



IP ユニキャスト ルーティングの設定

この章では、Catalyst 3750 スイッチに Internet Protocol Version 4 (IPv4) ユニキャスト ルーティングを設定する方法について説明します。特に明記しない限り、スイッチという用語はスタンドアロン スイッチおよびスイッチ スタックを意味します。スイッチ スタックは、ネットワーク内のそれ以外のルータに対して、単一のルータとして動作し、認識されます。スタティック ルーティングおよび Routing Information Protocol (RIP) などの基本的なルーティング機能は、IP ベース イメージと IP サービス イメージの両方で使用できます。先進のルーティング機能およびその他のルーティング プロトコルを使用するには、スタンドアロン スイッチやスタック マスターに IP サービス イメージをインストールする必要があります。



(注)

スイッチ スタックが拡張 IP サービス イメージを実行している場合、IP バージョン 6 (IPv6) ユニキャスト ルーティングもイネーブルにして IPv4 トラフィックに加えて IPv6 トラフィックを転送するようにインターフェイスを設定できます。スイッチに IPv6 を設定する手順については、[第 39 章「IPv6 ユニキャスト ルーティングの設定」](#)を参照してください。

IP ユニキャスト設定情報に関する詳細については、『Cisco IOS IP Configuration Guide, Release 12.2』を参照してください。これには、Cisco.com のホームページ ([Documentation] > [Cisco IOS Software] > [12.2 Mainline] > [Configuration Guides]) からアクセス可能です。このセクションで使用するコマンドの構文および使用方法の詳細については、それぞれのコマンド リファレンスを参照してください。これには、Cisco.com のホームページ ([Documentation] > [Cisco IOS Software] > [12.2 Mainline] > [Command References]) からアクセス可能です。

- 『Cisco IOS IP Command Reference, Volume 1 of 3: Addressing and Services, Release 12.2』
- 『Cisco IOS IP Command Reference, Volume 2 of 3: Routing Protocols, Release 12.2』
- 『Cisco IOS IP Command Reference, Volume 3 of 3: Multicast, Release 12.2』

この章で説明する内容は、次のとおりです。

- 「IP ルーティングの概要」(P.38-2)
- 「ルーティングを設定する手順」(P.38-5)
- 「IP アドレス指定の設定」(P.38-6)
- 「IP ユニキャスト ルーティングのイネーブル化」(P.38-20)
- 「RIP の設定」(P.38-21)
- 「OSPF の設定」(P.38-27)
- 「EIGRP の設定」(P.38-38)
- 「BGP の設定」(P.38-47)
- 「ISO CLNS ルーティングの設定」(P.38-69)

- 「マルチ VRF CE の設定」(P.38-79)
- 「プロトコル独立機能の設定」(P.38-94)
- 「IP ネットワークのモニタおよびメンテナンス」(P.38-110)



(注)

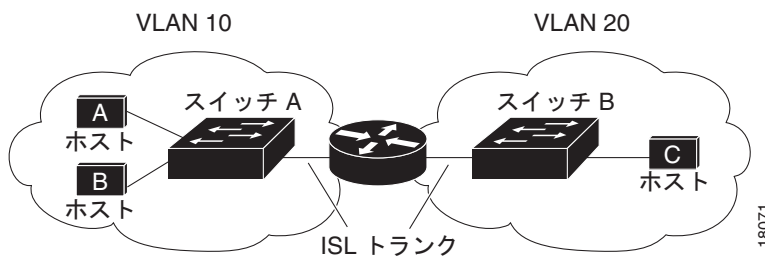
スイッチにルーティング パラメータを設定する場合、使用できるユニキャスト ルート数が最大となるようにシステム リソースを割り当てるには、**sdm prefer routing** グローバル コンフィギュレーション コマンドを使用し、ルーティング テンプレートに Switch Database Management (SDM; スイッチ データベース管理) 機能を設定します。SDM テンプレートの詳細については、第 8 章「SDM テンプレートの設定」、またはこのリリースのコマンド リファレンスの **sdm prefer** コマンドを参照してください。

