и спутниковой связью. Более того, он предназначался для работы только с одной сетью — ARPANET, а значит, не был способен осуществлять адресацию в других сетях и среди подключенных к ним компьютеров.

В это же время Роберт Кан (Robert Kahn), сотрудник агентства передовых исследовательских проектов DARPA, работал над концепцией открытой сетевой архитектуры. В рамках этой концепции независимые сети, различные по своей архитектуре и используемым технологиям, должны были свободно обмениваться данными. Требовалась лишь единая межсетевая модель для «прозрачного» обмена данными между компьютерами в различных сетях. Особенностью концепции Кана было то, что он рассматривал и функциональность беспроводных сетей пакетной коммутации. Поскольку радиосигнал может подвергаться искажениям до полной потери (например, при перемещении в туннеле), протокол должен был обеспечить надежную передачу данных независимо от качества сети.

В 1973 году Кан начал разработку протокола, который позволил бы передавать данные между хостами, используя любую коммуникационную технологию. Кан пригласил в свой проект Винтона Серфа (Vinton Cerf), в то время сотрудника Стэнфордского университета. Серф обладал необходимым опытом: ранее он участвовал в создании протокола NCP и разрабатывал сетевые интерфейсы к различным операционным системам. Благодаря совместным усилиям Кана и Серфа концепция нового протокола была представлена уже в сентябре 1973 года а годом позже, в декабре 1974-го, Серф вместе со своими аспирантами Йогеном Далалем (Yogen Dalal) и Карлом Саншайном (Carl Sunshine) опубликовал первую полную спецификацию протокола TCP. Аббревиатура означала Transmission Control Program, а сам протокол объединял в себе функции сегодняшних протоколов TCP и IP. Новейшая спецификация была зафиксирована в серии документов Request for Comments (RFC) под номером RFC 675.²

Интересно, что изначально архитектура протокола TCP предполагала использование 4 бит для адресации сети (допуская тем самым существование 16 сетей, из которых шесть были уже назначены: ARPANET, UCL, CYCLADES, NPL, CADG, EPSS), а также 16 бит для адресации хостов в сети (или «процессов TCP»). При этом заголовок пакета также содержал поле длины сетевого адреса, тем самым обеспечивая расширение адресного пространства при необходимости до 64 бит. Однако в следующей версии протокола TCP, опубликованной в 1977 году, уже использовались только адреса фиксированной длины. А ведь изначальная структура TCP предлагала более элегантный способ борьбы с нехваткой адресного пространства, нежели создание протокола IPv6, несовместимого с IPv4!

В том же 1977 году была проведена первая серьезная демонстрация работы «Интер-Нета»: три сети, использующие различные сетевые технологии — ARPANET, SATNET и сеть пакетного радио, — успешно обменивались данными по протоколу TCP. Так Интернет был рожден во второй раз.

Двумя месяцами позже в то время аспирант Калифорнийского университета в Лос-Анджелесе (UCLA) Джон Постел (Jon Postel) опубликовал статью, где предложил новый архитектурный взгляд на TCP — протокол, состоящий из двух компонентов. Первый компонент, который в последующей спецификации TCP получит название IP (Internet Protocol), отвечал только за передачу пакетов между узлами сети и маршрутизацию. Второй же компонент, TCP, обеспечивал сквозной поток данных между оконечными устройствами, контроль ошибок и повторную передачу потерянных данных.

В 1978 году Постел публикует четвертую версию протоколов IP и TCP. И наконец в 1980 году публикуется документ RFC 760³, содержащий спецификацию IPv4 и принципы архитектуры Интернета, какими мы их знаем сегодня.

² RFC 675: Specification of Internet Transmission Control Program, URL: <https://www.rfc-editor.org/rfc/rfc675>

³ RFC 760: DoD Standard Internet Protocol, URL: <https://www.rfc-editor.org/rfc/rfc760>

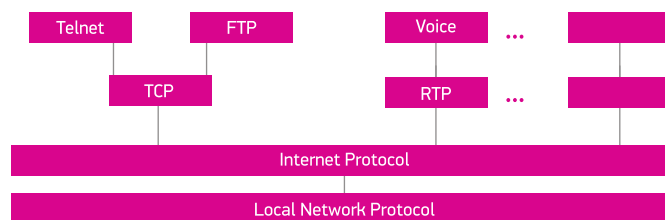


Рис. 1. Взаимодействие протоколов стека TCP/IP, определенное в спецификации RFC 760.

В рамках этой архитектуры IP отвечает за адресацию и фрагментацию пакетов при передаче от одного узла сети к другому. Адреса позволяют узлу принять решение, какому следующему узлу направить данные, а с помощью фрагментации данные можно передавать между сетями с различными допустимыми размерами пакета. Этим, собственно, функции протокола IP и ограничиваются.

В своей работе IP опирается на протоколы нижнего уровня, используемые в локальной сети, и транспортные протоколы, например, TCP. Сам же интернет-протокол не обеспечивает надежную передачу. В спецификации RFC 760 указано, что в протоколе IP «отсутствуют подтверждения, как сквозные, так и межузловые. Отсутствует контроль ошибок, за исключением контрольной суммы заголовка. Отсутствует функция повторной передачи. Отсутствует управление потоком данных».

Переход к семейству протоколов TCP/IP

Если вы думаете, что с переходом к протоколу IPv6 Интернет впервые переживает столь фундаментальное изменение базового протокола, то это не так. Сеть ARPANET конца 1970-х по-прежнему использовала протокол NCP, ограничивая возможности прозрачного обмена данными с другими сетями, например, с сетями пакетного радио или спутниковыми сетями. А в этом и заключалась основа концепции Интернета.

В ноябре 1981 года Джон Постел опубликовал план перехода ARPANET от протокола NCP к протоколам TCP/IP⁴. Учитывая, что новый протокол уже прошел успешное тестирование в различных конфигурациях, на переход отводился один год.

Тогдашнюю ARPANET невозможно сравнить с сегодняшней сетью Интернет — и по размеру, и по зависимости общества и экономики от ее функционирования, и по степени контроля и координации. Тем не менее, переход занял целый год

⁴ RFC 801: NCP/TCP transition plan, URL: <https://www.rfc-editor.org/rfc/rfc801>

и потребовал определенного количества напоминаний и увещеваний со стороны Джона Постела. Кроме того, на сутки был отключен протокол NCP по всему ARPANET'у, так что только узлы, поддерживавшие протокол TCP/IP, могли обмениваться данными.

Окончательный переход на TCP/IP произошел, как и было запланировано, 1 января 1983 года. Так Интернет был рожден в третий раз, теперь с протоколом IPv4.

Эволюция системы адресации: от протокола IPv4 к протоколу IPv6

В 1981 году трудно было представить, что 32 бита адреса IPv4, позволяющие присвоить уникальный номер четырем миллиардам систем (компьютеров, маршрутизаторов и т.п.), когда-либо станут реальным ограничением. Однако уже к 1992 году масштабируемость и ограниченность адресного пространства IPv4 встала на повестку дня.

Для поиска решения проблемы в ноябре 1991 года организация по стандартизации IETF сформировала специальную группу для «мозгового штурма» в области маршрутизации и адресации — ROAD (Routing and Addressing). Учеными было найдено краткосрочное решение проблемы: они предложили супернеты — концепцию, впоследствии переработанную в архитектуру CIDR (Classless Inter-Domain Routing, бесклассовая междоменная маршрутизация). Этот подход, который был стандартизован в 1993 году (RFC 1518⁵, RFC 1519⁶), позволил существенно замедлить расходование запаса доступных адресов. В чем заключалась суть концепции CIDR? Граница подсетей становилась подвижной в зависимости от фактического размера адресуемой сети. Вместо распределения сетей класса C (/24) фиксированного размера (254 устройства) стало возможным создавать сети /23, /22 и так далее. Внутри же сервис-провайдер мог создать структуру, более соответствующую реальной топологии, распределяя сети меньшего размера, например /25. CIDR предполагал изменения как в системе распределения адресного пространства (об этом мы поговорим позже, в разделе «Глобальная система администрирования адресного пространства»), так и в системе маршрутизации.

Последнее было связано с тем, что фактически произошел отказ от концепции классов сетей (A, B, C и D), в которой деление между сетевым адресом и адресом устройства в сети было предопределено. Для маршрутизации CIDR стало необходимым явно указывать, сколько битов IP-адреса относятся к адресу сети (это данные, которые носят название «сетевая маска»).

⁵ RFC 1518: An Architecture for IP Address Allocation with CIDR, URL: <https://www.rfc-editor.org/rfc/rfc1518>

⁶ RFC 1519: Classless Inter-Domain Routing (CIDR): an Address Assignment and Aggregation Strategy, URL: <https://www.rfc-editor.org/rfc/rfc1519>

Вот что отметили в то время члены руководящего комитета IETF — IESG⁷: «CIDR потребует изменения в политике [распределения адресных ресурсов], спецификации протоколов, разработке и внедрении ПО для маршрутизаторов, но не требуется изменение программного обеспечения конечных устройств». И действительно, новая архитектура была внедрена достаточно быстро. Этому содействовали и относительно небольшой размер Интернета, и его научно-исследовательский характер, и то, что опорная инфраструктура и протоколы находились в стадии разработки.

CIDR позволил избавиться от острых симптомов надвигающейся проблемы, но глобальное решение еще требовалось найти. Поэтому в начале 1994 года IETF начал работу над созданием новой версии протокола IP, позднее получившей название IPv6. Базовая спецификация была опубликована в 1998 году (RFC 2460⁸), а окончательная версия структуры адресации IPv6 — в 2006-м (RFC 4291⁹).

Основные отличия IPv6 от протокола предыдущего поколения — IPv4

Размер адресного пространства

Размер адреса IPv6 составляет 128 бит. Он позволяет адресовать 2^{128} узлов. Это огромное адресное пространство, и масштаб его поистине космический. Например, IPv6 позволяет присвоить 1027 адресов каждой из звезд Млечного Пути. При этом каждая звезда получит адресное пространство в 1018 раз больше, чем весь Интернет IPv4! Очевидно, что дефицита адресов IPv6 в обозримом будущем не предвидится.

IPv6 — это колоссальное количество доступных адресов, это возможность адресации любого мыслимого и немыслимого устройства. Эффективное использование возможностей нового протокола способно породить новый виток информационной революции. Достаточно посмотреть на текущий уровень распределения адресного пространства РИРами (региональными интернет-регистратурами): ясно видно, что их IPv6-пул далек от опустошения (рис. 2).

⁷ RFC 1380: IESG Deliberations on Routing and Addressing, URL: <https://www.rfc-editor.org/rfc/rfc1380>

⁸ RFC 2460: Internet Protocol, Version 6 (IPv6), URL: <https://www.rfc-editor.org/rfc/rfc2460>

⁹ RFC 4291: IP Version 6 Addressing Architecture, URL: <https://www.rfc-editor.org/rfc/rfc4291>

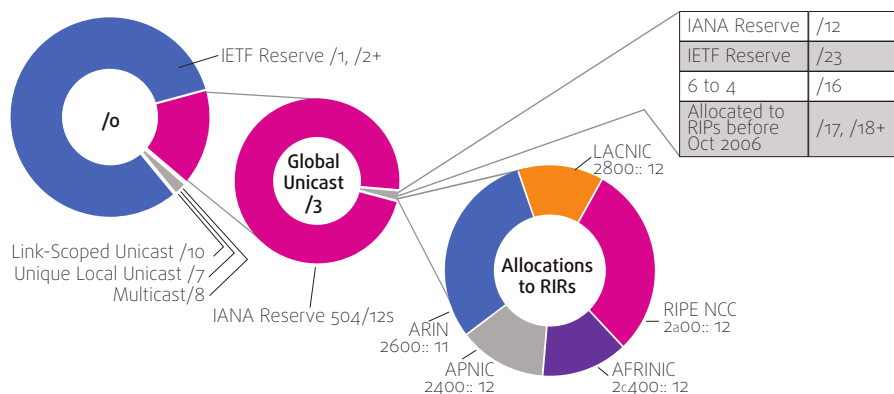


Рис. 2. Адресное пространство IPv6, предоставленное для распределения через региональные интернет-регистратуры. На настоящий момент IPv6-пул РИРов составляет чуть больше пяти блоков /12, и он далек от опустошения; эти блоки составляют ничтожный процент всего доступного в будущем адресного пространства.

Источник: статистика NRO (<http://www.nro.net/statistics>)

Помимо значительного увеличения адресного пространства изначально предполагалось, что IPv6 сможет поддерживать большее число уровней сетевой иерархии и обеспечит оптимальное распределение адресного пространства с точки зрения маршрутизации и конфигурации. Но в этом отношении ожидания создателей не оправдались: довольно жесткая иерархическая структура адресации была отвергнута операторами в пользу гибкой архитектуры CIDR. Также на сегодняшний момент IPv6 унаследовал многие «болячки» IPv4 (например, независимое от провайдера адресное пространство), которые не способствуют сдерживанию роста таблицы маршрутизации.

Расширяемость и дополнительные функции

При разработке протокола IPv6 особое внимание было уделено возможности добавления новых функций без потери эффективности обработки пакетов на сетевом уровне. IPv6 предполагает наличие дополнительных заголовков для различных расширений (extension header, EH) — например, для криптографической защиты данных (Authentication EH и Encapsulating Security Payload EH). В то же время базовый заголовок IPv6 содержит минимальное число полей и имеет фиксированный размер. В частности, в IPv6 маршрутизаторы не производят фрагментацию, поэтому поля, относящиеся к этой функции, перенесены в соответствующий заголовок расширений (Fragmentation EH).

Как видно из рис. 3, заголовки расширений связаны цепочкой указателей Next Header («Следующий заголовок»).

