

多重ルーティング型マルチホームアーキテクチャの提案

宇多 仁[†] 小柏 伸夫[†] 永見 健一^{††} 近藤 邦昭^{††} 中川 郁夫^{††}

篠田 陽一^{†††} 江崎 浩^{††††}

[†] 北陸先端科学技術大学院大学 情報科学研究科, 〒923-1292 石川県能美郡辰口町旭台 1-1.

^{††} (株) インテック・ネットコア, 〒136-0075 東京都江東区新砂 1-3-3.

^{†††} 北陸先端科学技術大学院大学 情報科学センター, 〒923-1292 石川県能美郡辰口町旭台 1-1.

^{††††} 東京大学 情報理工学系研究科, 〒113-8656 東京都文京区本郷 7-3-1.

E-mail: ^{††††}{zin,n-ogashi,shinoda}@jaist.ac.jp, ^{††}{nagami,kuniaki,ikuo}@inetcore.com, ^{††††}hiroshi@wide.ad.jp

あらまし インターネットの発展に伴い、重要度の高いネットワークへの到達性の向上を目指したマルチホーム接続が広く行われるようになった。しかし、現在用いられているマルチホーム方式は、インターネット経路制御に対するインパクトが大きく、さらなるマルチホーム利用者の増加に耐えられない。また、障害時の到達性向上は実現できるが、流入トラフィックの制御が困難であるために、正常時にマルチホームによる複数の接続を有効に使い分けることが困難である。我々は、本論文で、これらの問題点を解決する新たなマルチホームアーキテクチャとして多重ルーティングを用いたマルチホームアーキテクチャを提案する。本方式により、マルチホームの主目的である障害時の到達性の向上を実現しつつ、利用者数に対する対規模性の向上と、流入トラフィックの柔軟な制御を可能とする。

キーワード インターネット, マルチホーム, 多重ルーティング

A new multi-homing architecture based on overlay network

Satoshi UDA[†], Nobuo OGASHIWA[†], Kenichi NAGAMI^{††}, Kuniaki KONDO^{††},

Ikuo NAKAGAWA^{††}, Yoichi SHINODA^{†††}, and Hiroshi ESAKI^{††††}

[†] School of Information Science, Japan Advanced Institute of Science and Technology,
1-1 Asahidai, Tatsunokuchi, Ishikawa 923-1292, Japan.

^{††} Intec NetCore Inc., 1-3-3 Shinsuna, Koto-ku, Tokyo 136-0075, Japan.

^{†††} Center for Information Science, Japan Advanced Institute of Science and Technology,
1-1 Asahidai, Tatsunokuchi, Ishikawa 923-1292, Japan.

^{††††} Graduate School of Information Science and Technology, The University of Tokyo,
7-3-1 Hongo, Bunkyo-ku, Tokyo 113-8656, Japan.

E-mail: ^{††††}{zin,n-ogashi,shinoda}@jaist.ac.jp, ^{††}{nagami,kuniaki,ikuo}@inetcore.com, ^{††††}hiroshi@wide.ad.jp

Abstract As a result of the vast Internet growth, a multi-homing technology is being extensively utilized to keep a connectivity to important networks. However, the traditional multi-homing architecture severely impacts routing performance of the Internet making it non-scalable. Furthermore, although it can enhance a reachability in case of connectivity anomalies, it can not be used to enhance an utilization because it can not assign in-coming traffic to multiple links. In this paper, we propose a new multi-homing architecture which is based on the overlay networking model, providing solutions to these problems. Our architecture improves scalability for increasing multi-homing users and realizes a dynamic and flexible line selection for in-coming traffic, while retaining the original advantage of the multi-homing scheme, which is continuous connectivity to the Internet.

Key words Internet, Multi-Homing, Overlay network.