IP ルーティングの概要

一部のネットワーク環境で、VLAN（仮想 LAN）は各ネットワークまたはサブネットワークに関連付けられています。IP ネットワークで、各サブネットワークは 1 つの VLAN に対応しています。VLAN を設定すると、ブロードキャスト ドメインのサイズを制御し、ローカル トラフィックをローカル内にとどめることができます。ただし、異なる VLAN 内のネットワーク デバイスが相互に通信するには、VLAN 間でトラフィックをルーティング（VLAN 間ルーティング）するレイヤ 3 デバイス（ルータ）が必要です。VLAN 間ルーティングでは、適切な宛先 VLAN にトラフィックをルーティングするため、1 つまたは複数のルータを設定します。

図 38-1 に基本的なルーティング トポロジを示します。スイッチ A は VLAN 10 内、スイッチ B は VLAN20 内にあります。ルータには各 VLAN のインターフェイスが備わっています。

図 38-1 ルーティング トポロジの例



VLAN10 内のホスト A が VLAN10 内のホスト B と通信する場合、ホスト A はホスト B 宛にアドレス指定されたパケットを送信します。スイッチ A はパケットをルータに送信せず、ホスト B に直接転送します。

ホスト A から VLAN20 内のホスト C にパケットを送信する場合、スイッチ A はパケットをルータに転送し、ルータは VLAN10 インターフェイスでトラフィックを受信します。ルータはルーティング テーブルを調べて正しい発信インターフェイスを判別し、VLAN 20 インターフェイスを経由してパケットをスイッチ B に送信します。スイッチ B はパケットを受信し、ホスト C に転送します。

ここでは、ルーティングに関する次の内容について説明します。

- 「ルーティング タイプ」(P.38-3)
- 「IP ルーティングおよびスイッチ スタック」(P.38-3)

ルーティング タイプ

ルータおよびレイヤ 3 スイッチは、次の 3 つの方法でパケットをルーティングできます。

- デフォルト ルーティング
- 事前にプログラミングされているトラフィックのスタティック ルートの使用
- ルーティング プロトコルによるルートの動的な計算

デフォルト ルーティングとは、宛先がルータにとって不明であるトラフィックをデフォルトの出口または宛先に送信することです。

スタティック ユニキャスト ルーティングの場合、パケットは事前に設定されたポートから単一のパスを通り、ネットワークの内部または外部に転送されます。スタティック ルーティングは安全で、帯域幅をほとんど使用しません。ただし、リンク障害などのネットワークの変更には自動的に対応しないため、パケットが宛先に到達しないことがあります。ネットワークが拡大するにつれ、スタティック ルーティングの設定は煩雑になります。

ルータでは、トラフィックを転送する最適ルートを動的に計算するため、ダイナミック ルーティング プロトコルが使用されます。ダイナミック ルーティング プロトコルには次の 2 つのタイプがあります。

- ディスタンス ベクタ プロトコルを使用するルータでは、ネットワーク リソースの距離の値を使用してルーティング テーブルを保持し、これらのテーブルをネイバーに定期的に渡します。ディスタンス ベクタ プロトコルは 1 つまたは複数のメトリックを使用し、最適なルートを計算します。これらのプロトコルは、簡単に設定、使用できます。
- リンク ステート プロトコルを使用するルータでは、ルータ間の Link-State Advertisement (LSA; リンクステート アドバタイズメント) の交換に基づき、ネットワーク トポロジに関する複雑なデータベースを保持します。LSA はネットワークのイベントによって起動され、コンバージェンス時間、またはこれらの変更への対応時間を短縮します。リンクステート プロトコルはトポロジの変更に対応しますが、ディスタンス ベクタ プロトコルよりも多くの帯域幅およびリソースが必要になります。

スイッチでサポートされているディスタンス ベクタ プロトコルは、RIP および Border Gateway Protocol (BGP; ボーダー ゲートウェイ プロトコル) です。RIP は最適パスを決定するために単一の距離メトリック (コスト) を使用し、BGP はパス ベクタ メカニズムを追加します。また、Open Shortest Path First (OSPF) リンクステート プロトコル、および従来の Interior Gateway Routing Protocol (IGRP) にリンクステート ルーティング機能の一部を追加して効率化を図った Enhanced IGRP (EIGRP) もサポートされています。



(注)