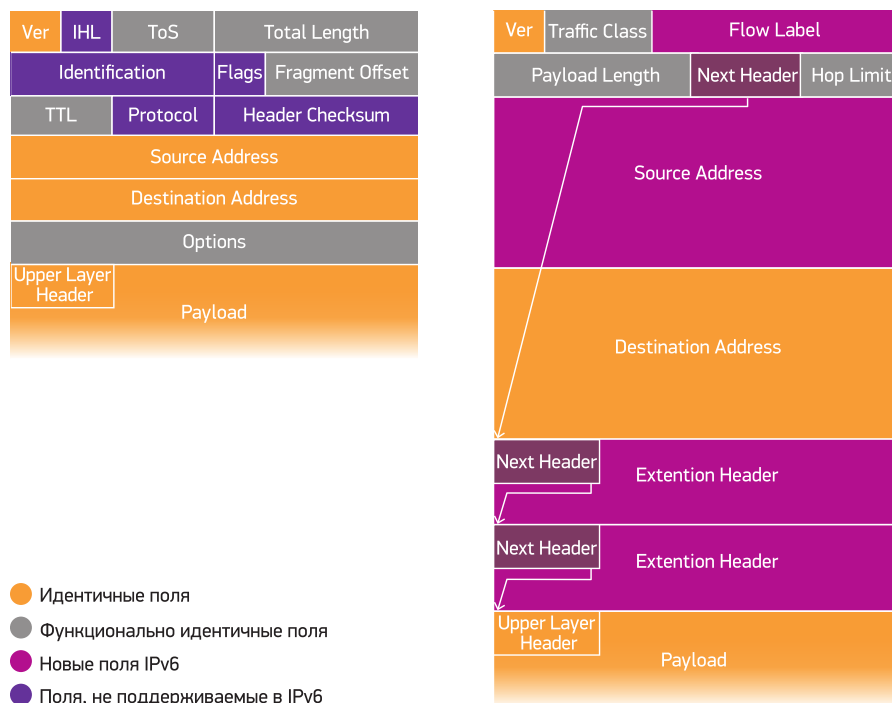


Рис. 3. Форматы пакетов IPv4 и IPv6.

Фрагментация

Как было упомянуто выше, протокол IPv6 иначе обрабатывает фрагментацию пакетов. В случае IPv4, когда маршрутизатор получает пакет, размер которого превышает предел передачи через интерфейс, маршрутизатор производит фрагментацию — дробление пакета на более мелкие части. В дальнейшем они консолидируются получателем в исходный пакет. Заголовок пакета IPv4 имеет соответствующее поле (Fragment Offset), поддерживающее эту функцию.

В IPv6 фрагментация промежуточными устройствами запрещена. Если пакет IPv6 превышает допустимый размер для последующей передачи, маршрутизатор генерирует сообщение ICMP «packet too big» («слишком большой пакет») и посылает его обратно отправителю. В зависимости от приложения отправитель либо выбирает размер пакета, который позволит ему на всем пути следовать без фрагментации, либо дробит пакет самостоятельно. Как и в случае IPv4, консолидация фрагментированных пакетов входит в задачу получателя. Как следствие, передача пакетов IPv6 требует меньших затрат от промежуточного сетевого оборудования.

Автоконфигурация

Для протокола IPv6 была разработана так называемая система автоконфигурации без сохранения состояния (Stateless Autoconfiguration). Данный протокол позволяет различным устройствам, присоединенным к сети IPv6, получить необходимые установки для доступа в Интернет без дополнительных средств — например, без сервиса DHCP (Dynamic Host Configuration Protocol). Суть подхода заключается в том, что устройство получает адрес, состоящий из префикса сети и идентификатора устройства, автоматически сгенерированного с использованием MAC-адреса.

Защита данных

В протокол IPv6 изначально включена система безопасности, основанная на технологии IPsec. Предусмотрено два режима работы: транспортный и туннельный. В транспортном режиме производится защита (шифрование) данных пакета, но не заголовка. С точки зрения маршрутизации такой IP-пакет выглядит вполне обычно, а в задачу получателя входит декодирование содержимого пакета. При использовании туннельного режима данные всего пакета, включая заголовок, шифруются и инкапсулируются в новый пакет. Получатель, указанный в этом новом пакете, является окончанием защищенного канала, или туннеля, и в его задачу входит извлечение изначального пакета и последующая обработка. Дополнительно пакет IPv6 содержит заголовок аутентификации (Authentication EH) для определения подлинности и отсутствия модификации данных пакета.

Мобильность

Поддержка мобильности в протоколах IP означает, что оконечное устройство может изменить свое местоположение в сети и IP-адрес без потери существующих связей, которые соответствуют потокам передачи данных. Для этого мобильные устройства используют отдельные IP-адреса, по которым устройства всегда доступны при передаче данных. За авторизацию мобильного устройства в сети и обеспечение соответствия между реальным и мобильным IP-адресами отвечает «Домашний агент» — устройство, расположенное в «домашней» сети мобильного пользователя. Реализация мобильности в протоколах IPv4 и IPv6 различается. В случае IPv4 передача данных также производится (туннелируется) через «Домашнего агента», в то время как в IPv6 «Домашний агент» обеспечивает только контролирующие функции (авторизацию и обеспечение соответствия между реальным и мобильным адресами). При этом передача данных производится между отправителем и получателем напрямую. Такой подход оптимизирует маршрутизацию данных и, как следствие, повышает качество передачи.

Приведенные особенности протокола IPv6 призваны улучшить производительность, качество и защиту передачи данных. Однако опыт практического внедрения протокола IPv6 показывает, что указанные улучшения весьма незначительны и во многих случаях не используются. Напротив, операторы зачастую прибегают к проверенным методам, разработанным для сетей IPv4. Так, для конфигурации подключенных устройств используется система DHCP, а

в области защиты данных технология IPsec может быть использована в IPv4 почти так же эффективно, как и в IPv6. Эффективная поддержка multihoming (подключения клиента к нескольким сервис-провайдерам для повышения надежности) в IPv6 потребовала отдельного решения и существенно усложнила элегантную структуру маршрутизации, считающуюся одним из преимуществ IPv6. В результате на практике multihoming реализуется аналогично IPv4, что приводит к неоправданному росту таблиц маршрутизации.

Неудивительно, что в среде сетевых операторов существует мнение, что основное преимущество IPv6 — только лишь расширение доступного адресного пространства.

Практика и проблемы внедрения протокола IPv6

Стратегия развития: сосуществование IPv4 и IPv6

Основная проблема перехода от IPv4 к IPv6 — несовместимость двух протоколов. Клиент IPv6 не может напрямую общаться с клиентом, поддерживающим только IPv4.

Изначально представлялось, что эту проблему решит внедрение «двойного стека» — когда компьютеры сети поддерживают оба протокола и подключены как к сети IPv4, так и к сети IPv6. Данное разделение является логическим, а физически используется одна и та же сетевая инфраструктура. Для доступа к ресурсам IPv4 используется протокол IPv4, а к ресурсам IPv6 — протокол IPv6. Все достаточно просто, но... Темпы внедрения IPv6 оказались недостаточными. План «двойного стека» мог сработать, если бы подавляющее большинство компьютеров Интернета имело доступ как к IPv4, так и к IPv6 до того, как пул адресов IPv4 опустошился. В таком случае можно было бы просто отключить поддержку IPv4 и — чудо! — Интернет просто перешел бы на новый протокол. Однако реальность оказалась сложнее.

Сложность внедрения протокола IPv6 во многом связана с так называемым сетевым эффектом. Этот экономический термин описывает явление, когда ценность технологии зависит от числа игроков, ее использующих. Действительно, возможность обмениваться трафиком IPv6 с парой других энтузиастов, как это было в начале 2000-х, с практической точки зрения не представляет особого интереса. Даже при сегодняшнем уровне использования IPv6 (по оценкам Google почти 40% запросов используют этот протокол¹⁰) большая часть Интернета по-прежнему доступна только через протокол IPv4. Размер этой части Интернета определяет значимость протокола IPv4 и, в обратной пропорции, протокола IPv6 для сервис-провайдеров.

¹⁰ Процент пользователей, использующих IPv6 для доступа к услугам Google, <https://www.google.com/intl/en/ipv6/statistics.html#tab=ipv6-adoption>

Каждый новый подключенный клиент должен иметь возможность обмениваться данными с Интернетом по протоколу IPv4, что требует предоставления ему адреса IPv4. Скажем прямо, для растущих сервис-провайдеров, возможно, более приоритетным является решение проблемы отсутствия адресов IPv4, чем внедрение IPv6. В то же время важно отметить, что обсуждаемая стратегия и динамика сосуществования двух протоколов основана на предположении, что инфраструктура сервис-провайдера обеспечивает полноценную поддержку IPv6.

Динамика потребности в адресном пространстве IPv4 по мере глобального внедрения IPv6 показана на рис. 4. На нем фиолетовой линией обозначен рост глобального Интернета. По мере внедрения протокола IPv6 доля Интернета, доступного только по IPv4, будет неуклонно уменьшаться (оранжевая кривая). Синяя линия отображает размер сервис-провайдера, характеризуемый, например, числом подключенных пользователей. В данном случае рассматривается растущий провайдер. Наконец, потребность в адресах IPv4 показана кривой зеленого цвета.

По мере расширения клиентской базы провайдера пропорционально увеличивается потребность в дополнительных адресах IPv4. В то же время все большая и большая часть Интернета становится доступной по протоколу IPv6, что выражается в обратной тенденции, когда все меньшее число пользовательских соединений основано на протоколе IPv4. Соответственно, потребность в адресах IPv4 снижается. Наконец, когда подавляющее большинство ресурсов Интернета станет доступным по IPv6, потребность в IPv4 станет ничтожной. Таким образом, завершится фаза перехода Интернета на протокол IPv6. Продолжительность этой фазы может составить несколько лет. Не исключена, правда, вероятность, что данная фаза не закончится никогда, но об этом — чуть позже.

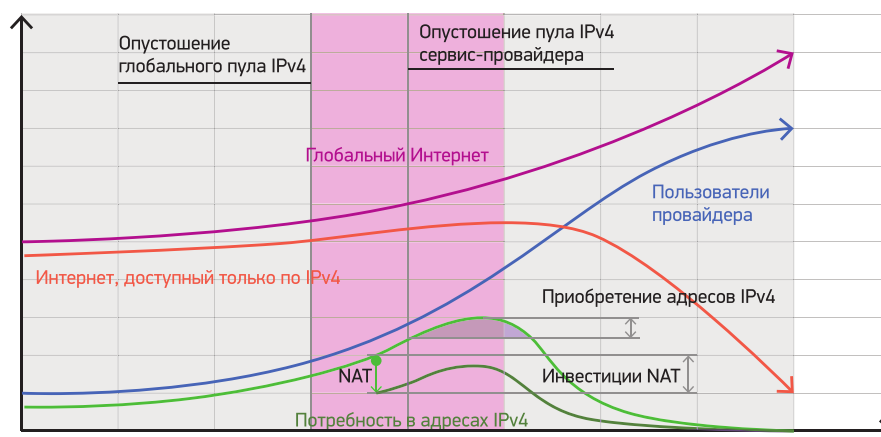


Рис. 4. Динамика сосуществования IPv4 и IPv6.

Как видно из графика, наиболее критичной фазой для сервис-провайдера является промежуток времени с момента опустошения глобального свободного пула IPv4 до момента, когда потребность в дополнительных адресах IPv4 начнет уменьшаться. Эта фаза отмечена на графике сиреневым цветом.

Надо заметить, что высота порога, образуемого зеленой кривой, для разных провайдеров отличается. Также различен момент завершения свободных адресов в собственном пуле провайдера (вторая вертикальная синяя линия). Другими словами, умеренно растущий провайдер с достаточным запасом свободных адресов имеет шансы «перезимовать» переходный период без особых ухищрений. Важно отметить, что и в этом случае необходимой является полноценная поддержка IPv6 в инфраструктуре провайдера и неуклонное массовое распространение IPv6 в глобальном Интернете. Все большее число сервис-провайдеров страдают от проблемы нехватки IPv4.

Существует два способа решения этой проблемы. Первый – это получение дополнительных адресов IPv4. Однако в соответствии с текущей политикой распределения оставшегося адресного пространства IPv4¹¹ максимум, на что может рассчитывать провайдер, – это одноразовый блок размером /24, да и то придется подождать, пока такой блок появится, например, вследствие возврата адресов закрывшегося оператора. Рассчитывать на это не стоит – на февраль 2023 года список ожидания превысил 1200 претендентов, и уже больше года свободных блоков не появлялось.

Поэтому более практичный вариант получения дополнительных адресов – это их покупка на рынке IPv4. Уже несколько лет на рынке работают так называемые брокеры, связывающие желающих продать и купить. Передача адресов от одной организации к другой регламентирована соответствующими политиками RIPE – «Передача цифровых интернет-ресурсов и изменение официального юридического имени члена»¹² и «Передача цифровых интернет-ресурсов между региональными интернет регистратурами»¹³. Однако следует иметь в виду, что покупка адресов – дорогостоящая операция. За последние два года цена в расчете на один адрес удвоилась и достигла \$45-55. Хотя наблюдается некоторая стабилизация и даже охлаждение рынка в 2022, скорее всего, это явление временное, и цены продолжают расти, см. рис. 5.

¹¹ IPv4 Address Allocation and Assignment Policies for the RIPE NCC Service Region, секция 5.1, <https://www.ripe.net/publications/docs/ripe-733>

¹² Transfer of Internet Number Resources and Change of a Member's Official Legal Name, RIPE-758, 30 марта 2021 г., <https://www.ripe.net/publications/docs/ripe-758>

¹³ Inter-RIR Transfer of Internet Number Resources, RIPE-769, 24 ноября 2021 г., <https://www.ripe.net/publications/docs/ripe-769>

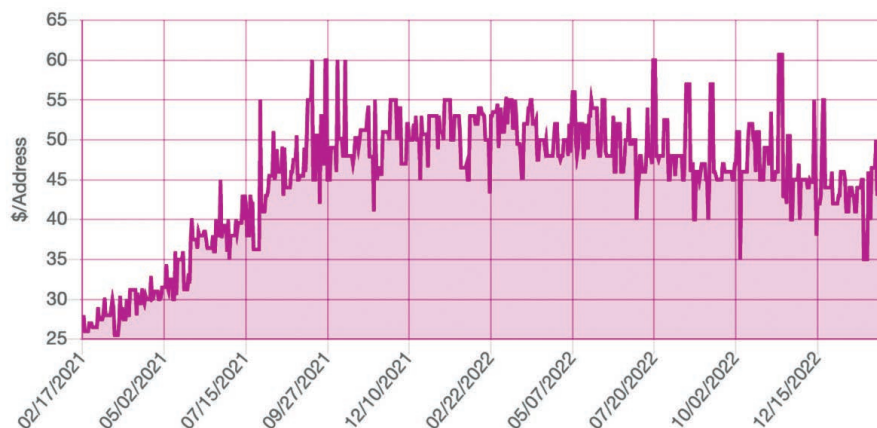


Рис 5. Колебание цен на адреса IPv4 в пересчете на один адрес.