1. はじめに

インターネットは急激に拡大を遂げ、現在では、情報通信における社会基盤の1つとなっている。このように成長したインターネットにおいては、接続されるネットワークの重要度も様々である。重要度の高いネットワークにおいては、接続回線の障害や輻輳などによるサーバへの到達性の喪失などの問題を回避するために、複数の接続事業者 (ISP) 等と複数の回線で接続するマルチホーム接続 [1], [2] が広く利用されている。

現在広く用いられているマルチホーム接続は、利用者網に対する経路到達性を経路制御システムの作用で冗長化し実現するものが一般的である。この方式では、特に、利用者網宛トラフィックの制御が難しいために、複数の接続をより柔軟かつ効果的に利用したいといった利用者の要求に応えることが困難である。さらには、インターネット上の経路情報数の増加といった大きな問題を抱えており、さらに多くの利用者がマルチホーム接続を利用することは困難である。

本論文では、既存マルチホーム技術の問題点をまとめた上で、指摘した問題点を解決する新たなマルチホームアーキテクチャとして多重ルーティング技術を用いたマルチホーム接続を提案する。そのアーキテクチャに基づいてプロトタイプ実装および運用実験を行なった経験をまとめる。

2. マルチホーム技術の現状と問題点

現在のマルチホーム技術は、利用者網を複数の ISP へ接続し、経路制御システムの作用により利用者網に対する経路到達性を向上させるものが一般的である。マルチホーム接続利用者網は、そのアドレス空間を経路制御プロトコルを用いインターネットへ広告する。該当利用者網宛のパケットは、この経路情報に基づき、いずれかの接続 ISP を介し該当利用者網へ到達する。

マルチホーム接続の増加に伴い、この方式に基づくマルチホーム接続では、次の問題が表面化しつつある。

- インターネット上の経路情報の増加
- 複数接続の柔軟な制御と有効利用の難しさ
- マルチホーム利用者網の運用の難しさ

まず、インターネット上の経路情報の増加に関する問題について述べる。2003 年現在のインターネット上の経路数は 12 万経路を越え、今なお増加を続けている。経路数の増加は、バックボーンルータのメモリ空間を圧迫するため、大きな問題となっている。このような経路数の増加を抑制するため、CIDR (Classless Inter-Domain Routing) を用いたプロバイダ集約が広く用いられている [3], [4]。マルチホーム接続では、複数の接続を介しての利用者網への経路到達性を確保するため、利用者網のアドレス空間を利用者網が接続する全ての ISP から広告する必要がある。利用者網の接続先 ISP が単一 ISP であれば該当 ISP における経路集約が可能であるが、冗長性確保の点で不安が残るため、このような形態は極めて稀である。つまり、多くの場合、各 ISP は、マルチホーム接続する利用者のアドレス空間を経路集約することなくインターネットへ広告する必要がある (図 1)、マルチホーム接続利用者毎にインターネット上の経路数が 1 つずつ

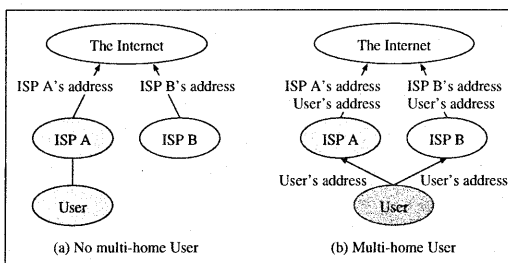


図1 既存マルチホーム接続における経路広告

P.len	Ents	P.len	Ents	P.len	Ents	P.len	Ents
1	0	9	5	17	1,666	25	78
2	0	10	8	18	2,991	26	87
3	0	11	13	19	8,446	27	13
4	0	12	53	20	8,458	28	19
5	0	13	98	21	6,011	29	31
6	0	14	256	22	8,997	30	1
7	0	15	469	23	8,498	31	0
8	21	16	7,400	24	59,378	32	20

表1 インターネット上の経路数 (プレフィックス長毎)

増加することになる。現在では、インターネット上の経路数の半分以上である約 6 万経路を、プレフィックス長 24 ビット以上の小さなアドレス空間の経路が占めるに至っている (表 1)。これは、マルチホームをはじめとする小さなアドレス空間の利用者の経路がインターネットの経路の半分以上を占めていることを示す。