スイッチ スタックでサポートされるプロトコルは、スタック マスター上で稼動しているソフトウェアによって決まります。スタック マスター上で IP ベース イメージが稼動している場合は、デフォルトのルーティング、スタティック ルーティング、および RIP だけがサポートされます。その他のすべてのルーティング プロトコルには、IP サービス イメージが必要です。

IP ルーティングおよびスイッチ スタック

スタック内のどのスイッチがルーティング ピアに接続されているかに関係なく、ネットワークは Catalyst 3750 スイッチ スタックを単一ルータとして認識します。スイッチ スタックの動作の詳細については、[第 5 章「スイッチ スタックの管理」](#)を参照してください。

スタック マスターは、次に示す機能を実行します。

- ルーティング プロトコルを初期化し、設定します。
- ルーティング プロトコル メッセージおよびアップデートを他のルータに送信します。

- ・ ピア ルータから受信したルーティング プロトコル メッセージおよびアップデートを処理します。
- ・ **distributed Cisco Express Forwarding (dCEF)** データベースを生成および維持し、すべてのスタック メンバーに配信します。このデータベースに基づいて、スタック内のすべてのスイッチにルートがプログラムされます。
- ・ スタック マスターの MAC (メディア アクセス制御) アドレスはスタック全体のルータ MAC アドレスとして使用され、すべての外部デバイスはこのアドレスを使用して IP パケットをスタックに送信します。
- ・ ソフトウェア転送またはソフトウェア処理を必要とするすべての IP パケットは、スタック マスターの CPU を通ります。

スタック メンバーは、次に示す機能を実行します。

- ・ ルーティング スタンバイ スイッチとして機能します。スタック マスターに障害が発生し、新規スタック マスターとして選択された場合に、処理を引き継ぐことができます。
- ・ ルートをハードウェアにプログラムします。スタック メンバーによってプログラムされたルートは、dCEF データベースの一部としてスタック マスターがダウンロードしたルートと同じです。

スタック マスターに障害が発生すると、スタックはスタック マスターがダウンしていることを検出し、スタック メンバーの 1 つを新規スタック マスターとして選択します。この期間中に、ハードウェアは一時的な中断を除き、アクティブなプロトコルがない状態で、パケットの転送を継続します。

ただし、スイッチ スタックで障害発生後にハードウェア ID を保持していても、スタック マスターがリスタートする前の短時間の中断中に、ルータ ネイバーのルーティング プロトコルがフラップすることもあります。OSPF や EIGRP などのルーティング プロトコルは、ネイバー トランジションを認識する必要があります。ルータでは、2 つのレベルの **Nonstop Forwarding (NSF)** を使用してスイッチオーバーを検出し、ネットワーク トラフィックの転送を継続し、ピア デバイスからのルート情報を回復しています。

- ・ NSF 認識ルータは、ネイバー ルータの障害を許容しています。ネイバー ルータが再起動した後、NSF 認識ルータが要求に応じてステートとルートの隣接関係に関する情報を提供します。
- ・ NSF 対応ルータは NSF をサポートします。スタック マスターの変更を検出すると、NSF 認識ネイバーまたは NSF 対応ネイバーからルーティング情報を再構築し、再起動まで待機しません。

スイッチ スタックは OSPF および EIGRP の NSF 対応ルーティングをサポートしています。詳細については、「[OSPF NSF 機能](#)」(P.38-30) および「[EIGRP NSF 機能](#)」(P.38-41) を参照してください。

新規スタック マスターは、選択されたときに次の機能を実行します。

- ・ ルーティング アップデートの生成、受信、および処理を開始します。
- ・ ルーティング テーブルを構築し、CEF データベースを生成して、スタック メンバーに配信します。
- ・ ルータ MAC アドレスとして自身の MAC アドレスを使用します。新規 MAC アドレスのネットワーク ピアを通知するために、新規ルータ MAC アドレスを使用して **Gratuitous Address Resolution Protocol (ARP)** 応答を定期的に (5 分間、数秒おきに) 送信します。



(注)

固定 MAC アドレス機能をスタックに設定してスタック マスターを変更した場合、スタック MAC アドレスは設定された期間変更されません。この期間に前のスタック マスターがメンバー スイッチとしてスタックに復帰する場合、スタック MAC アドレスは前のスタック マスターの MAC アドレスのままになります。「[固定 MAC アドレスのイネーブル化](#)」(P.5-20) を参照してください。