Источник: <https://ipv4.global/>

В этой связи второй способ — повышение эффективности использования адресного пространства с помощью технологии NAT (Network Address Translation) — является более реальной альтернативой или дополнительным решением. Этот сценарий показан на графике кривой розового цвета.

Поскольку мы заговорили о технологии NAT, пожалуй, стоит остановиться на ней поподробнее. Ведь эта технология является ключевой в моделях сосуществования IPv4 и IPv6.

Техническое отступление: как происходит передача данных в Интернете

Прежде чем перейти непосредственно к разговору о будущем Интернета и перспективах IPv6, давайте совершим краткий экскурс в техническую область и в общих чертах рассмотрим, как же происходит передача данных в Интернете и какую роль играют адреса. Работа Интернета основана на технологии пакетной коммутации без установления соединения. Структура пакета определена протоколом IP, при этом каждый пакет содержит IP-адрес отправителя и получателя. В задачу каждого узла сети (также называемого маршрутизатором) входит передача пакета, полученного от соседнего узла, к последующему узлу. Выбор каждого следующего узла происходит с помощью системы маршрутизации. Благодаря этой системе маршрутизатор знает, какому из своих соседей следует передать пакет с конкретным IP-адресом получателя.

Однако с точки зрения пользователя передача данных происходит между его приложением и приложением получателя. Например, между веб-браузером и веб-сайтом. Поэтому можно представить, что существует виртуальное соединение между этими приложениями: по нему и происходит передача данных. Помимо IP-адреса отправителя (в данном случае — компьютера пользователя) и IP-

адреса получателя (веб-сервера) это соединение характеризуется дополнительными параметрами — так называемыми портами получателя и отправителя. Их можно рассматривать как локальные идентификаторы конкретных приложений на компьютере. Наконец, транспортный протокол (например, TCP или UDP) является пятым параметром, однозначно определяющим поток данных в Интернете в пределах ограниченного времени.

Таким образом, отправитель и получатель данных в действительности каждый адресуются парой {IP-адрес, порт}. Именно эта особенность используется в технологии NAT (Network Address Translation), или более точно — NAPT (Network Address & Port Translation). С помощью одного IP-адреса можно теоретически адресовать 65 535 «соединений» — число, значительно превышающее потребности единичного пользователя. В этом случае устройство NAT для внешней сети будет выглядеть как компьютер с очень большим числом одновременно работающих приложений. Хотя на самом деле устройство NAT при передаче пакетов подставляет вместо порта и собственного IP-адреса (как адреса получателя с точки зрения внешних приложений) порт и локальный IP-адрес реального получателя. Обычно для адресации конечных устройств локальной сети, расположенной за устройством NAT, используется специальное зарезервированное адресное пространство. Схема работы NAT показана на рис. 6.

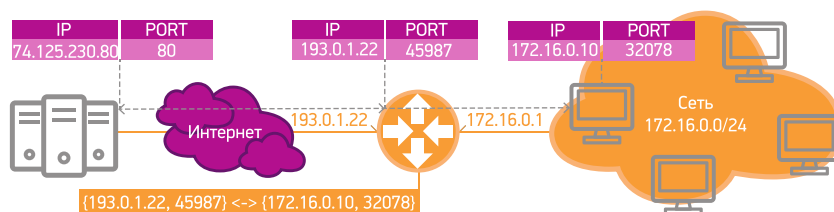


Рис. 6. Схема работы NAT.

Насколько эффективен NAT? Это зависит от характера приложений, работающих на конечных устройствах, и интенсивности их взаимодействия с глобальным Интернетом. На сегодня компьютер обычного пользователя во время работы в сети создает от 60 до 100 соединений с различными ресурсами глобального Интернета. Цифра может показаться большой, но ведь многие приложения открывают более одного соединения — так функционирует большинство веб-приложений. Например, Google Maps одновременно использует несколько десятков соединений. Но даже если эта цифра на порядок крупнее, трудностей не возникает: технология NAT позволяет совместно использовать один и тот же IP-адрес более чем 60 пользователям.

Звучит очень привлекательно — но, к сожалению, в реальности все не так радужно. Технология NAT содержит ряд серьезных недостатков, о которых мы поговорим позже. Здесь же отметим, что NAT нарушает принцип «прозрачности»

соединений между любыми конечными устройствами в Интернете. Помимо усложнения архитектуры сети, для полноценной работы некоторых приложений требуются дополнительные средства, такие как STUN¹⁴, ICE¹⁵, TURN¹⁶. Использование каскадов NAT, когда в сети за устройством NAT расположены еще и NAT со «вложенными» сетями, только усугубляет эти проблемы.

Переходные технологии сосуществования

Итак, технология NAT-мультиплексирования — еще один метод решения проблемы сосуществования двух протоколов, позволяющий бороться с острой нехваткой адресов IPv4. Однако по-прежнему одним из основных препятствий перехода к IPv6 является его несовместимость со своим предшественником — протоколом IPv4. Устройство, поддерживающее только IPv6, не может непосредственно обмениваться данными с устройством IPv4. Виной этому является, скорее, протокол IPv4, который был разработан для адресации нескольких десятков, может быть — сотен или тысяч устройств Сети и не предусматривал способа расширения.

Переходный план «двойного стека» предполагал отсутствие устройств, «говорящих» только на одном из протоколов, другими словами — глобальное двуязычие. Соответственно, он основан на предположении, что все рассматриваемые устройства имеют адреса IPv4. Опустошение пула адресов IPv4 существенно ограничило сферу применения этого подхода. Поэтому для обмена данными между устройствами и сетями разных протоколов необходимо применение дополнительных технологий — так же как мы прибегаем к услугам переводчика для преодоления языкового барьера.

За последние два десятилетия было предложено множество решений, но не все оказались эффективными и устойчивыми, а многие вообще не прижились.

Давайте посмотрим, что же имеется в арсенале сервис-провайдеров на сегодняшний день. Помимо «двойного стека», предполагающего прозрачную связность, переходные технологии делятся на два типа: туннелирование и трансляцию.

Технологии туннелирования

Технологии туннелирования приходят на помощь, когда инфраструктура сервис-провайдера не поддерживает один из протоколов.

¹⁴ RFC 5389: Session Traversal Utilities for NAT (STUN),
URL: <https://www.rfc-editor.org/rfc/rfc5389>

¹⁵ RFC 5245: Interactive Connectivity Establishment (ICE): A Protocol for Network Address Translator (NAT) Traversal for Offer/Answer Protocols,
URL: <https://www.rfc-editor.org/rfc/rfc5245>

¹⁶ Traversal Using Relay NAT,
URL: https://ru.wikipedia.org/wiki/Traversal_Using_Relay_NAT;
RFC 5766: Traversal Using Relays around NAT (TURN): Relay Extensions to Session Traversal Utilities for NAT (STUN),
URL: <https://www.rfc-editor.org/rfc/rfc5766>

6rd

«rd» в названии этой технологии означает «rapid deployment», или быстрое развертывание. И действительно, эта технология позволила одному из крупнейших французских провайдеров Free в течение пяти недель осуществить поддержку IPv6 для пользователей сети. Технология была после этого стандартизована в IETF¹⁷. 6rd делает доступным Интернет IPv6 пользователям провайдера широкополосного доступа, не требуя при этом поддержки IPv6 в сети самого провайдера.

Во-первых, сеть 6rd использует собственное адресное пространство IPv6, полученное от региональной интернет-регистратуры. Это позволяет сервис-провайдеру анонсировать реальные IPv6-префиксы и, таким образом, более точно определять собственную политику маршрутизации.

Во-вторых, вся зона функционирования 6rd ограничена сетью сервис-провайдера. Так называемые шлюзы 6rd встроены в оконечное оборудование клиента (CPE, customer premise equipment), а релеи являются частью инфраструктуры сервис-провайдера.



Рис. 7. Схема работы технологии 6rd.

DS-Lite

DS-Lite¹⁸ в некотором смысле является зеркальной технологией по отношению к 6rd. DS-Lite предполагает, что сеть провайдера полностью поддерживает IPv6, а туннели используются для передачи трафика IPv4 от сети пользователя к устройствам NAT сервис-провайдера. Также подразумевается, что устройства сети пользователя поддерживают «двойной стек», а именно оба протокола IPv4 и IPv6.

Суть метода заключается в одновременном применении технологий туннелирования (инкапсуляция трафика IPv4 в пакеты IPv6) и централизованного NAT,

¹⁷ RFC 5969: IPv6 Rapid Deployment on IPv4 Infrastructures (6rd) - Protocol Specification, URL: <https://www.rfc-editor.org/rfc/rfc5969>

¹⁸ RFC 6333: Dual-Stack Lite Broadband Deployments Following IPv4 Exhaustion, URL: <https://www.rfc-editor.org/rfc/rfc6333>

или CGN (Carrier Grade NAT, также называемого LSN, Large Scale NAT). Благодаря этому ограниченный пул адресов IPv4 совместно используется всеми пользователями сервис-провайдера. Обмен трафиком с интернет-ресурсами IPv4 происходит с использованием протокола IPv4, а с ресурсами IPv6 — с использованием IPv6. Эта схема не предусматривает трансляции протоколов.

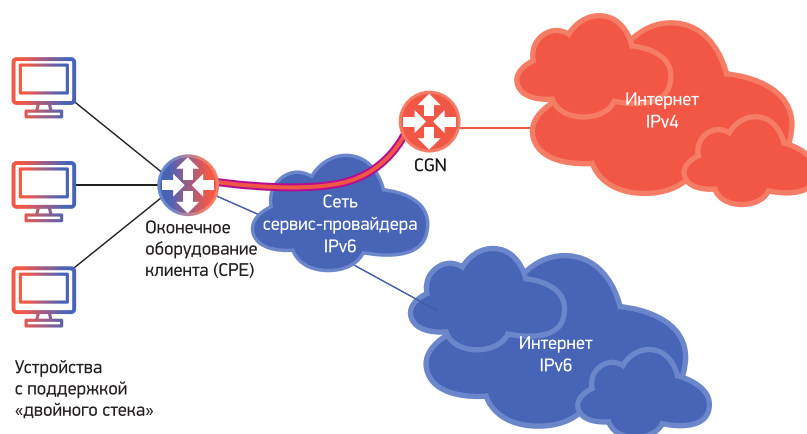


Рис. 8. Схема работы DS-Lite.

На рис. 8 представлена схема работы DS-Lite. Как можно заметить, обмен трафиком с ресурсами IPv6 происходит непосредственно, без использования каких-либо промежуточных технологий, например, туннелей.

В отношении IPv4 ситуация гораздо сложнее. Нехватка адресного пространства IPv4 — это серьезная проблема для растущего числа сетей. Поэтому схемы, предусматривающие назначение каждому абоненту публичного адреса IPv4, используемого устройством NAT-пользователя для построения домашней локальной сети, имеют все более ограниченное применение.

Возможным решением этой проблемы (кстати, уже применяемым некоторыми сервис-провайдерами) является создание еще одного уровня NAT в сети сервис-провайдера. Такая схема работает в общем случае, но результат ее применения — существенные ограничения для многих сегодняшних и будущих приложений, а также сложность обслуживания.

Задача DS-Lite — исключить каскадирование устройств NAT, когда все устройства пользователей непосредственно взаимодействуют с центральным устройством NAT сервис-провайдера. В этом случае оконечное устройство пользователя не выполняет никаких функций NAT, а вместо этого обеспечивает создание туннелей к центральному NAT для каждого нового соединения между приложениями пользователя и сервисами Интернета.

Таким образом, все пользовательские соединения, так же, как и в схеме каскадирования NAT, отображаются центральным CGN. Однако значительно по-

вышается прозрачность архитектуры, растет эффективность использования адресного пространства IPv4.

Кстати, о прозрачности. Одна из основных проблем, связанных с применением NAT, — это контроль приложений за значениями порта и IP-адреса соединений, поскольку устройство NAT заменяет их на динамически присваиваемые. От этого зависит нормальное функционирование некоторых приложений, например, большинства мультимедийных интерактивных программ. На сегодняшний день разработано несколько механизмов решения этой проблемы — такие как STUN, ICE и TURN. Но очевидно, что каскадирование устройств NAT усложняет ситуацию.

В то же время, поскольку централизованная трансляция адресов осуществляется для каждого сетевого потока, DS-Lite сложно масштабировать, особенно в сетях крупных провайдеров. Чтобы решить эту проблему, в IETF было предложено расширение DS-Lite, получившее название «Lightweight 4over6», или lw4ob. Новый подход требует поддержки состояния не для каждого потока, а только для каждого абонента и перемещает функцию NAT обратно на клиентское оборудование (CPE). Технология lw4ob была также стандартизована в IETF¹⁹.

Отметим, что технологии DS-Lite и lw4ob не предусматривают поддержку устройств, работающих только по протоколу IPv6. Для этого используются технологии трансляции.

Технология трансляции: NAT64 + DNS64

Логично предположить, что в недалеком будущем появятся устройства, поддерживающие только IPv6. Если мы говорим о масштабных мобильных, сенсорных или RFID-сетях, необходимость поддержки двух протоколов усложнит и удорожит такие устройства.