次に、複数接続の有効利用の難しさに関して述べる。既存マルチホーム技術の大きな問題は利用者網へ流入するパケットの制御が難しいことである。マルチホーム利用者網と ISP との間の経路制御には AS (Autonomous System) 間経路制御プロトコルである BGP-4 (Border Gateway Protocol version 4) [5] を用いるが、BGP-4 では、パケット受信側から流入パケットの制御に関する細かな要求を伝えることは難しい。例えば、利用者網への流入トラフィックをマルチホーム接続する複数のリンクに均等に割り振るといった制御は非常に難しい。また、流入パケットの制御が経路制御の枠内での制御に限られると、複数のリンクの使い分けは、パケットの始点と宛先の情報のみをもとにしか行なうことができない。例えば、マルチホーム接続において各接続 ISP の特性に違いがある場合 (一方が広帯域・低品質、もう一方が狭帯域・高品質など)、利用アプリケーション毎にこれらのリンクを使い分けたいといった要求が生じる。既存のマルチホーム技術では、このような要求に応えることは困難である。

最後に、マルチホーム接続利用者網の運用の難しさに関して述べる。BGP-4 による経路制御は、インターネットのバックボーン上で用いられる高機能で複雑な経路制御プロトコルであり、この運用には高い技術を要する。既存マルチホーム技術では、各利用者網においても、上流 ISP との接続に BGP-4 を用いる必要がある。さらに、複数の接続のより有効な利用には、この BGP-4 による経路交換において繊細なパラメータ調整が必

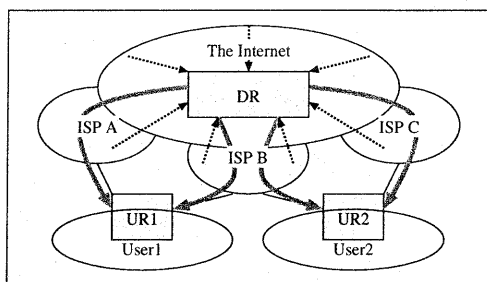


図2 提案マルチホームアーキテクチャの概要

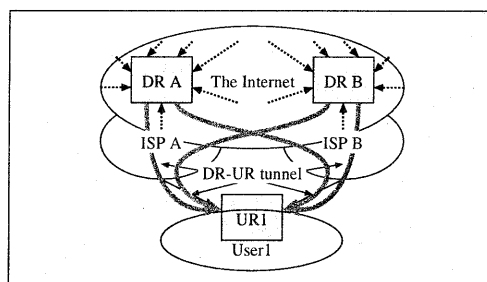


図3 複数のDRを用いた構成

要がある。BGP-4による経路制御を用いたトラフィック制御は、今日でも技術者の勘に頼る部分も大きく、高い技術力と経験が必要である。

3. 多重ルーティング型マルチホーム

前章で述べた既存マルチホーム技術の問題点を解決するために、我々は経路制御技術だけに頼らない、多重ルーティング技術にもとづいたマルチホームアーキテクチャを提案する。

3.1 提案アーキテクチャ概要

既存のマルチホーム方式において、流入トラフィックの制御が困難な点に着目すると、その原因として、利用者網へパケットを送出する側のルータが利用者の制御可能範疇にないことが挙げられる。そこで、提案方式では、トラフィック分散制御装置 (DR) と呼ぶ装置をインターネットバックボーン上に導入し、利用者網宛のトラフィックをDRを介して転送することとした。DRでは、利用者網宛パケットを解析し利用者のポリシーに応じたパケット配送経路選択等を実現する(図2)。

利用者網とDRの間は、マルチホーム接続毎にトンネル等の仮想接続を用いて接続する。また、利用者網側では利用者側終端装置 (UR) と呼ぶ装置を導入し、仮想接続を終端する。例えば、2つのISP(ISP-A, ISP-B)を用いてマルチホーム接続する利用者網は、2種類 (ISP-A 経由およびISP-B 経由) の仮想接続でDRと結ばれることとなる。URにおいてこの仮想接続を終端する端点アドレスには、各接続ISPから割り当てられたアドレスを用いる。つまり、ISP-A側のリンクに割り当てられたアドレス、および、ISP-B側のリンクに割り当てられたアドレスである。このようにISP-AおよびISP-Bからみると利用者網は、それぞれのISPのアドレスのみを用いた通信を行っているように見え、また、インターネット上のノードから見て、利用者網はDRの先に接続されているように見える。