- ・ ARP 要求をプロキシ ARP IP アドレスに送信し、ARP 応答を受信して、各プロキシ ARP エントリの到達可能性を判別しようとします。到達可能なプロキシ ARP IP アドレスごとに、新規ルータ MAC アドレスを使用して **Gratuitous ARP** 応答を生成します。このプロセスは、新規スタック マスターが選択された後、5 分間繰り返されます。



(注)

スタック マスターで IP サービス イメージが稼働中の場合、スタックは OSPF、EIGRP、BGP などのすべてのサポートするプロトコルを実行することができます。スタック マスターに障害が発生し、新規に選択されたスタック マスター上で IP ベース イメージが稼働している場合、これらのプロトコルはスタック内で稼働しなくなります。



注意

スイッチ スタックを複数のスタックに分割すると、ネットワークが適切に動作しなくなる場合があります。

ルーティングを設定する手順

スイッチ上で、IP ルーティングはデフォルトでディセーブルとなっています。ルーティングを行う前に、IP ルーティングをイネーブルにする必要があります。IP ルーティング設定情報に関する詳細については、『Cisco IOS IP Configuration Guide, Release 12.2』を参照してください。これには、Cisco.com のホームページ ([Documentation] > [Cisco IOS Software] > [12.2 Mainline] > [Configuration Guides]) からアクセス可能です。

次の手順では、次に示すレイヤ 3 インターフェイスの 1 つを指定する必要があります。

- ルーテッド ポート : **no switchport** インターフェイス コンフィギュレーション コマンドを使用し、レイヤ 3 ポートとして設定された物理ポートです。
- Switch Virtual Interface (SVI; スイッチ仮想インターフェイス) : **interface vlan vlan_id** グローバル コンフィギュレーション コマンドによって作成された VLAN インターフェイス。デフォルトではレイヤ 3 インターフェイスです。
- レイヤ 3 モードの EtherChannel ポート チャンネル: **interface port-channel port-channel-number** グローバル コンフィギュレーション コマンドを使用し、イーサネット インターフェイスをチャンネルグループにバインドして作成されたポートチャンネル論理インターフェイスです。詳細については、「[レイヤ 3 EtherChannel の設定](#)」(P.36-16) を参照してください。



(注)

スイッチは、ユニキャスト ルーテッド トラフィックのトンネル インターフェイスをサポートしません。

ルーティングが発生するすべてのレイヤ 3 インターフェイスに、IP アドレスを割り当てる必要があります。「[ネットワーク インターフェイスへの IP アドレスの割り当て](#)」(P.38-7) を参照してください。



(注)

レイヤ 3 スイッチでは、ルーテッド ポートおよび SVI ごとに IP アドレスを 1 つ割り当てることができます。ソフトウェアに、設定できるルーテッド ポートおよび SVI の個数制限はありません。ただし、ハードウェアによって制限されるため、設定できるルーテッド ポートおよび SVI の個数と、実装されている機能の組み合わせによっては、CPU 利用率が影響を受けることがあります。システム メモリをルーティング用に最適化するには、**sdm prefer routing** グローバル コンフィギュレーション コマンドを使用します。

ルーティングを設定するための主な手順は次のとおりです。

- VLAN インターフェイスをサポートするために、スイッチ スタックで VLAN を作成および設定し、レイヤ 2 インターフェイスに VLAN メンバシップを割り当てます。詳細は、[第 13 章「VLAN の設定」](#)を参照してください。
- レイヤ 3 インターフェイスを設定します。

- ・ スイッチ上で IP ルーティングをイネーブルにします。
- ・ レイヤ 3 インターフェイスに IP アドレスを割り当てます。
- ・ 選択したルーティング プロトコルをスイッチ上でイネーブルにします。
- ・ ルーティング プロトコル パラメータを設定します (任意)。

IP アドレス指定の設定

IP ルーティングを設定するには、レイヤ 3 ネットワーク インターフェイスに IP アドレスを割り当ててインターフェイスをイネーブルにし、IP を使用するインターフェイスを経由してホストとの通信を許可する必要があります。ここでは、さまざまな IP アドレス機能の設定方法について説明します。IP アドレスをインターフェイスに割り当てる手順は必須ですが、その他の手順は任意です。