Для взаимодействия таких сетей с Интернетом IPv4 необходимо применение трансляции адресов IPv6 в адреса IPv4 и обратно. Ввиду недостатка ресурсов IPv4 здесь, как и в случаях, рассмотренных выше, необходимо применение мультиплексирования потоков. По существу, нужно использовать технологию централизованного NAT с внедрением дополнительной функции трансляции протоколов. Этот компонент еще называют NAT64²⁰. Взаимодействие с другими сетями IPv6 происходит прозрачно: эта архитектура показана на рис. 9.

Однако в данной схеме есть одна особенность, а именно необходимость дополнительной поддержки одного из наиболее критических приложений Интернета — системы доменных имен DNS.

¹⁹ RFC7596: Lightweight 4over6: An Extension to the Dual-Stack Lite Architecture, URL: <https://www.rfc-editor.org/rfc/rfc6146>

²⁰ RFC 6146: Stateful NAT64: Network Address and Protocol Translation from IPv6 Clients to IPv4 Servers, URL: <https://datatracker.ietf.org/doc/rfc6146>

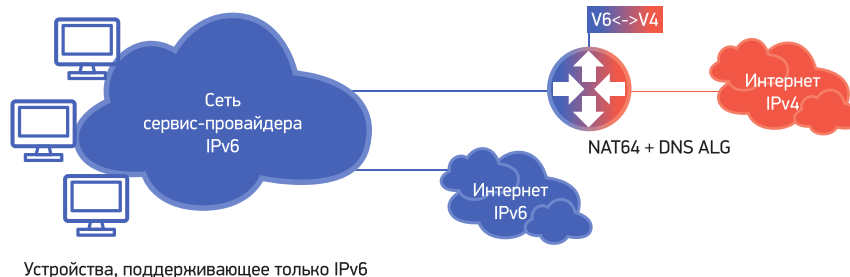


Рис. 9. Архитектура системы трансляции протоколов NAT64.

Дело в том, что для большей части ресурсов Интернета запрос DNS вернет адрес IPv4. Поскольку сети, о которых идет речь, поддерживают только IPv6, такой ответ DNS вряд ли окажется полезным. Для решения этой проблемы используется дополнительный компонент — шлюз приложений (Application Layer Gateway, ALG). Суть его заключается в замещении адреса IPv4 в ответе DNS на синтезированный адрес IPv6, который понятен и клиенту, и транслятору протоколов NAT64.

Работа DNS ALG происходит следующим образом. Как обычно, перед началом связи клиент посылает запрос локальному DNS-серверу. В нашем случае его роль выполняет ALG. Он производит разрешение запроса и, допустим, получает IPv4-адрес искомого ресурса. Но в ответ клиенту ALG подставляет синтезированный адрес IPv6. По существу, этот адрес состоит из предустановленного префикса (известного и ALG, и NAT64), а также из IPv4-адреса ресурса. Теперь, когда клиент попытается установить связь с ресурсом, NAT64 поймет, что клиент использует синтезированный адрес, и преобразует его в исходный IPv4-адрес получателя. Схематически это показано на рис. 10.

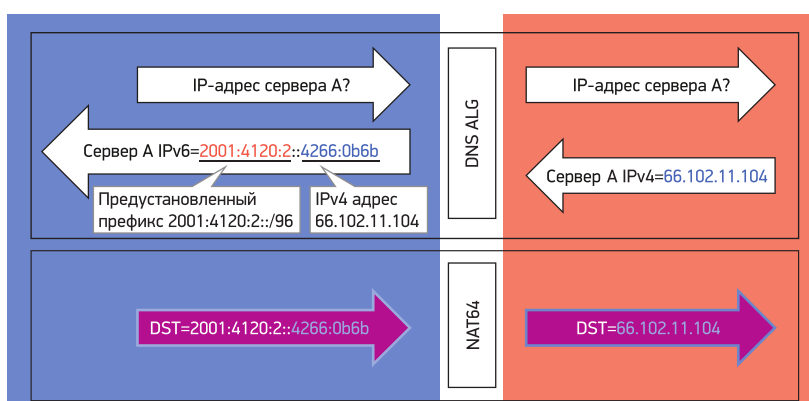


Рис. 10. Схема работы DNS ALG.

Вопросы внедрения IPv6 в мобильных сетях

Сегодня разговор об эволюции Интернета немаловажен без взгляда на мобильные сети. Здесь мы наблюдаем наиболее стремительный рост как по количеству абонентов, так и по возможностям, которые они открывают для пользователей. Архитектурные решения, принимаемые при разработке или модернизации мобильных сетей, определяют архитектуру будущего Интернета. Поставим вопрос более остро: останется ли Интернет уникальной коммуникационной средой с колоссальным инновационным потенциалом или наше информационное пространство будут определять закрытые и ограниченные платформы мобильных приложений, такие как Apple Store или Google Play?

Развитие мобильных сетей началось с сетей мобильной телефонии, основанных на телефонных стандартах и технологии коммутации каналов. Передача данных была внедрена позже как отдельная подсистема, существенно отличающаяся как по архитектуре, так и по используемым технологиям. Так, для обеспечения услуг на основе пакетной передачи — в первую очередь для доступа к Интернету — в сетях 2G и 3G была разработана система GPRS (General Packet Radio Service). Для возможности предоставления услуг голосовой связи на основе протокола IP в 2002 году была разработана система IMS (IP Multimedia System). Сегодняшние 3G-сети предоставляют услуги передачи данных и доступа в Интернет в качестве стандартного пакета, однако для осуществления голосовой связи, как правило, по-прежнему используются сети коммутации каналов.

Появление сетей следующего поколения LTE/4G существенно изменило ситуацию. Действительно, эти сети используют исключительно технологию пакетной передачи на основе протокола IP. Для оператора это означает возможность унификации передачи голоса и данных. И хотя для связи с традиционными телефонными сетями необходимы шлюзы, связь между абонентами собственной сети и сетями партнеров, также использующих эти технологии, а также предоставление доступа в Интернет осуществляется унифицированной инфраструктурой на основе пакетной коммутации IP.

В архитектуре LTE опорная сеть, так называемая EPC (Evolved Packet Core), представляет собой нормальную сеть пакетной коммутации на основе IP. Дополнительные устройства и шлюзы необходимы для контроля доступа к услугам, поддержки мобильности и роуминга, а также биллинга. Схематично структура сети LTE приведена на рис. 11.

Так, узел управления мобильностью MME (Mobility Management Entity) осуществляет контроль доступа к сети LTE, производит аутентификацию пользователя и поддерживает функции мобильности. Обслуживающий шлюз SGW (Serving Gateway) является маршрутизатором доступа, он поддерживает в том числе и переход мобильного терминала между базовыми станциями (eNodeB). Наконец, пакетный шлюз (PGW, Packet Data Network Gateway) является граничным маршрутизатором, обеспечивающим связность с другими системами: IMS и внешними сетями, например Интернетом.

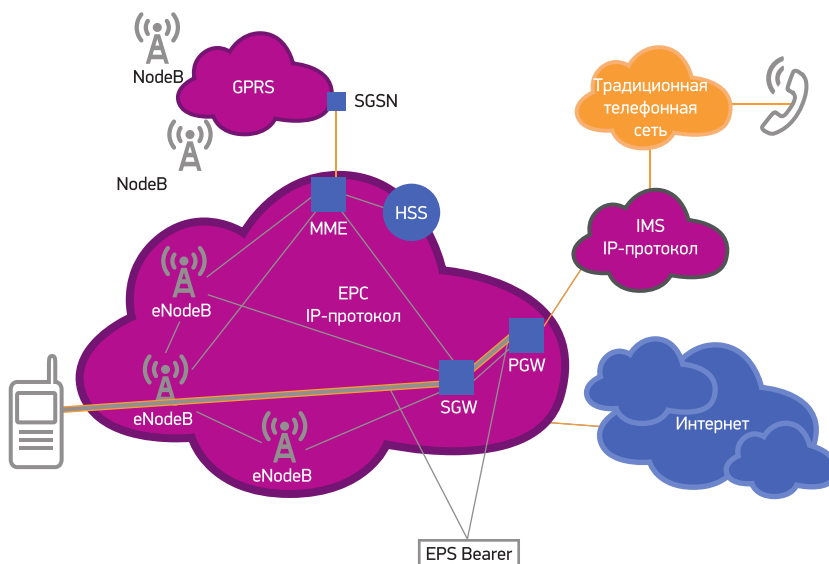


Рис. 11. Архитектура сети LTE.

Здесь стоит отметить одну архитектурную особенность мобильной сети. Передача данных во внутренней инфраструктуре происходит по виртуальному каналу, или туннелю, соединяющему мобильное пользовательское устройство, например телефон, с пакетным шлюзом PGW. В сетях LTE такие туннели называются носителями EPC, а в сетях GPRS использовался термин «PDP-контекст». Специальные протоколы отвечают за создание виртуального канала и резервирования определенных ресурсов сети, в том числе самого ценного ресурса — радиоканала.

Поскольку стандарт LTE поддерживает только пакетную коммутацию IP, потребовалась новая структура сети передачи голоса. Ведь, как уже упоминалось, в мобильных сетях предыдущих поколений голосовая связь была основана на коммутации каналов, а передача данных являлась сопутствующей подсистемой.

Одним из наиболее перспективных подходов обеспечения голосовой связи является архитектура VoLTE (Voice over LTE — «Голос поверх LTE»). Этот подход использует подсистему IMS (IP Multimedia Subsystem), полностью основанную на IP и использующую стандарты передачи голоса в IP-сетях — SIP и RTP. В результате голосовые и управляющие соединения представляют собой не что иное, как обычные потоки данных сети LTE.

Итак, на повестку дня достаточно остро встает задача внедрения IMS. Но есть и еще один важный вопрос: а какой из протоколов применять — IPv4 или IPv6? С одной стороны, так же, как и в сетях фиксированной связи, протокол IPv4 хорошо отработан и поддерживается всеми производителями оборудования.

С другой стороны, нехватка адресного пространства IPv4 заставляет операторов применять более сложные системы трансляции и динамического назначения адресов. К тому же следует учитывать колоссальные масштабы мобильных сетей, насчитывающих миллионы абонентов. Ведь даже использование зарезервированных адресных блоков, например, 10.0.0.0/8, наталкивается на серьезные ограничения. Действительно, в рамках этого блока, широко используемого в сетях с трансляцией адресов, максимальное число адресуемых устройств составляет всего 16,7 миллиона — для многих операторов это меньше, чем число абонентов. Все перечисленные факторы заставляют мобильных операторов серьезно задуматься об использовании протокола IPv6.

Но, как мы уже обсуждали, внедрение IPv6 само по себе не является панацеей — на сегодняшний день Интернет в основном доступен по протоколу IPv4. То есть, хотя внедрение IPv6 является разумной долгосрочной стратегией, на текущем этапе необходимо также обеспечить доступ к сетям и ресурсам IPv4, а значит — использовать технологии сосуществования.

Внедрение IPv6 в мобильных сетях на основе технологии трансляции XLAT

В предыдущих разделах мы рассмотрели различные технологии сосуществования — изначальную стандартную архитектуру «двойного стека», технологии туннелирования и трансляции. До реализации возможности создания виртуальных каналов (PDP-контекст в GPRS и носитель EPC в LTE), поддерживающих «двойной стек», данная архитектура требовала дублирования туннелей и, соответственно, сетевых ресурсов. Поэтому до появления релизов R8 и R9, в которых такие соединения были определены для EPC и GPRS соответственно, этот подход считался неприемлемым. Сегодня — это рабочая альтернатива, серьезно рассматриваемая некоторыми операторами.

Однако поддержка одновременно обоих протоколов в базовой системе, включая EPC (в сетях 3G — GPRS) и IMS, означает существенное удорожание инфраструктуры и обслуживания. Также серьезной проблемой остается нехватка адресов в зарезервированных блоках, требующая разбиения сети на отдельные домены и контроля отсутствия адресных конфликтов. Более подробное обсуждение вопросов внедрения архитектуры «двойного стека» в мобильных сетях и связанных с этим проблем можно найти в RFC 6459²¹.

В связи с этим более реалистичным представляется подход, основанный на архитектуре трансляции, при котором опорная сеть и базовые системы поддерживают только один протокол — IPv6. Мобильные устройства при этом поддерживают оба протокола, а поддержка приложений IPv4 происходит с помощью трансляции IPv4 в IPv6, передачи данных опорной сетью и обратной трансляции шлюзом взаимодействия с сетями IPv4.

²¹ RFC 6459: IPv6 in 3rd Generation Partnership Project (3GPP) Evolved Packet System (EPS), URL: <https://www.rfc-editor.org/rfc/rfc6459>

Эта архитектура, получившая название 4G4XLAT и описанная в RFC 6877²², во многом похожа на уже рассмотренную нами архитектуру DS-Lite. Основным отличием является то, что при передаче данных по IPv6-сети провайдера происходит адресная трансляция, а не туннелирование. Данный подход отличается относительной простотой, а кроме этого, более эффективно используется полоса пропускания — важный фактор в беспроводных сетях. Схема этой архитектуры приведена на рис. 12.

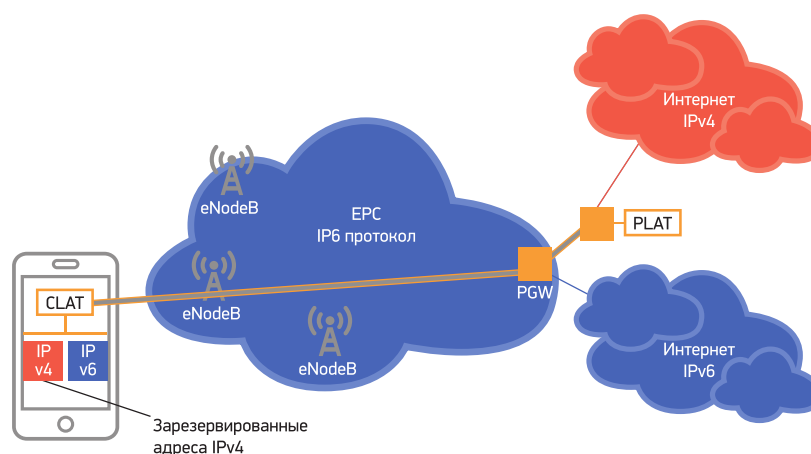


Рис. 12. Архитектура 4G4XLAT с опорной сетью, поддерживающей только IPv6.