利用者網宛トラフィックは、DRにおいてパケット解析とポリシーに応じたパケット配送経路選択が行われ、利用者側仮想接続のいずれかを用いてURへと転送される。DR-UR間の仮想接続は、利用者網が物理的に接続している接続に対応しており、このDRの選択機能により、利用者網宛各パケットの通るべき物理回線等を選択することが可能である。この経路選択機構においては、一般的な経路選択手法 (つまり宛先アドレスのみにもとづいた経路選択) を越えて、より詳細なパケットの特性ま

で解析した経路選択が可能である。例えば、リアルタイムアプリケーションの通信に対しては低ジッタの経路選択を行なうようなことも可能となる。また、各物理リンクのトラフィック状況を計測し、その結果を回線選択ポリシーにフィードバックすることで、複数のリンクのトラフィックを同等にし回線利用率を高めるような制御も可能である。

もちろん、マルチホーム接続において一部接続の障害時にも利用者網への到達性を維持することは最も重要なことである。本アーキテクチャにおいては、経路選択機構と利用者網を接続する仮想接続の状態を常時監視する。経路選択機構では、障害による仮想接続の切断を検出し、該当仮想接続を利用しない経路選択へと動的に経路選択ルールを変更する。これにより、障害のある経路をパケット転送に用いなくなり、利用者網への接続性は維持される。

3.2 利用者網経路情報の広告

提案方式においては、利用者網が利用するアドレス空間は、DRを運用する事業者に対して割り当てられたCIDRブロックから割り当てる。利用者網の経路情報は、複数の利用者網の経路情報を経路集約した上で広告する。

これは、提案方式では、各利用者網が物理的にどのようなISP等を介して接続されているかにかかわらず、DRの先に同等に接続されているかのように見えるため、DRにおける経路集約が可能となるのである。これにより、既存マルチホーム技術の問題点の1つとして挙げた、経路数の増加による経路制御システムへの影響を大幅に改善することが可能である。

また、利用者網の物理接続先に変更が生じた場合においても、本方式では、仮想接続の利用者網側端点の変更となるのみで、アドレス・リネーミングなどの必要はない。これも、仮想接続の利用者網側端点の変更後も、利用者網は変更前と同等にDRの先に接続されているかのように見えるためである。

3.3 DRの複数化

上述のアーキテクチャにより、利用者網への流入トラフィックの制御は可能となる。また、マルチホーム接続の本来の目的である、一部接続の障害時における利用者網への到達性の維持も可能である。しかし、利用者網宛トラフィックが全てDRを介して配送されるアーキテクチャを採用しているため、DRが単一点障害 (a single point of failure) の要因となる。

この回避のために、複数のDRをインターネット上に分散配

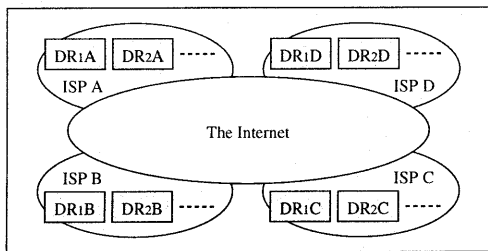


図4 複数のDR群を用いた構成

置することとする。全DRは利用者網のURとの間に仮想接続を確立し、利用者網宛のトラフィックは、発信点から最も近いDRを経由し、仮想接続を通して各利用者網に設置されたURへと配送する。つまり、分散配置したDRは通常時は全DRが利用者網宛のパケット配送を行うこととなる（図3）。このように、全DRが同等なパケットの中継処理等を行うために、DR間で利用者網などに関する情報を常に共有する。