- ・ [「アドレス指定のデフォルト設定」 \(P.38-6\)](#)
- ・ [「ネットワーク インターフェイスへの IP アドレスの割り当て」 \(P.38-7\)](#)
- ・ [「アドレス解決方法の設定」 \(P.38-10\)](#)
- ・ [「IP ルーティングがディセーブルの場合のルーティング支援機能」 \(P.38-13\)](#)
- ・ [「ブロードキャスト パケットの処理方法の設定」 \(P.38-15\)](#)
- ・ [「IP アドレスのモニタおよびメンテナンス」 \(P.38-19\)](#)

アドレス指定のデフォルト設定

[表 38-1](#) に、アドレス指定のデフォルト設定を示します。

表 38-1 アドレス指定のデフォルト設定

機能	デフォルト設定
IP アドレス	未定義。
ARP	Address Resolution Protocol (ARP; アドレス解決プロトコル) キャッシュに永続的なエントリはありません。 カプセル化：標準イーサネット形式の ARP。 タイムアウト：14400 秒 (4 時間)。
IP ブロードキャスト アドレス	255.255.255.255 (すべて 1)。
IP クラスレス ルーティング	イネーブル。
IP デフォルト ゲートウェイ	ディセーブル。
IP 指定ブロードキャスト	ディセーブル (すべての IP 指定ブロードキャストがドロップされます)。
IP ドメイン	ドメイン リスト：ドメイン名は未定義。 ドメイン検索：イネーブル。 ドメイン名：イネーブル。

表 38-1 アドレス指定のデフォルト設定 (続き)

機能	デフォルト設定
IP 転送プロトコル	ヘルパー アドレスが定義されているか、または UDP フラッディングが設定されている場合、デフォルト ポートでは UDP 転送がイネーブルとなります。 ローカル ブロードキャスト：ディセーブル。 Spanning-Tree Protocol (STP; スパニング ツリー プロトコル)：ディセーブル。 ターボフラッディング：ディセーブル。
IP ヘルパー アドレス	ディセーブル。
IP ホスト	ディセーブル。
ICMP Router Discovery Protocol (IRDP)	ディセーブル。 イネーブルの場合のデフォルト： - ブロードキャスト IRDP アドバタイズメント。 - アドバタイズメント間の最大インターバル：600 秒。 - アドバタイズメント間の最小インターバル：最大インターバルの 0.75 倍。 - 初期設定：0。
IP プロキシ ARP	イネーブル。
IP ルーティング	ディセーブル。
IP サブネットゼロ	ディセーブル。

ネットワーク インターフェイスへの IP アドレスの割り当て

IP アドレスは IP パケットの送信先を特定します。一部の IP アドレスは特殊な目的のために予約されていて、ホスト、サブネット、またはネットワーク アドレスには使用できません。RFC 1166『Internet Numbers』には IP アドレスに関する公式の説明が記載されています。

インターフェイスには、1 つのプライマリ IP アドレスを設定できます。マスクは、IP アドレスのネットワーク番号を表すビットを特定します。マスクを使用してネットワークをサブネット化する場合、そのマスクをサブネット マスクと呼びます。割り当てられているネットワーク番号については、インターネット サービス プロバイダーにお問い合わせください。

IP アドレスおよびネットワーク マスクをレイヤ 3 インターフェイスに割り当てるには、特権 EXEC モードで次の手順を実行します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	interface interface-id	インターフェイス コンフィギュレーション モードを開始し、設定するレイヤ 3 インターフェイスを指定します。
ステップ 3	no switchport	レイヤ 2 コンフィギュレーション モードからインターフェイスを削除します (物理インターフェイスの場合)。
ステップ 4	ip address ip-address subnet-mask	IP アドレスおよび IP サブネット マスクを設定します。
ステップ 5	no shutdown	インターフェイスをイネーブルにします。
ステップ 6	end	特権 EXEC モードに戻ります。

	コマンド	目的
ステップ 7	<code>show interfaces [interface-id]</code> <code>show ip interface [interface-id]</code> <code>show running-config interface [interface-id]</code>	設定を確認します。
ステップ 8	<code>copy running-config startup-config</code>	(任意) コンフィギュレーション ファイルに設定を保存します。