В то время как передача данных IPv6 происходит в сетях 4G4XLAT абсолютно прозрачно, для передачи трафика IPv4 производится двойная трансляция. Сначала производится трансляция IPv4 в IPv6 на стороне клиента. Для этого используется так называемое устройство CLAT (customer-side translator, транслятор со стороны клиента). В мобильной сети функция CLAT реализована в пользовательском оборудовании — мобильном телефоне. CLAT является транслятором без сохранения состояния, что существенно упрощает его реализацию и работу. Дело в том, что все адресное пространство IPv4 может быть однозначно отображено в часть адресного пространства IPv6, а точнее, в пространство, определяемое IPv6-префиксом длиной 96 бит (/96).

На другой стороне IPv6-сети провайдера производится обратная трансляция. При этом используется тот же алгоритм преобразования, только в обратную сторону. Другими словами, если IPv6-адрес получателя 2001: DB8: AAAA::198.51.100.1, то соответствующий ему адрес IPv4 — 198.51.100.1. Этой задачей занимается устройство PLAT (provider-side translator, транслятор со стороны провайдера).

²² RFC 6877: 4G4XLAT: Combination of Stateful and Stateless Translation, URL: <https://www.rfc-editor.org/rfc/rfc6877>

Во многих случаях PLAT выполняет также функцию стандартного транслятора NAT. Дело в том, что для IPv4-адресации мобильных терминалов многие операторы используют зарезервированные адресные блоки (например, 10.0.0.0/8). При передаче пакета в Интернет такие адреса транслируются в глобальные адреса из пула провайдера. Напомним, что NAT является устройством, сохраняющим состояние, и поэтому более сложным и дорогостоящим, чем CLAT.

В архитектуре 4G/LTE ограничение размера пула зарезервированных адресов не является проблемой, поскольку каждый CLAT может быть однозначно идентифицирован уникальным IPv6-префиксом, назначенным ему оператором, например, при подключении к сети.

Хотя существуют убедительные примеры использования этой архитектуры в сетях мобильных операторов, многих останавливает недостаточная ее поддержка ведущими разработчиками операционных систем для смартфонов. В настоящее время только Android продолжает поддерживать трансляцию XLAT.

Вопросы роуминга при внедрении IPv6

У различных мобильных операторов уровень поддержки и стратегия внедрения IPv6 могут существенно отличаться. Поэтому обеспечение роуминга требует тщательного анализа вероятных проблемных ситуаций.

Роумингом называется возможность предоставления услуг сотовой связи абоненту вне зоны обслуживания его «домашней» сети другим оператором, так называемой гостевой сетью. При этом абоненту не требуется заключать договор с принимающим оператором, а плата за услуги взимается «домашним» оператором. Услуга роуминга требует предварительной взаимной договоренности между операторами.

В общих чертах роуминг осуществляется следующим образом. При включении мобильного устройства вне домашней сети оно просканирует все радиоканалы в поиске сети, к которой можно подключиться. При подключении узел MME (или узел SGSN в случае 3G/GPRS сети) сначала сделает запрос в домашнюю сеть пользователя к серверу HSS (или HLR в случае 3G) для получения профиля абонента и последующей его аутентификации. Профиль абонента, помимо прочего, содержит информацию о варианте маршрутизации и доступных типах PDP-контекстов (носителей EPC в сетях LTE). По завершении процесса регистрации возможно создание PDP-контекста. Здесь, в зависимости от конфигурации абонента, возможны два варианта: маршрутизация через домашнюю сеть и маршрутизация с местным выходом.

В первом случае при активации контекста PDP устройству абонента будет назначен IP-адрес из домашней сети. Маршрутизация всего трафика будет проходить через домашнюю сеть. Во втором варианте IP-адрес назначается из гостевой сети, соответственно, и трафик во внешние сети, например

Интернет, будет передаваться из гостевой сети, без захода в домашнюю сеть. Тем самым может быть достигнут оптимальный маршрут. Различие между этими двумя вариантами показано на рис. 13.

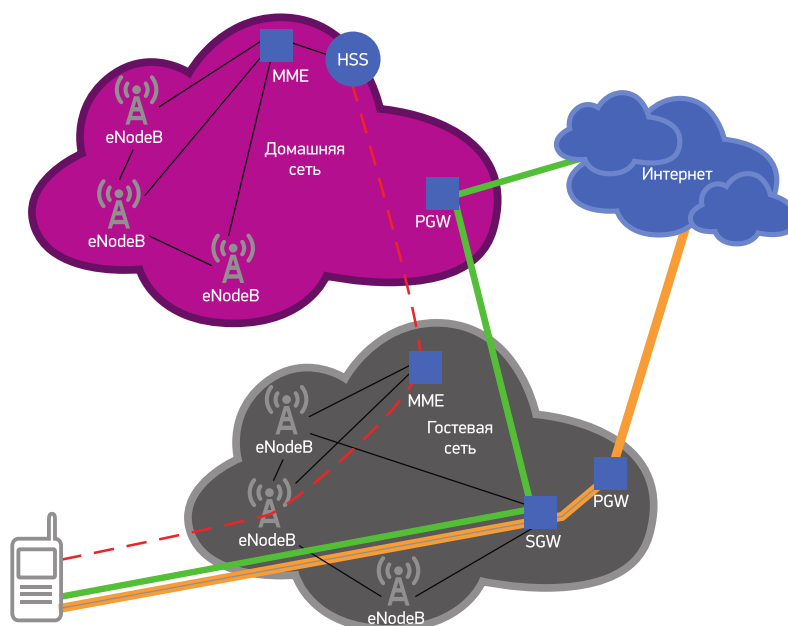


Рис. 13. Два варианта маршрутизации трафика при роуминге: зеленым цветом показан PDP-контекст при маршрутизации через домашнюю сеть, оранжевым — при маршрутизации с местным выходом; красная пунктирная линия отображает канал сигнализации.

Стандарты 3GPP определяют три типа PDP-контекста: PDP/PDN типа IPv4, PDP/PDN типа IPv6 и PDP/PDN типа IPv4v6. Последний тип контекста был введен в спецификации начиная с Релиза 9 для поддержки архитектуры «двойного стека».

Проблема при подключении абонента к гостевой сети может возникнуть, если данная сеть не поддерживает контекста IPv4v6 и, соответственно, не может правильно интерпретировать соответствующую запись в профиле абонента. В этом случае абоненту будет отказано в подключении. Одним из решений этой проблемы является определение нескольких доступных типов PDP-контекста для абонента, например, IPv4v6 и IPv4. В ответ на запрос на роуминг из гостевой сети для обеспечения максимальной совместимости оператор передаст профиль абонента только с типом PDP-контекста IPv4. В то же время в домашней сети абонент может пользоваться более эффективным PDP-контекстом — IPv4v6. Оператор может использовать и так называемые белые листы — списки гостевых сетей, для которых поддержка определенных PDP-контекстов заведомо известна.

Другой класс проблем связан с назначением IP-адреса мобильному устройству при маршрутизации с местным выходом. В этом случае возможно несоответствие между запрашиваемым и доступным PDP-контекстами, между возможностями приложений и подсистем, например IMS, и созданными типами каналов, а также между функциями устройства и соответствующими функциями сети, например при использовании модели 4G/LTE, описанной в предыдущем разделе. Решениями данных проблем является запрещение варианта маршрутизации с местным выходом или адаптация профиля абонента в соответствии с возможностями гостевой сети. В последнем случае могут также применяться белые списки.

Вопросы безопасности, связанные с IPv6

IPv6 является относительно новым протоколом, во многом отличным от своего предшественника — IPv4. Соответственно, с внедрением IPv6 связаны дополнительные риски.

Приведем такой пример. Изначально спецификация IPv6 определяла заголовок расширений маршрутизации (Routing EH), а также его подтип Routing Header Type 0, или RHo. Поле RHo может содержать множество адресов промежуточных узлов, через которые должна пройти передача пакета, причем один и тот же адрес может быть указан более одного раза. А значит, есть вероятность того, что пакеты будут осциллировать между двумя узлами, тем самым вызывая перегрузку канала между ними. Эта возможность может быть использована атакующим для создания атаки отказа в обслуживании (Denial of Service, DoS) с усилением. Поэтому в 2007 году IETF исключил данную функциональность из спецификации IPv6.²³

В целом, риски можно разделить на следующие категории:

Риски, связанные с недостаточной подготовленностью персонала

Очевидно, что недостаточный опыт и подготовка персонала в обнаружении и решении проблем, связанных с IPv6, а также недостаточно хорошее понимание различных новых функций являются существенными рисками безопасности.

Риски, связанные с неадекватной политикой безопасности в отношении IPv6

Многие организации по-прежнему рассматривают протокол IPv6 как экспериментальный, даже если его внедрение происходит в рабочей инфраструктуре. Как следствие, политика безопасности зачастую разрабатывается и исполняется менее строго. Это усугубляется тем, что во многих случаях механическая репликация существующей политики для IPv4 невозможна — это связано как с различиями в семантике IPv6, так и с возможностями оборудования.

²³ RFC 5095: Deprecation of Type 0 Routing Headers in IPv6,
URL: <https://www.rfc-editor.org/rfc/rfc5095>

Один из примеров — неэффективность использования экранов безопасности без сохранения состояния (stateless firewalls). Дело в том, что заголовок IPv6 не содержит поля, указывающего на протокол верхнего уровня, например, TCP. В IPv6 протокол верхнего уровня определяется последним заголовком расширений — заголовком верхнего уровня (Upper Layer Header, ULH). Экраны безопасности обычно содержат правила, основанные как на информации интернет-уровня (IP), так и на протоколах верхнего уровня (например, TCP). Поскольку для фрагментов информация о последнем будет отсутствовать, возникает неопределенность в обработке таких пакетов экраном. Но даже если пакет доставлен целиком, для получения информации о протоколе верхнего уровня устройству потребуется проанализировать все заголовки расширений IPv6, которые в пакете представлены в виде связанного списка. Очевидно, что в некоторых случаях это может привести к значительным дополнительным затратам на обработку, уменьшая производительность устройства.

Риски, связанные с новой функциональностью IPv6 и новыми векторами атаки

Выше мы упомянули обнаруженную уязвимость протокола при применении специальных расширений, связанных с маршрутизацией пакета. Но атакующий может использовать и малтикаст-адреса пакетов для сканирования локальной сети на предмет присутствующих там устройств. Были обнаружены новые векторы атак, связанные и с системой автоконфигурации. Подробное обсуждение этих проблем можно найти в следующих документах IETF²⁴.

Уже упомянутые заголовки расширений ввиду их многообразия и сложности разбора также несут существенные риски, начиная от меньшей эффективности устройств безопасности и заканчивая возможностью их использования в качестве атакующего средства. Более подробно эта проблематика рассмотрена в документе IETF²⁵.

Текущее состояние внедрения и использования IPv6

С момента публикации спецификации IPv6 прошло уже больше 25 лет, но пока рано говорить о полном замещении протокола IPv4 протоколом IPv6. В то же время тенденция в целом положительная. Например, процент пользователей, использующих IPv6 для доступа к услугам Google, в феврале 2024 года составил 40-45%. Стоит заметить, что 10 лет назад эта цифра составляла меньше 3% (см. рис. 14).

²⁴ RFC 3756: IPv6 Neighbor Discovery (ND) Trust Models and Threats,
URL: <https://www.rfc-editor.org/rfc/rfc3756>;
RFC 6104: Rogue IPv6 Router Advertisements,
URL: <https://www.rfc-editor.org/rfc/rfc6104>

²⁵ RFC 9098: Operational Implications of IPv6 Packets with Extension Headers,
URL: <https://www.rfc-editor.org/rfc/rfc9098>

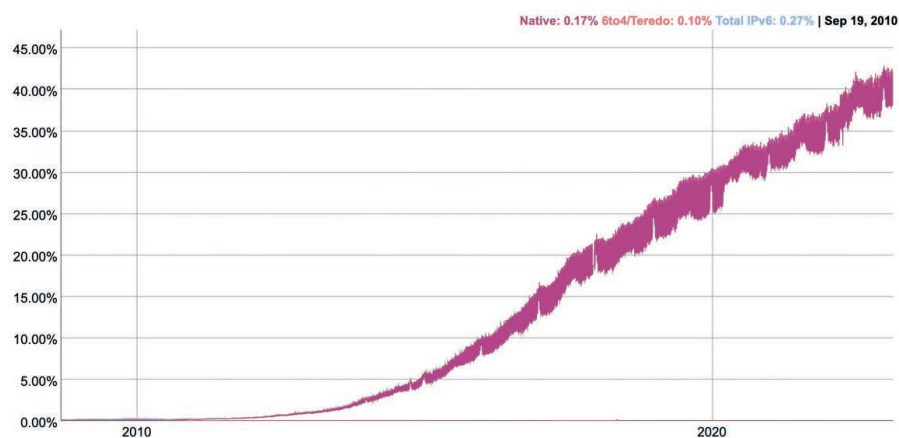


Рис. 14. Рост процента пользователей, использующих IPv6.

Источник: <https://www.google.com/intl/en/ipv6/statistics.html>

Поставщики контента немного отстают — процент веб-сайтов из списка 1000 наиболее популярных сайтов, доступных по IPv6, достиг почти 30%.

Однако приводить конкретные цифры в книге не имеет особого смысла — эти данные довольно быстро устареют. Вместо этого давайте рассмотрим основные аспекты готовности и соответствующие информационные ресурсы, где вы сможете найти актуальные данные.

Готовность инфраструктуры

Первый индикатор, который приходит в голову, — конечно, распределение адресного пространства IPv6. Эта цифра дает представление о максимальном числе провайдеров, внедряющих IPv6 в свою инфраструктуру, поскольку наличие адресного пространства является необходимым, но недостаточным условием его использования. Эти данные доступны на сайте NRO²⁶.

Более точным индикатором внедрения является процент сетей, анонсирующих адресное пространство IPv6 в Интернет. Эта информация доступна на сайте RIPE²⁷.