DRの障害時、あるいは、DRの保守作業時には、DRからの利用者網経路情報広告を停止することで、該当DRによる利用者網宛パケットの中継処理を安全に停止することが可能である。該当DRからの経路広告が停止すると、それまで該当DRで処理されていたトラフィックは、経路制御システムの作用により自動的に次候補のDRにおいて処理されることとなる。

また、この構成により、DRの障害に対する耐故障性の向上のみならず、パケット配送遅延の低減と並列処理によるトラフィック集中の回避といった利点も挙げられる。先に述べた通り、提案アーキテクチャでは、利用者網宛トラフィックは常にDRを介して配送されるが、この影響で利用者網に対して直接配送を行った場合に対しての配送遅延が生じる。ここで、複数のDRがインターネット上に分散配置されることで、この配送遅延を小さなものとできる。また、単一DR構成では利用者網宛のトラフィックが全て単一のDRを介して配送されるために、DRの負荷増大が懸念されるが、DRの複数化によりDRでのパケット転送処理を分散化できる。

3.4 対規模性の向上手法

前節で述べたDRの複数設置により、DRでのパケット転送処理を分散化することは可能である。しかし、全DRが各利用者網の仮想接続を保持しているために、利用者数の増加により利用者や仮想接続の管理負荷が増大する。

このような利用者数の増加に対して、前節で述べたDR群を複数用意することで対規模性の向上が可能である。最も単純な手法としては、あらかじめ利用者を複数のグループに分割し、各グループに対応するDR群を設置する。この場合、各DR群は群間で完全に独立して動作可能であり、利用者の増加に対し、利用者グループおよびそれに対応するDR群を追加設置することでより多くの利用者の収容が可能となる。実際には、インターネット上に分散してDR設置箇所を確保し、各箇所利用者グループ数ずつのDRを設置することとなる（図4）。

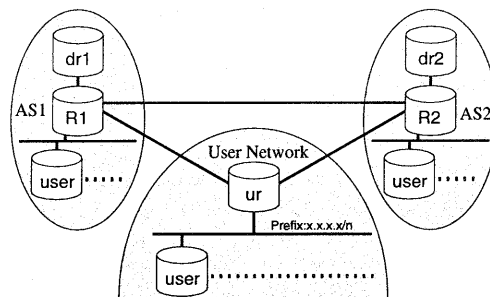


図5 実験ネットワークトポロジ

4. プロトタイプ実装と動作実験

本提案の多重ルーティング型マルチホームアーキテクチャの動作実験のため、プロトタイプ実装を行なった。本プロトタイプは、既存方式と比較し提案方式で新たに実現可能となった機能である利用者網への流入トラフィックの細かな制御の実証を目的とした。本プロトタイプで実現した機能は、

- DR-UR間の仮想接続の確立
- DRでの利用者網宛パケット中継
- DRからの利用者網経路広告
- DRにおける細粒度トラフィック制御

である。

本実装は、URおよびDRの実装から構成される。DR-UR間の仮想接続技術としては、IP over IP トンネリング技術の一つであるGRE (Generic Routing Encapsulation) [6] トンネルを採用した。DRでは、利用者網アドレス空間の経路広告機能、トラフィック制御に用いる仮想接続の管理、複数の仮想へのトラフィック振り振り、仮想接続の状態監視機能を実現する機能を実装した。URでは、仮想接続の確立処理、仮想接続を介したパケットの受信処理を実現する機能を実装した。

URおよびDRの実装は、ソフトウェアアルタとしても広く用いられているNetBSDオペレーティングシステムに対して行なった。この実装は、NetBSD 1.6 RELEASEに対するカーネル拡張とユーザランドアプリケーションの追加から成る。

この実装を用いて、図5に示す実験網を構築し、動作実験を行なった。本実験では、実験網上の任意のノードからの利用者網宛パケットが、提案アーキテクチャの設計通り、DRを介し、仮想接続（ここではGREトンネル）を通り、URへ伝送される基本動作の確認ができた。さらに、DRにおいて、利用者網宛トラフィックの処理ポリシーを設定し、アプリケーション（ポート番号）毎に異なった接続を介してパケットを配送することが可能となっている点も確認した。