サブネット ゼロの使用

サブネット アドレスがゼロであるサブネットを作成しないでください。同じアドレスを持つネットワークおよびサブネットがある場合に問題が発生することがあります。たとえば、ネットワーク 131.108.0.0 のサブネットが 255.255.255.0 の場合、サブネット ゼロは 131.108.0.0 と記述され、ネットワーク アドレスと同じとなってしまいます。

すべてが 1 のサブネット (131.108.255.0) は使用可能です。また、IP アドレス用にサブネット スペース全体が必要な場合は、サブネット ゼロの使用をイネーブルにできます (ただし推奨できません)。

サブネット ゼロをイネーブルにするには、特権 EXEC モードで次の手順を実行します。

	コマンド	目的
ステップ 1	<code>configure terminal</code>	グローバル コンフィギュレーション モードを開始します。
ステップ 2	<code>ip subnet-zero</code>	インターフェイス アドレスおよびルーティングの更新時にサブネット ゼロの使用をイネーブルにします。
ステップ 3	<code>end</code>	特権 EXEC モードに戻ります。
ステップ 4	<code>show running-config</code>	設定を確認します。
ステップ 5	<code>copy running-config startup-config</code>	(任意) コンフィギュレーション ファイルに設定を保存します。

デフォルトに戻して、サブネット ゼロの使用をディセーブルにするには、`no ip subnet-zero` グローバル コンフィギュレーション コマンドを使用します。

クラスレス ルーティング

ルーティングを行うように設定されたスイッチで、クラスレス ルーティング動作はデフォルトでイネーブルとなっています。クラスレス ルーティングがイネーブルの場合、デフォルト ルートがないネットワークのサブネット宛パケットをルータが受信すると、ルータは最適なスーパーネット ルートにパケットを転送します。スーパーネットは、単一の大規模アドレス スペースをシミュレーションするために使用されるクラス C アドレス スペースの連続ブロックで構成されています。スーパーネットは、クラス B アドレス スペースの急速な枯渇を回避するために設計されました。

図 38-2 では、クラスレス ルーティングがイネーブルになっています。ホストがパケットを 128.20.4.1 に送信すると、ルータはパケットを廃棄せずに、最適なスーパーネット ルートに転送します。クラスレス ルーティングがディセーブルの場合、デフォルト ルートがないネットワークのサブネット宛パケットを受信したルータは、パケットを廃棄します。

図 38-2 IP クラスレス ルーティングがイネーブルの場合

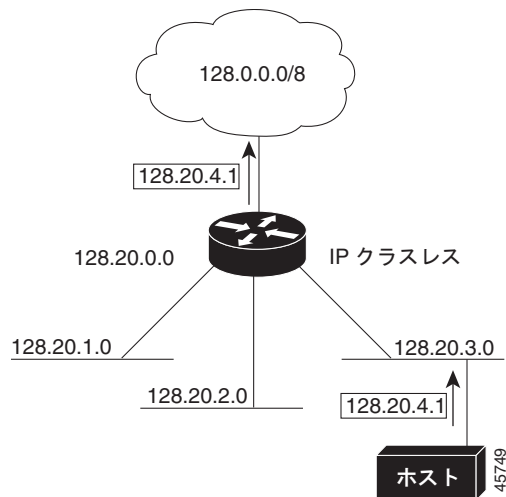
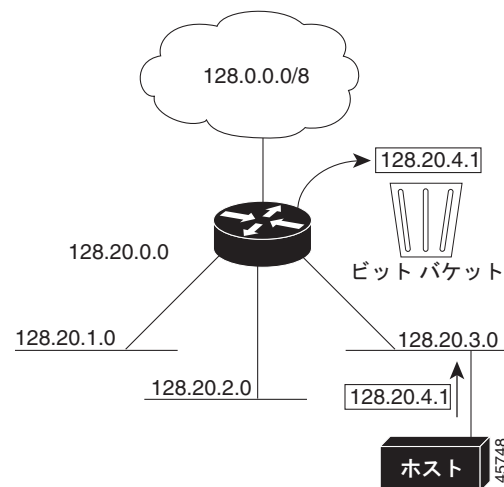