Информационные ресурсы

В конечном итоге пользователям не важно, какой протокол сетевого уровня они используют для доступа к ресурсам Интернета. Но если необходимый ресурс доступен по протоколу IPv6, а сеть провайдера и конечное оборудование (операционная система) пользователя поддерживают IPv6, то, скорее всего, для доступа будет использован именно этот протокол. Если одно из этих условий не выполняется, по-прежнему работу выполнит протокол IPv4.

²⁶ <https://www.nro.net/statistics>

²⁷ <https://www.ripe.net/analyse/statistics/?tags=ipv6>

В этом смысле доступность информационных ресурсов по протоколу IPv6 является важным индикатором готовности Интернета к новому протоколу.

Сайт World IPv6 Launch²⁸ до июня 2022 года отслеживал процент от 1000 самых популярных веб-сайтов, доступных по IPv6 (с июня 2022 года сайт перестал обновляться). Этот сайт также содержит список других индикаторов и измерений степени внедрения IPv6.

Фактическое использование IPv6

Наконец, для получения полной картины необходимо взглянуть на использование протокола конечными пользователями.

Одним из наиболее популярных графиков использования IPv6 пользователями Интернета является статистика Google²⁹. Это неудивительно, ведь колоссальное число пользователей услуг Google и YouTube позволяет составить очень реалистичную картину.

Интересно также взглянуть на измерения исследователей APNIC, исследующих процент пользователей, которые могут использовать IPv6³⁰.

Сопутствующим индикатором является объем трафика, передаваемого по протоколу IPv6 в Интернете. Здесь стоит посмотреть на статистику точек обмена трафиком (например, AMS-IX³¹) или сетей распределения контента CDN (например, Akamai³²).

Глобальная система администрирования адресного пространства

Присвоением числовых идентификаторов также занимается Джон. Если вы разрабатываете протокол или приложение, которые предполагают использование идентификатора линка, сокета, порта, протокола или сети, пожалуйста, обратитесь к Джону за присвоением числового идентификатора.

RFC 790³³, «Присвоенные номера», 1981 г.

²⁸ <https://www.worldipv6launch.org/measurements>

²⁹ <https://www.google.com/intl/en/ipv6/statistics.html>

³⁰ <https://stats.labs.apnic.net/ipv6>

³¹ <https://stats.ams-ix.net/sflow/ipv6.html>

³² <https://www.akamai.com/internet-station/cyber-attacks/state-of-the-internet-report/ipv6-adoption-visualization>

³³ RFC 790: ASSIGNED NUMBERS, URL: <https://www.rfc-editor.org/rfc/rfc790>

Итак, в рамках модели IP каждое устройство, а точнее, сетевой интерфейс каждого устройства, подключенного к Интернету, имеет уникальный IP-адрес. Для обеспечения уникальности присвоения IP-адресов необходима система учета и распределения адресных ресурсов — система администрирования адресного пространства. Начиная разговор об администрировании адресного пространства, стоит заметить, что в конце 1970-х гг. Интернет полностью относился к министерству обороны США, к тем университетам и научным центрам, которые вели работы в рамках DARPA. Более широкое подключение университетов к ARPANET и создание научно-образовательных сетей общего назначения (CSNET, а затем NSFNET) началось только в 80-х гг. прошлого столетия.

Неудивительно, что организации, обеспечивавшие координацию Интернета, осуществляли эти функции по контрактам с министерством обороны США.

Начиная с ARPANET за распределение различных цифровых идентификаторов, включая доменные имена верхнего уровня, параметры протоколов, IP-адреса и номера автономных систем, отвечала организация IANA (Internet Assigned Numbers Authority, в переводе — Администрация присвоенных номеров Интернета). Ее функции до 1999 года выполнял Институт информатики (Information Sciences Institute) Университета Южной Калифорнии (USC). Фактически же распределение адресов и номеров автономных систем было делегировано сетевому информационному центру DDN-NIC компании SRI International, который обслуживал так называемую интернет-регистрацию, или IP (Internet Registry, IR).

Конец 80-х гг. прошлого века ознаменовался быстрым развитием компьютерных сетей, основанных на протоколе IP. И не только в США, но и за их пределами, особенно в Европе. Надо сказать, что в то время IP в Европе был своего рода гадким утенком. Существующие телефонные компании, монополисты в своей стране, продвигали сети коммутации пакетов, основанные на сетевой модели OSI (Open System Interconnection — взаимодействие открытых систем) и связанной с ней системе протоколов. Эти протоколы, большинство из которых кануло в Лету, были документированы в тщательно разработанных стандартах, основывались на семиуровневой модели (от физического до уровня приложений) и позволяли сетям и приложениям различных операторов взаимодействовать друг с другом. Все это было хорошо, за исключением того, что работа носила теоретический характер, пыталась предвидеть и решить все будущие потребности пользователей и делала это в рамках существующих телекоммуникационных моделей.

С другой стороны, практические требования существующих пользователей сетей передачи данных, в основном университетов и научно-исследовательских центров, были достаточно просты: предоставьте нам канал передачи данных за разумную цену, а с протоколами мы сами разберемся. В большинстве случаев для передачи данных использовалась более простая система TCP/IP,

хорошо зарекомендовавшая себя в научно-исследовательских сетях США. Понятно, что такие запросы не встречали радушного отклика со стороны телекоммуникационных компаний и организаций, отвечающих за стандартизацию, таких как Международная организация по стандартизации (ИСО) и Международный союз электросвязи (МСЭ).

Можно сказать, что развитие академических сетей в Европе во многом проходило под знаком борьбы демократичного и прагматичного TCP/IP с жесткой и дорогостоящей системой OSI. По своему характеру этот процесс очевидным образом соотносился с либерализацией, происходившей в Европе: окончание холодной войны и падение железного занавеса, движения за независимость и демократию. В это время создаются различные организации и сообщества для финансирования и координации развития академических сетей. Некоторые из них имеют французские названия с английскими акронимами — дань моде того времени. Например, RARE — Réseaux Associés pour la Recherche Européenne (Европейское сообщество научно-исследовательских сетей). Или RIPE — Réseaux IP Européens (Европейские IP-сети).

Интернационализация Интернета

В ответ на стремительное развитие и интернационализацию Интернета в августе 1990 года IAB (в то время Internet Activities Board — Совет, отвечающий за «определение технического направления создания стандартов и разрешение проблем в Интернете»,³⁴ опубликовал документ за авторством Винта Серфа (Vint Cerf) «Рекомендуемая политика распределения адресных идентификаторов Интернета» — RFC 1174³⁵.

Суть предлагаемых изменений:

- возможность делегирования распределения адресного пространства и номеров автономных систем другим организациям. Предполагалось, что эти организации будут утверждены Комитетом CCIRN (Coordinating Committee for Intercontinental Research Networking) — группой, созданной в 1988 году для координации глобального развития Интернета, изначально преимущественно между США и Европой. В рамках этой модели предполагалось, что IP сохранит свою центральную функцию и продолжит осуществлять надзор над распределением адресного пространства во всем мире;
- предложение о прекращении использования статуса «подключенный» («connected») при распределении блоков адресов.

³⁴ RFC 1118: The Hitchhikers Guide to the Internet, URL: <https://www.rfc-editor.org/rfc/rfc1118>

³⁵ RFC 1174: IAB Recommended Policy on Distributing Internet Identifier Assignment and IAB Recommended Policy Change to Internet "Connected" Status, URL: <https://www.rfc-editor.org/rfc/rfc1174>

На последнем изменении стоит остановиться подробнее.

Суть статуса «подключенный» заключалась в разрешении обмена трафиком с сетью NSFNET — опорной научно-образовательной сетью США, которая, собственно, и являлась Интернетом. Поэтому быть «подключенным» к NSFNET означало быть подключенным к Интернету.

Проблема заключалась в том, что поддержка работы NSFNET осуществлялась на государственные деньги США. Изначально подключиться к NSFNET могли только государственные организации США или организации, спонсором которых являлось какое-либо американское государственное агентство.

Однако распространение сетевых технологий привело к тому, что сети стали создаваться самыми разнообразными организациями, включая негосударственные и коммерческие. К тому же правительство США старалось ограничить бюджетное финансирование подключения к NSFNET, соответственно, стимулируя подключенные сети к переходу в режим самоокупаемости. А наиболее очевидной стратегией в данном случае было обслуживание в том числе и коммерческих организаций.

Таким образом, в практике «подключения» появились нюансы — важным стало не то, какой тип организации к какой сети подключен, а то, какого типа трафиком обмениваются эти сети. Например, считалось недопустимым использование бюджетных сетей для обмена коммерческим трафиком, однако обмен трафиком между коммерческой организацией и университетом в рамках научно-исследовательского проекта был вполне возможен.

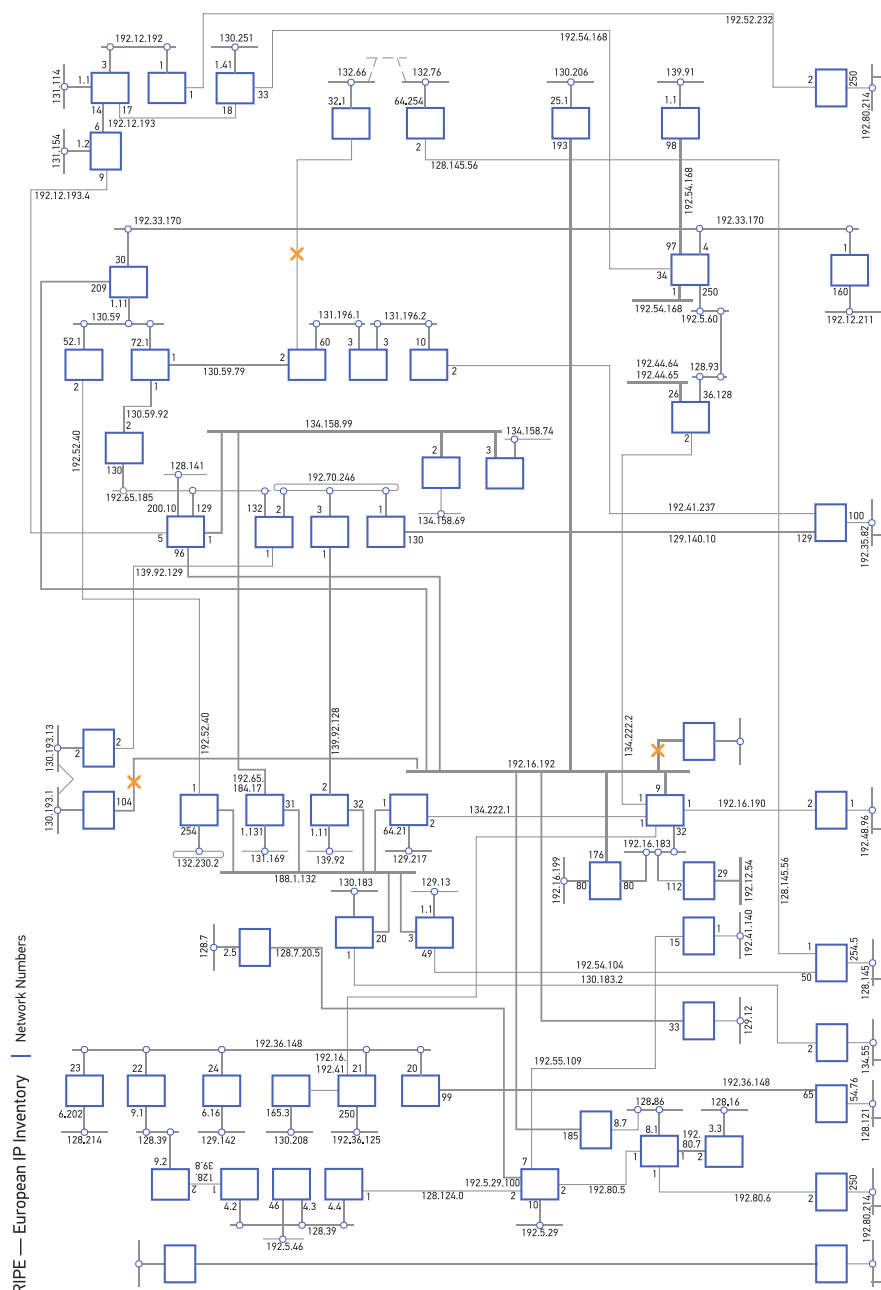
Единовременно присуждаемый статус более не являлся работающим критерием, и требовался более тонкий контроль за соблюдением так называемых правил, или политики допустимого использования (AUP — Acceptable Use Policy). Вместо бинарного статуса IAB предлагал сбор информации о политике использования подключенных сетей. Контроль за соответствием политики маршрутизации политике использования оставался за NSFNET, которая вела соответствующую базу данных (Policy Routing Database, PRDB, позже трансформированную в регистратуру маршрутизации RADB).

Краткая история политики распределения адресных ресурсов в Европейском регионе

В том же августе 1990 года на рассмотрение участников 6-го совещания RIPE было вынесено предложение о создании Сетевого координационного центра RIPE (RIPE NCC) со следующими задачами:

- создание и обслуживание Европейской IP-регистратуры в рамках архитектуры, предложенной в RFC 1174;³⁵
- информационное обслуживание сетей;
- административная поддержка RIPE.

For sure you will find something we missed to include into this map.
Please send corrections and additions to the given address.



Источник: <ftp://ftp.ripe.net/ripe/docs/ripe-026.ps>

Назначение Даниела Карренберга (Daniel Karrenberg) генеральным директором RIPE NCC, а также размещение центра в нидерландском Национальном институте ядерной физики и физики высоких энергий NIKHEF было утверждено в январе 1992 года советом RARE — организации, которая обеспечивала юридическую платформу для нового координационного центра. В апреле того же года RIPE NCC был формально представлен участникам XII совещания RIPE. Центр арендовал две комнаты в NIKHEF, а его штат насчитывал три человека.

Одной из задач RIPE NCC являлось создание «делегированной» интернет-регистратуры в рамках концепции, предложенной в RFC 1174.³⁵ Уже в сентябре 1992 г. RIPE NCC начал обслуживать запросы на получение адресного пространства от европейских организаций.