また、一部の接続の障害時における障害検出と経路切替に関しても実験を行なった。利用者網の接続のうち実際にトラフィックが流れている側の接続を切断したところ、約200msでそのトラフィックが他方の接続を介しての配送へと自動変更されることを確認した。実アプリケーションにおける実験として、DVによる動画伝送アプリケーションを用いた利用者網への動画配

	既存方式	提案方式
障害時の到達性回復	良	優
経路集約による経路数増加の抑制	困難	可
流入トラフィックの制御	困難	可
利用者網側の運用コスト	高	低
トンネルを用いることによる影響	無	有

表 2 既存方式と提案方式の比較

信中の回線切断実験を行ったが、回線切り替えの影響は目視上ほとんど気にならない範囲内であり、アプリケーションのセッション断なども生じないことを確認した。

さらに、本提案システムの実運用における知見の獲得を目的とし、本実装の実運用網上で運用する広域運用実験を予定している。

5. 考察と課題

5.1 既存方式との比較

本節では、我々の提案する多重ルーティング型マルチホームアーキテクチャの評価として、既存方式に基づくマルチホーム方式との比較を行う。ここでは、マルチホームの主目的である障害時の利用者網への到達性向上効果とともに、2. 章で既存マルチホームの問題点とした挙げた以下の項目に着目し比較した。

- インターネット上の経路数に与える影響
- 複数のリンクの有効な使い分けの容易さ
- 利用者網側での運用の容易さ

まず、マルチホーム接続の主目的である、障害時の対象ネットワークへの到達性の向上効果に関して述べる。既存マルチホームアーキテクチャでは、障害時には、(1) 障害箇所での BGP-4 による経路交換が停止し、(2) 障害箇所を経由する利用者網宛の経路情報が消失し、(3) 他の網を経由する経路が選択され、(4) 利用者網への経路到達性が回復する。ところが、そもそも BGP-4 は頻繁な経路変更を想定したプロトコルではないため、障害を検出した後、経路が安定し到達性が回復するまでにはネットワークの構成にも依るが 10 秒から数分程度の時間を要する。我々の提案方式では、DR-UR 間の生存確認により障害を検出されると、該当仮想接続を指す DR の経路表エントリを非障害トンネル側に変更するのみで、経路切り替えが完了し利用者網への経路到達性が回復する。このため、生存確認間隔にも依存するが、1～数秒程度の時間で経路切り替えが完了し到達性が回復する。

次に、マルチホーム利用者数の増加がインターネット上の経路数に与える影響については、既存方式が 1 ユーザあたり経路数を 1 エントリ増加させるのに対し、提案方式では経路選択機構が経路広告時に複数利用者の経路を集約可能となっており、利用者数の増加が経路数に与える影響は限定的なものとなる。

また、マルチホーム接続時の複数接続の柔軟な制御という点では、BGP-4 の広告コスト調整のみでしか流入トラフィックを制御できない既存方式は、複数のリンクにトラフィックを均等に割り振るといった制御や、アプリケーションの要求に応じた品質のリンクを使用する制御などは不可能である。提案方式で

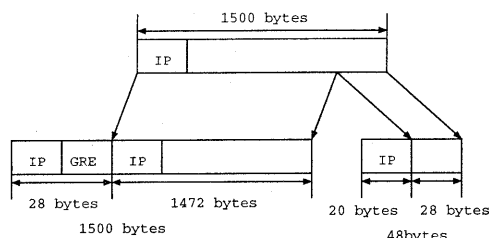


図 6 カプセル化によって生じる IP 断片化

は、経路選択機構による柔軟な経路選択 (ポリルーティング) により利用者の要求に柔軟に応じられ、UR 側トラフィック監視機能との連携により先に挙げたようなトラフィック割り振りなどが実現可能である。