図 38-3 では、ネットワーク 128.20.0.0 のルータはサブネット 128.20.1.0、128.20.2.0、128.20.3.0 に接続されています。ホストがパケットを 128.20.4.1 に送信すると、ネットワークのデフォルト ルートが存在しないため、ルータはパケットを廃棄します。

図 38-3 IP クラスレス ルーティングがディセーブルの場合



認識されないサブネット宛のパケットが最適なスーパーネット ルートに転送されないようにするには、クラスレス ルーティング動作をディセーブルにします。

クラスレス ルーティングをディセーブルにするには、特権 EXEC モードで次の手順を実行します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	no ip classless	クラスレス ルーティング動作をディセーブルにします。
ステップ 3	end	特権 EXEC モードに戻ります。

	コマンド	目的
ステップ 4	show running-config	設定を確認します。
ステップ 5	copy running-config startup-config	(任意) コンフィギュレーション ファイルに設定を保存します。

デフォルトに戻して、デフォルト ルートがないネットワークのサブネット宛パケットが最適なスーパーネット ルートに転送されるようにするには、**ip classless** グローバル コンフィギュレーション コマンドを使用します。

アドレス解決方法の設定

インターフェイス固有の IP 処理方法を制御するには、アドレス解決を行います。IP を使用するデバイスには、ローカル セグメントまたは LAN 上のデバイスを一意に定義するローカル アドレスまたは MAC (メディア アクセス制御) アドレスと、デバイスが属するネットワークを特定するネットワーク アドレスがあります。



(注) Catalyst 3750 スイッチ スタックでは、スタックの単一の MAC アドレスおよび IP アドレスを使用して、ネットワーク通信を行います。

ローカルアドレス (MAC アドレス) は、パケット ヘッダーのデータ リンク層 (レイヤ 2) セクションに格納されて、データ リンク (レイヤ 2) デバイスによって読み取られるため、データ リンク アドレスと呼ばれます。ソフトウェアがイーサネット上のデバイスと通信するには、デバイスの MAC アドレスを学習する必要があります。IP アドレスから MAC アドレスを学習するプロセスを、**アドレス解決**と呼びます。MAC アドレスから IP アドレスを学習するプロセスを、**逆アドレス解決**と呼びます。

スイッチでは、次の形式のアドレス解決を行うことができます。

- **ARP** : IP アドレスを MAC アドレスと関連付けるために使用されます。ARP は IP アドレスを入力と解釈し、対応する MAC アドレスを学習します。次に、IP アドレス/MAC アドレスの関連を ARP キャッシュに格納し、すぐに取り出せるようにします。その後、IP データグラムがリンク レイヤ フレームにカプセル化され、ネットワークを通じて送信されます。イーサネット以外の IEEE 802 ネットワークにおける IP データグラムのカプセル化、および ARP 要求や応答については、Subnetwork Access Protocol (SNAP) で規定されています。
- **プロキシ ARP** : ルーティング テーブルを持たないホストで、他のネットワークまたはサブネット上のホストの MAC アドレスを学習できるようにします。スイッチ (ルータ) が送信元と異なるインターフェイス上のホストに宛てた ARP 要求を受信した場合、そのルータに他のインターフェイスを経由してそのホストに至るすべてのルートが格納されていれば、ルータは自身のローカル データ リンク アドレスを示すプロキシ ARP パケットを生成します。ARP 要求を送信したホストはルータにパケットを送信し、ルータはパケットを目的のホストに転送します。

スイッチでは、ARP と同様の機能 (RARP パケットがローカル MAC アドレスでなく IP アドレスを要求する点を除く) を持つ Reverse Address Resolution Protocol (RARP; 逆アドレス解決プロトコル) を使用することもできます。RARP を使用するには、ルータ インターフェイスと同じネットワーク セグメント上に RARP サーバを設置する必要があります。サーバを識別するには、**ip rarp-server address** インターフェイス コンフィギュレーション コマンドを使用します。

RARP に関する詳細については、『Cisco IOS Configuration Fundamentals Configuration Guide, Release 12.2』を参照してください。これには、Cisco.com のホームページ ([Documentation] > [Cisco IOS Software] > [12.2 Mainline] > [Configuration Guides]) からアクセス可能です。