Двумя месяцами раньше, в июле, был опубликован документ (RIPE NCC Internet Numbers Registration Procedures, *ripe-065*³⁶), который можно считать первой политикой распределения ресурсов. Данная процедура регистрации адресов содержала ряд требований:

- адресное пространство выделялось только сервис-провайдерам, которые, в свою очередь, распределяли его индивидуальным организациям;
- распределенное адресное пространство подлежало регистрации в базе данных RIPE;
- дополнительное обоснование требовалось для запроса сети класса В (это было время «классовых» сетей, хотя уже применялась концепция супернетов — объединение нескольких последовательных адресных блоков класса С в сеть большего размера, — для распределения адресных ресурсов, технология CIDR еще не была внедрена);
- для сетей класса С устанавливалось требование поддержки супернетов, а именно — требование резервирования сервис-провайдерами соседних блоков класса С с целью возможности их последующего объединения.

Процедура являлась достаточно простой и неформальной. Немногим позже она была слегка формализована: появилась форма запроса, в которой, помимо контактных данных, заявитель предоставлял сведения о размере сети и перспективах ее роста на ближайшие два года. Обсуждение этих документов происходило в рамках образованной рабочей группы Local IR.

В то время разработка системы и политики распределения адресных ресурсов происходила в рамках IETF. Помимо упомянутого основополагающего документа RFC 1174³⁵, в мае 1993 года вышел более полный документ «Руководство по управлению адресным пространством» — RFC 1466.³⁷ Этот документ, который активно обсуждался и был согласован с RIPE, явился основой политики распределения адресных ресурсов в Европе.

³⁶ <https://www.ripe.net/publications/docs/ripe-065>

³⁷ RFC 1466: Guidelines for Management of IP Address Space,
URL: <https://www.rfc-editor.org/rfc/rfc1466>

Совместно с системой CIDR эти документы обозначили становление новой иерархической системы распределения адресного пространства, верхний уровень которой базировался на геополитическом принципе, а последующие — на системе взаимоотношений провайдер–клиент. Последняя предусматривала, что адресное пространство, полученное сервис-провайдером от региональной интернет-регистратуры (РИР), будет далее иерархически распределено этим провайдером его клиентам и т.д. Последствием внедрения такой системы явилась зависимость адресации сетей клиентов от конкретного провайдера и требование переадресации в случае, когда сеть меняет своего провайдера.

Существовала, правда, возможность получить непосредственно от РИР адресное пространство, которое являлось независимым от провайдера. Такие ресурсы так и назывались — независимые от провайдера (Provider Independent, PI). Все остальные получили название «агрегируемые» (Provider Aggregatable, PA). Независимые ресурсы являлись более «дорогостоящими» для Интернета, поскольку не могли быть агрегированы в большие блоки и тем самым оказывали большую нагрузку на систему маршрутизации. Надо сказать, что в то время ресурсы маршрутизаторов были достаточно ограничены и рост маршрутизационных таблиц представлял серьезную проблему. Было даже предложено остановить выдачу ресурсов PI, однако RIPE выступил против, опасаясь последствий в виде усиления регулирующих функций RIPE NCC. В ответ было решено усилить разъяснительную работу среди запрашивающих ресурсы PI, акцентируя внимание на то, что данные ресурсы могут иметь проблемы с глобальной маршрутизацией.

Летом 1995 года концепция «агрегируемых» (Provider Aggregatable) и «независимых» от провайдера (Provider Independent) адресных ресурсов была включена в политику распределения адресных ресурсов и документирована в RIPE-127³⁸.

В это же время перед Интернетом стояла еще одна проблема — назревающая нехватка адресного пространства. Хотя внедрение CIDR отсрочило опустошение пула свободных адресов, внимание общественности было привлечено к более бережливому распределению конечного ресурса. Надо отметить, что решение задачи сохранения адресных ресурсов усугубляет проблему роста таблиц маршрутизации — и наоборот. Дело в том, что для сохранения адресного пространства желательно выделять как можно меньшие блоки адресов, минимизируя резервирование, в то время как для оптимальной маршрутизации важно, чтобы адресные блоки сервис-провайдера были максимально агрегируемы. Вопрос верного баланса между этими противоречащими друг другу целями впервые появился на повестке дня XXII конференции RIPE (январь 1996 года) и с тех пор является ключевым в обсуждении различных правил и параметров распределения. Осенью 1996 года были опубликованы два документа. Один, традиционно подготовленный в рамках IETF и опубликованный под номером

³⁸ <ftp://ftp.ripe.net/ripe/docs/ripe-127.txt>

RFC 2050,³⁹ назывался «Руководство по распределению IP-адресов интернет-регистратурой». По существу, данный RFC документировал существовавшую до этого времени практику распределения, которая являлась основой политик RIPE. Второй документ назывался «Политика и процедуры европейской интернет-регистратуры» (European Internet Registry Policies and Procedures, RIPE-140⁴⁰) и являлся обобщением принципов, правил и процедур, связанных с распределением адресного пространства RIPE NCC. Этот документ стал плодом многомесячного обсуждения в рамках рабочей группы Local IR и обозначил образование независимой региональной политики RIPE. С этого момента процесс разработки политик, связанных с распределением адресных ресурсов, происходит в списках рассылки рабочей группы Local IR.

Принципы распределения адресного пространства

Документ RIPE-140 представил основные принципы распределения адресного пространства.

Уникальность — основополагающее требование для глобальной системы распределения адресов. Каждый присвоенный адрес должен быть уникальным в глобальной сети Интернет.

Агрегируемость — иерархическое распределение адресов, позволяющее оптимизировать глобальную систему маршрутизации. Распределение адресных ресурсов, учитывающее топологию сети и взаимоотношения провайдер–клиент, позволяет оптимизировать маршруты и, как следствие, уменьшить нагрузку на глобальную систему маршрутизации.

Сохранение — распределение ресурсов «по потребностям», минимизация неиспользуемых запасов.

Регистрация — регистрация распределенных и присвоенных адресов в общедоступной базе данных для поддержки уникальности и решения сетевых проблем на любом уровне.

К этим принципам добавился принцип «Справедливости» — все политики и практики, связанные с использованием пространства публичного оповещения, должны справедливо и равноправно применяться ко всем существующим и потенциальным членам интернет-сообщества, независимо от их местонахождения, национальности, размера или любого другого фактора.

Принцип «Сохранения» для адресного пространства IPv4 особого смысла не имеет и поэтому из текущей политики⁴¹ исключен.

³⁹ RFC 2050: Internet Registry IP Allocation Guidelines,
URL: <https://www.rfc-editor.org/rfc/rfc2050>

⁴⁰ <ftp://ftp.ripe.net/ripe/docs/ripe-140.txt>

⁴¹ RIPE-733, URL: <https://www.ripe.net/publications/docs/ripe-733>

Хотя для адресного пространства IPv6 принцип «Сохранения» кажется менее важным, чем в свое время в отношении IPv4, политика распределения⁴² стремится обеспечить оптимальный баланс между противоречащими принципами агрегируемости и сохранения. По существу этих параметров три:

- **Минимальный размер распределяемого пространства.** Чем больше минимальный размер, тем выше вероятность неиспользуемых запасов. Для адресного пространства IPv6 этот параметр равен /32.
- **Временной интервал планирования** для демонстрации потребности в ресурсах. Чем больше этот интервал, тем больше последовательного адресного пространства может получить сервис-провайдер, тем меньше различных блоков необходимо анонсировать провайдеру, тем меньше записей в глобальной таблице маршрутизации. Для адресов IPv6 этот параметр определен менее четко, чем в случае IPv4 (для IPv4 он был равен 12 месяцам), в силу значительно большего размера адресного пространства, получаемого изначально. В принципе, горизонт планирования составляет два года.
- **Процент использования распределенных ресурсов (утилизация).** Для оценки утилизации адресного пространства IPv6 используется коэффициент HD-ratio (Host-Density Ratio), документированный в RFC 3194⁴³. Для получения последующего блока адресов провайдер должен продемонстрировать утилизацию HD-ratio не менее 0,94.

Современная система

Вслед за созданием координационного центра RIPE NCC и признанием его в качестве региональной регистратуры в 1994 года последовало официальное признание IANA второй регистратуры — APNIC, отвечавшей за распределение адресного пространства и номеров автономных систем в Азиатско-Тихоокеанском регионе.

Создание третьей региональной интернет-регистратуры, ARIN, было связано с процессом приватизации Интернета в США. Изначально интернет-регистратура обслуживалась исследовательским институтом SRI International. Учитывая стремительное развитие Интернета, в том числе и в коммерческом секторе, правительство США и NSF приняли решение об изменении существующей структуры ИР, финансируемой министерством обороны. Начиная с 1991 года обслуживание ИР, которая теперь стала называться InterNIC и занималась в том числе и распределением доменных имен, переходит к небольшой (в то время) компании Network Solutions Incorporated (NSI).

Региональные регистратуры RIPE NCC и APNIC эффективно обслуживали соответствующие сообщества операторов в соответствии с ими же разработанными политиками. Этот успех явился одним из факторов разделения функции распределения доменных имен и создания в США отдельной регистратуры для номеровных

⁴² RIPE-738, URL: <https://www.ripe.net/publications/docs/ripe-738>

⁴³ RFC 3194: The Host-Density Ratio for Address Assignment Efficiency: An update on the H ratio, URL: <https://www.rfc-editor.org/rfc/rfc3194>

ресурсов. В декабре 1997 года была создана Американская регистратура интернет-номеров ARIN.

В результате мир был разделен на три сферы обслуживания. Европейские страны, страны бывшего СССР, Ближнего Востока, а также Северной Африки обслуживались RIPE NCC; Азиатско-Тихоокеанский регион обслуживался APNIC, офис которого к тому времени переместился из Японии в Австралию, г. Брисбен. Наконец регистратура ARIN обслуживала Северную Америку и все остальные регионы — страны Латинской Америки, Карибского бассейна и южной части Африки.

Когда в начале XXI века стало очевидно, что нужна отдельная организация LACNIC (Латиноамериканская и Карибская региональная регистратура) для обслуживания латиноамериканских сетей, Интернет уже давно вырос из научно-образовательной сети США и стал глобальной сетью, обладающей громадным экономическим и социальным потенциалом. Формализованы были процессы принятия важных решений, в них участвовало международное сообщество. Процесс создания новых региональных регистратур был установлен в документе ICP-2 «Критерии создания новых региональных интернет-регистратур»⁴⁴, созданном в процессе консультаций между RIPE NCC, APNIC, ARIN и ICANN, а также соответствующими сообществами. Одними из важнейших критериев являлись нейтралитет, техническая экспертиза и широкая поддержка регионального сообщества сетевых операторов.

После годового пробного срока 7 ноября 2002 года IANA официально объявила о признании LACNIC четвертой РИР. Часть зоны обслуживания ARIN отошла к новой РИР.

За LACNIC последовало создание AfriNIC (в апреле 2005 года) под чью ответственность попал весь Африканский континент, до этого обслуживаемый ARIN и RIPE NCC.

В октябре 2003 года APNIC, ARIN, LACNIC и RIPE NCC заключили соглашение об образовании Организации номерных ресурсов⁴⁵. В 2005 году к ним присоединился AfriNIC.

Задачи NRO — обеспечение единого интерфейса взаимодействия системы РИР с внешним миром, а также координация различных совместных проектов регистратур.

Руководством деятельности NRO занимается исполнительный комитет, состоящий из пяти руководителей РИР (формально члены исполкома назначаются советом каждого РИР).

NRO участвует в разработке политик через Номерной комитет (NRO Number Council) — но, по существу, это независимая от NRO структура, тождественная Совету поддерживающей организации по адресации ICANN (Address Supporting

⁴⁴ <https://www.icann.org/resources/pages/new-rirs-criteria-2012-02-25-en>

⁴⁵ Number Resource Organization, NRO, URL: <http://www.nro.net>

Organization, ASO). Каждое из региональных сообществ RIPE выбирает двух членов для включения в состав Номерного комитета (и, соответственно, в Совет ASO).

Процесс разработки политики распределения номерных ресурсов Интернета

Хотя RIPE-140 и определил основные принципы и правила распределения номерных ресурсов (адресных блоков IPv4 и номеров автономных систем), жизнь вносила свои коррективы, возникали новые требования. До некоторого времени процесс разработки новых правил распределения адресов носил достаточно неформальный характер. Работа над определенной политикой была открытой, и любой желающий мог принять в ней участие. Как правило, процесс начинался с подготовки проекта в свободном формате. Он описывал проблему — например, отсутствие правил для определенной ситуации — и выдвигал предложение по ее решению.

В результате последующего обсуждения предложение либо браковалось (например, из-за отсутствия интереса к представленной проблеме), либо дорабатывалось и, возможно, принималось.

Хотя процесс не был формализован, он основывался на трех принципах:

- **Открытость.** Любой желающий может предложить политику и участвовать в ее обсуждении. При этом в расчет принимается аргументация, а не ранг участника или его работодателя.
- **Прозрачность.** Обсуждения и результаты обсуждений документированы и свободно доступны каждому.
- **Консенсус.** Решения принимаются на основе консенсуса.

Обсуждение предложения в основном происходило в соответствующих списках рассылки, а конференции RIPE использовались для более интерактивного обмена мнениями. Зачастую предварительное «зондирование почвы» на предмет интереса к предполагаемой проблеме происходило также на конференции RIPE.

Такая система прекрасно работала, пока сообщество RIPE было небольшим и однородным. Однако по мере роста сообщества усиливались требования по структуризации и формализации процесса. Например, не все предложения и обсуждения носят одинаковый характер. Часть из них является просто обсуждением рекомендаций и технических практик. Другие же, наоборот, ставят своей целью разработку требований или запроса к RIPE NCC. Наконец, ряд дискуссий в действительности затрагивают принципиальные моменты, являющиеся частью существующей или новой политики RIPE. Не все требуют одинакового «обхождения» и форумов для обсуждения.