さらに、マルチホーム利用者網側の運用コストに関しては、既存方式に高度な経路制御知識が必要であったのに比べ、提案方式では、経路制御などの複雑な運用は DR の運用者 (マルチホームサービスの提供者) が行うこととなる。利用者側は、UR の設置・運用を行うことになるが、UR では経路制御システム等を動作させる必要はなく、既存方式でのマルチホーム運用に比べ運用は容易となる。

その一方で、提案方式では、利用者網と外部ネットワークとの間のトラフィックがトンネル等の仮想接続を介して転送されるため、トンネルによる転送オーバーヘッド等が問題となる可能性が挙げられる。これに関しては、次節で述べる。

ここまで述べた項目について、表 2 にまとめた。これから、我々の提案手法は、既存手法の問題点を解決しつつ基本的機能である障害時の到達性回復時間の短縮も実現してと言える。

5.2 仮想接続の利用による影響

前節で述べた通り、提案方式によるマルチホーム接続では、トンネル等の仮想接続を介して転送を行うことによる影響が生じる。これは、パケットがトンネル中を配送される際、トンネル自体を示す IP ヘッダとトンネルヘッダが元の IP パケットに付加されることによる。例えば、今回のプロトタイプ実装のようにトンネル方式として GRE トンネルを用いた場合、最低 28 バイト (IP ヘッダ 20 バイトと GRE ヘッダ 8 バイト) が付加される。

インターネット上の各リンクには、そのデータリンクの特性等に応じて MTU (Maximise Transfer Unit) が定められている。ここで、トンネル両端点間の MTU (Maximise Transfer Unit) を 1500 と仮定すると、上記 GRE トンネルによる仮想リンクの MTU は 28 バイトの縮減により 1472 バイトとなる。ところが、実際のインターネット上には、パケットフィルタリング等の設定の誤りから MTU 超過の際の動作を正常に行えないネットワークが数多く存在する。これは大きな問題であるが、MTU 超過処理を正しく行えないネットワークとの通信を維持するために、通常、構築するネットワークの MTU を 1500 未満にしないという運用がなされている。先の GRE トンネルの例では、仮想リンクの MTU として 1500 を維持すると、1472

転送される パケット長	カプセル化前		カプセル化後	
	構成比	(帯域)	構成比	(帯域)
20 - 47	24.15	(2.16)	1.92	(0.13)
48	2.89	(0.27)	20.57	(1.94)
49 - 255	32.23	(5.28)	58.30	(10.66)
256 - 511	7.31	(5.72)	6.24	(4.77)
512 - 1023	6.03	(7.73)	7.42	(9.32)
1024 - 1499	7.29	(19.67)	5.51	(14.60)
1500	20.10	(59.17)	22.06	(64.94)
計	100.00	(100.00)	122.02	(106.36)

表3 カプセル化前後のパケット長分布の比較

バイトを超過するサイズのパケットをカプセル化した場合、トンネルの両端点間ではIP断片化された状態で配送される。この場合、1500バイト長の1つのパケットをトンネルを介して送達する際には、1500バイト長のパケットと48バイト長のパケットの2つのパケットに断片化し送達される(図6)。

表3は、インターネットバックボーン上でトラフィックを24時間観測し、パケット総数と消費帯域をそれぞれ100として、パケット長毎の比率をまとめた。さらに、このトラフィックを提案手法におけるトンネルを介して送った場合、つまり、48バイト長のカプセル化ヘッダを付加し必要に応じてIP断片化処理を行った場合のパケット数および消費帯域を示している。

この結果から、カプセル化処理により、パケット数が約22%、帯域が約6%追加消費されていることが分かる。さらに、パケット長毎のパケット数分布にも変化が見られる。まず、パケット長48バイトのパケットの増加が顕著だが、これはパケット長1500バイトのパケットがカプセル化後に断片化された結果生じたものである。さらに、パケット長1500バイトのパケットの増加、パケット長48バイト未満のパケットの急減が見られる。

このように、提案アーキテクチャがトンネルを用いたインターネットと利用者網の接続を行っていることによりある程度の影響が生じることが観測された。この影響は、パケット数にしてみると20%を超える増加と大きなものであるが、帯域に対する影響は高々6%程度である。この高々6%の帯域増に関しては、提案アーキテクチャが提供するリンクの有効利用とトラフィック制御に関する利点と比較するとその影響は限定的なものであると言える。

この問題に対する改善手法としては、カプセル化時のパケットサイズに対する影響の小さな方式を用いた仮想接続を用いる方法が挙げられる。例えば、DR-UR間の網においてMPLSによるLSPの確立が可能であれば、LSPを用いた接続も利用可能である。さらに、我々は、トンネリングによるMTU縮退の影響を回避するために、MTU値に影響を与えないトンネリング手法の開発を試みている。これは、利用者網のアドレス空間が比較的狭域であることなどを利用し、カプセル化の際、IPヘッダ圧縮技術等を用いることでMTU縮退を防ぐものである。

5.3 DRの設置状況と配送遅延

本アーキテクチャでは、利用者網宛のパケットは、常にDRを経由して配送される。このため、利用者網への最短経路によ

る配送と比較すると、DRへと迂回するためにパケット配送時間の増加が懸念される。しかし、DRはインターネットバックボーン上に分散配置されるため、十分な数のDRを分散配置することで、その送達遅延の増加は軽微なものとできると考える。具体的には、主要IXの近傍や本マルチホーム利用者の多数接続している網内にDRを設置することで、送達遅延に関する問題は殆んど生じないと考える。

そこで、この送達遅延時間の増加問題に関しては、DRの配置状況と遅延時間に関し、予定している実ネットワーク上での広域運用実験を通して評価する予定である。

5.4 セキュリティー

提案アーキテクチャでは利用者網のインターネットとの接続に仮想接続を用いているが、仮想接続の確立時やDRの制御時には、インターネット上での利用者認証が必要である。特に、提案アーキテクチャでは、DRが真正なURによってのみ正しく制御されることが重要である。DRの制御権あるいはDR-UR間の仮想接続が第三者によって奪取されると、該当利用者網全体が第三者の制御下に陥る恐れがある。このため、UR-DR間の制御プロトコルの設計においては、制御の正当性検証などセキュリティの面での考察が重要である。現在、このような考察のもと、DR-UR間の制御プロトコル設計と実装を進めている。

6. まとめ

本論文では、既存マルチホーム技術の問題点をまとめ、その解決として多重ルーティング型マルチホームアーキテクチャを提案した。このアーキテクチャは、インターネット上の経路数の増加といった既存マルチホームの大きな問題を解決するばかりでなく、利用者アプリケーションの要求特性に応じた回線選択といった新たなサービス体系の実現への可能性を持っている。

さらに、特に利用者網への流入トラフィック制御を実現するプロトタイプ実装を行ない、その動作を確認した。また、既存方式との比較により提案方式の有効性を示した上で、広域運用実験へ向けて現在進めている実装の拡張や制御プロトコルの設計に関して議論した。

謝辞 本研究は、総務省 戦略的情報通信研究開発推進制度「国際技術獲得型研究開発」の助成によって行なわれている。関係者各位に感謝する。

文 献

- [1] T. Bates, and Y. Rekhter, RFC 2260: Scalable Support for Multihomed Multi-provider Connectivity, IETF, January 1998.
- [2] J. Hagino, and H. Snyder, RFC 3178: IPv6 Multihoming Support at Site Exit Routers, IETF, October 2001.
- [3] Y. Rekhter, and T. Li, RFC 1518: An Architecture for IP Address Allocation with CIDR, IETF, September 1993.
- [4] V. Fuller, T. Li, J. Yu, and K. Varadhan, RFC 1519: Classless Inter-Domain Routing (CIDR), IETF, September 1993.
- [5] Y. Rekhter, and T. Li, RFC 1771: A Border Gateway Protocol 4 (BGP-4), IETF, March 1995.
- [6] D. Farinacci, T. Li, S. Hanks, D. Meyer, and P. Traina, RFC 2784: Generic Routing Encapsulation (GRE), IETF, March 2000.