Еще одним недостатком существовавшего процесса являлось отсутствие четких сроков подготовки и обсуждения предложения. В результате процесс мог либо затянуться на неопределенное время, либо закончиться неожиданно быстро для тех, кто не особенно внимательно следил за его развитием.

Хотя все обсуждения и решения документировались (архивы списков рассылки, протоколы совещаний на конференциях RIPE, документы RIPE), отсутствовала документация самого процесса и его фаз. В результате отслеживание разработки политик порой требовало значительных усилий.

В сентябре 2005 года был опубликован документ «Процесс разработки политики в RIPE» (Policy Development Process in RIPE) под авторством Роба Блокзайла (Rob Blokzijl), описывающий основные принципы и элементы этого процесса. С тех пор документ прошел через несколько ревизий, но его основа сохранилась прежней.⁴⁶

По существу, этот документ — сжатая инструкция для разработчика политики. За отправную точку взят следующий принцип: единственным требованием для участия в процессе разработки политики, или ПРП — от изначальной идеи, обсуждения предложения до согласованной политики — является доступ к электронной почте и подписка на список рассылки соответствующей рабочей группы, в рамках которой предполагается обсуждение предложения. Для участия не обязательно быть членом RIPE NCC или участником конференций RIPE. Однако важно иметь предложение, которое решало бы насущную проблему.

Таким предложением может быть изменение или дополнение в существующую политику либо совершенно новая политика RIPE. Если такое предложение имеется, необходимо соответствующим образом оформить его (структура проекта политики приведена в документе). Проект политики обычно представляется через председателя соответствующей рабочей группы. В случае, когда неочевидно, к какой рабочей группе относится предложение, проект можно представить через председателя RIPE по адресу policyproposal@ripe.net. Кстати, RIPE NCC может оказать помощь в подготовке проекта.

После подачи проект должен пройти три стадии до возможного утверждения в качестве новой политики. Диаграмма на рис. 16 показывает эти стадии на временной оси.

Сначала предложение проходит стадию **обсуждения** (Discussion phase). Эта фаза начинается с анонсирования нового предложения. Для этого используется список policy-announce@ripe.net. Тогда же сообщается, в какой рабочей группе будет происходить обсуждение. Председатель рабочей группы устанавливает продолжительность периода обсуждения, это как минимум четыре недели. По окончании данной стадии предложение может перейти в следующую стадию, отправлено на доработку с повторным обсуждением или вообще исключено. Это зависит от характера замечаний и комментариев и решается по согласованию между подателем предложения и председателем рабочей группы.

Если проект переходит в следующую фазу — **рецензирования** (Review phase), — то в течение четырех недель необходимо подготовить начальную версию документа RIPE — официального способа публикации утвержденных политик.

⁴⁶ Текущая версия — <https://www.ripe.net/publications/docs/ripe-642>

Фаза	Действие	Неделя
Проект	Предложение новой политики	0
Фаза рассмотрения проекта (9 недель)	Изначальное обсуждение проекта (как минимум четыре недели)	1
		2
		3
		4
	Перейти к документации?	5
	Документ оформлен и опубликован	6
		7
		8
		9
Фаза рецензирования (5 недель)	Комментарии и рецензирование (не более четырех недель)	10
		11
		12
		13
	Консенсус?	14
Заключительная фаза (5 недель)	Последний звонок (точно четыре недели)	15
		16
		17
		18
	Консенсус?	19
	Объявление решения	

Рис. 16. Схема процесса разработки политики (ПП) RIPE.

Источник: <https://www.ripe.net/publications/docs/ripe-642>

В стадии рецензирования работа идет над фактическим текстом политики — таким, каким он появится в окончательном варианте. Максимально отведенное время — также четыре недели. По прошествии этого времени председатель решает, был ли достигнут консенсус в отношении проекта. В случае положительного решения проект переходит в **заключительную** стадию (Conclusion phase). В противном случае председатель может полностью исключить проект, отправить на доработку предложения или текст проекта (в зависимости от этого процесс начинается либо со стадии обсуждения, либо рецензирования).

Заключительная стадия начинается с объявления «последнего звонка» (Last Call) в списках рабочей группы и policy-announce@ripe.net. В течение последующих четырех недель любой член сообщества может направить свои комментарии.

Цель — дать возможность противникам политики обозначить свою позицию, если они по каким-либо причинам не сделали это на предыдущих стадиях. По прошествии четырех недель председатели всех рабочих групп совместно решают, был ли достигнут консенсус.

В случае положительного решения объявляется новая действующая политика. Если, по мнению председателей, консенсус не был достигнут, председатели могут исключить проект или направить его на повторное обсуждение.

Все текущие предложения представлены на сайте RIPE⁴⁷. Там же можно узнать, в какой стадии находится то или иное предложение, дату окончания дискуссий и рабочую группу, в которой ведется обсуждение.

В ПРП существует несколько моментов, когда принимается решение о дальнейшей судьбе предложения. Возможно, что участники обсуждения не согласны с решением председателя. Для таких случаев предусмотрена специальная процедура разрешения споров. Надо отметить, что к рассмотрению принимаются претензии только по решению о достижении или отсутствии консенсуса. Разногласия относительно самого предложения, его технических и прочих качеств должны разрешаться в ходе самого обсуждения.

Возможен вариант, при котором — с точки зрения участника работы над предложением — его взгляды не были приняты к рассмотрению в стадии обсуждения. Или решение о достижении консенсуса по окончании стадии рецензирования кажется неправомерным. В таком случае участник должен попытаться разрешить этот вопрос с председателем. Если это невозможно, задача разрешения спора решается председателями всех рабочих групп совместно путем голосования. Председатель, вовлеченный в спор, от голосования воздерживается. Это решение председателей является окончательным.

Другой пример спорной ситуации — несогласие с решением о достижении консенсуса в заключительной стадии, которое совместно приняли председатели всех рабочих групп. В этом случае высшей инстанцией разрешения спора является председатель RIPE. Опять же, его решение в этом споре является окончательным.

Решение об исключении не означает, что предложение похоронено навсегда. В любой момент процесс может быть начат заново — бывает, что новую жизнь обретает даже то же самое предложение. Например, если обнаружили новые убедительные аргументы в его пользу или возникло новое предложение, основанное на изначальной идее.

Наиболее важные политики

За годы существования RIPE сообществом были разработаны десятки различных политик, начиная с общих политик распределения номерных ресурсов (адресов и номеров автономных систем) до специальных случаев и рекомендаций. Дюжина документов регламентирует современный процесс распределения номерных ресурсов в RIPE. Все эти документы доступны на сайте RIPE⁴⁸.

Текущие политики распределения номерных ресурсов можно условно разделить на несколько категорий: глобальные, общие и специальные региональные политики.

⁴⁷ <https://www.ripe.net/participate/policies/current-proposals/current-policy-proposals>

⁴⁸ <https://www.ripe.net/publications/docs/ripe-policies/ripe-policies>

Глобальные политики

Пример глобальных политик — распределение блоков IPv4, IPv6 и номеров автономных систем от IANA региональным интернет-регистратурам. К глобальным относится политика, достигшая консенсуса во всех регионах (РИРax) и ICANN и требующая участия IANA или какой-либо внешней организации, связанной с ICANN, для ее осуществления. Процесс обсуждения и достижения консенсуса происходит во всех регионах более или менее независимо, следуя своим региональным правилам. Например, в RIPE обсуждение глобальной политики должно следовать РПП. По достижении консенсуса во всех регионах проект политики рассматривается Советом ASO — комитета ICANN по вопросам номерных ресурсов. Если, по мнению Совета ASO, результирующий текст адекватно представляет обсуждавшийся проект, процесс достижения консенсуса не был нарушен и мнения основных заинтересованных сторон были учтены, текст направляется в Совет ICANN для ратификации. При положительном решении новая глобальная политика вступает в действие. Схема процесса показана на рис. 17, а его полное описание можно найти на сайте NRO⁴⁹.

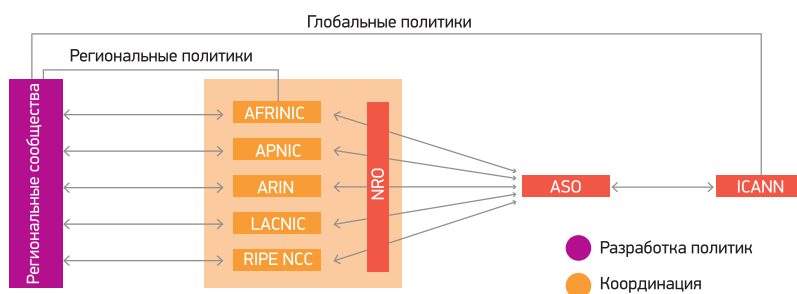


Рис. 17. Процесс разработки и принятия глобальной политики.

Источник: <http://www.nro.net/policies/global-policies-development-process>

Перечислим текущие глобальные политики.

В настоящий момент существуют три глобальные политики, определяющие распределение адресных ресурсов от IANA региональным интернет-регистратурам:

- Политика IANA распределения блоков номеров автономных систем региональным интернет-регистратурам⁵⁰.
- Политика IANA распределения адресных блоков IPv6 региональным интернет-регистратурам⁵¹.

⁴⁹ <https://www.nro.net/global-policy-development-process>

⁵⁰ Internet Assigned Numbers Authority (IANA) Policy for Allocation of ASN Blocks to Regional Internet Registries, <https://www.icann.org/resources/pages/global-policy-asn-blocks-2010-09-21-en>

⁵¹ Internet Assigned Numbers Authority (IANA) Policy for Allocation of IPv6 Blocks to Regional Internet Registries, <https://www.icann.org/resources/pages/allocation-ipv6-rirs-2012-02-25-en>

- Глобальная политика, определяющая механизмы распределения оставшихся адресов IPv4⁵².

Последняя политика заменила глобальную политику распределения адресного пространства IPv4 после опустошения пула свободных адресов IPv4 IANA.

Общие региональные политики распределения IP

Эти политики были утверждены сообществом RIPE и определяют принципы и правила, которым следует RIPE NCC при распределении ресурсов локальным интернет-регистраторам — ЛИРам. К этим политикам относятся:

- политика распределения адресного пространства IPv6⁵³;
- политика распределения адресного пространства IPv4⁵⁴;
- политика распределения номеров автономных систем⁵⁵.

Заметим, что подобные политики существуют во всех регионах. Ежеквартально NRO подготавливает обзорный документ, сравнивающий политики и практики РИРОВ. Эти документы доступны на сайте NRO⁵⁶.

Специальные региональные политики распределения IP

Эти политики охватывают специальные случаи. Например:

- распределение номерных ресурсов для использования самим RIPE NCC — Allocating/Assigning Resources to the RIPE NCC; данная политика определяет, как сам RIPE NCC может запросить и, возможно, получить ресурсы от RIPE NCC;
- требования к договорным отношениям для держателей ресурсов, не зависящих от провайдера — Contractual Requirements for Provider Independent Resource Holders in the RIPE NCC Service Region; данная политика требует обязательного наличия договорных отношений между держателями таких ресурсов и «регистратором» этих ресурсов; в качестве регистратора может выступать либо ЛИР (в большинстве случаев это ЛИР, через которую и были получены данные ресурсы), либо RIPE NCC;
- специальный случай получения ресурсов IPv6 для точек обмена трафиком — IPv6 Address Space Policy for Internet Exchange Points;
- специальный случай получения ресурсов IPv6 для серверов корневой зоны DNS — IPv6 Addresses for Internet Root Servers in the RIPE Region.

⁵² Global Policy for Post Exhaustion IPv4 Allocation Mechanisms by the IANA, <https://www.icann.org/resources/pages/allocation-ipv4-post-exhaustion-2012-05-08-en>

⁵³ IPv6 Address Allocation and Assignment Policy, <https://www.ripe.net/publications/docs/ripe-738>

⁵⁴ IPv4 Address Allocation and Assignment Policies for the RIPE NCC Service Region, <https://www.ripe.net/publications/docs/ripe-733>

⁵⁵ Autonomous System (AS) Number Assignment Policies, <https://www.ripe.net/publications/docs/ripe-679>

⁵⁶ <https://www.ripe.net/publications/docs/ripe-804>

Некоторые общие политики описывают и специальные случаи. Например, выдача адресных блоков операторам доменов верхнего уровня и ENUM в случае использования технологии аникаст (anycast) рассматривается в политике распределения адресного пространства IPv6.

Рекомендации, процедуры и формы

В то время как политики определяют основополагающие принципы и правила, их конкретное воплощение отражается в разработанных RIPE NCC процедурах и формах. Эти документы можно найти на сайте RIPE⁵⁷ в разделе «Request Forms & Supporting Notes».

Заключение

Протокол IP в сетевой модели TCP/IP не напрасно называется уровнем Интернета. Можно сказать, что он приводит к общему знаменателю всю структуру Всемирной сети. Именно на уровне протокола IP взаимодействуют разнородные по своей архитектуре технологии и топологии сети, а на более высоком уровне на плодородной почве IP бурно развиваются и транспортные протоколы, и особенно протоколы приложений. Протокол IP является единственным универсальным требованием для подключения к Интернету. «Подключения» в широком смысле этого слова, ведь, подключаясь к Интернету, сеть или устройство по определению становится частью Интернета, расширяя его связность и функциональность.

IP можно по праву назвать универсальным коннектором. Современные технологии цифровой передачи данных обеспечивают немыслимую ранее пропускную способность, каждый день нас удивляют новейшие приложения для ПК и мобильных устройств — но все эти впечатляющие изменения происходят на уровнях ниже и выше IP. На самом же уровне IP время как будто остановилось. Новая версия протокола IPv6, открывающая новые возможности роста и инноваций, внедряется недостаточно быстро. Это неудивительно: обновление фундамента, да еще в такой самоорганизующейся среде, как Интернет, — задача чрезвычайно сложная, требующая усилий многих заинтересованных сторон, а также огромных затрат времени.

Времени до того момента, когда переход на IPv6 станет не опцией, а единственным возможным путем развития Интернета, к сожалению, остается все меньше.

⁵⁷ <https://www.ripe.net/publications/docs/ripe-documents>