アドレス解決を設定するために必要な作業は次のとおりです。

- 「[スタティック ARP キャッシュの定義](#)」(P.38-11)

- 「ARP カプセル化の設定」(P.38-12)
- 「プロキシ ARP のイネーブル化」(P.38-12)

スタティック ARP キャッシュの定義

ARP および他のアドレス解決プロトコルを使用すると、IP アドレスと MAC アドレス間を動的にマッピングできます。ほとんどのホストでは動的なアドレス解決がサポートされているため、通常の場合、スタティック ARP キャッシュ エントリを指定する必要はありません。スタティック ARP キャッシュ エントリを定義する必要がある場合は、グローバルに定義できます。グローバルに定義すると、IP アドレスを MAC アドレスに変換するために使用される永続的なエントリを、ARP キャッシュに確保できます。また、指定された IP アドレスがスイッチに属する場合と同じ方法で、スイッチが ARP 要求に応答するように指定することもできます。ARP エントリを永続的なエントリにしない場合は、ARP エントリのタイムアウト期間を指定できます。

IP アドレスと MAC アドレスの間でスタティック マッピングを行うには、特権 EXEC モードで次の手順を実行します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	arp ip-address hardware-address type	ARP キャッシュ内で IP アドレスを MAC (ハードウェア) アドレスにグローバルに関連付け、次に示すカプセル化タイプのいずれかを指定します。 • arpa : ARP カプセル化 (イーサネット インターフェイス用) • snap : SNAP カプセル化 (トークン リングおよび FDDI インターフェイス用) • sap : HP の ARP タイプ
ステップ 3	arp ip-address hardware-address type [alias]	(任意) 指定された IP アドレスがスイッチに属する場合と同じ方法で、スイッチが ARP 要求に応答するように指定します。
ステップ 4	interface interface-id	インターフェイス コンフィギュレーション モードを開始し、設定するインターフェイスを指定します。
ステップ 5	arp timeout seconds	(任意) ARP キャッシュ エントリがキャッシュに保持される期間を設定します。デフォルトは 14400 秒 (4 時間) です。指定できる範囲は 0 ~ 2147483 秒です。
ステップ 6	end	特権 EXEC モードに戻ります。
ステップ 7	show interfaces [interface-id]	すべてのインターフェイスまたは特定のインターフェイスで使用する ARP のタイプおよびタイムアウト値を確認します。
ステップ 8	show arp または show ip arp	ARP キャッシュの内容を表示します。
ステップ 9	copy running-config startup-config	(任意) コンフィギュレーション ファイルに設定を保存します。

ARP キャッシュからエントリを削除するには、**no arp ip-address hardware-address type** グローバル コンフィギュレーション コマンドを使用します。ARP キャッシュから非スタティック エントリをすべて削除するには、**clear arp-cache** 特権 EXEC コマンドを使用します。

ARP カプセル化の設定

IP インターフェイスでは、イーサネット ARP 形式の ARP カプセル化 (**arpa** キーワードで表される) がデフォルトでイネーブルに設定されています。ネットワークの必要性に応じて、カプセル化方法を SNAP に変更できます。

ARP カプセル化タイプを指定するには、特権 EXEC モードで次の手順を実行します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	interface interface-id	インターフェイス コンフィギュレーション モードを開始し、設定するレイヤ 3 インターフェイスを指定します。
ステップ 3	arp {arpa snap}	ARP カプセル化方法を指定します。 • arpa : ARP • snap : SNAP
ステップ 4	end	特権 EXEC モードに戻ります。
ステップ 5	show interfaces [interface-id]	すべてのインターフェイスまたは指定されたインターフェイスの ARP カプセル化設定を確認します。
ステップ 6	copy running-config startup-config	(任意) コンフィギュレーション ファイルに設定を保存します。

カプセル化タイプをディセーブルにするには、**no arp arpa** または **no arp snap** インターフェイス コンフィギュレーション コマンドを使用します。

プロキシ ARP のイネーブル化

デフォルトでは、プロキシ ARP が使用されます。ホストが他のネットワークまたはサブネット上のホストの MAC アドレスを学習できるようにするためです。

ディセーブルになっているプロキシ ARP をイネーブルにするには、特権 EXEC モードで次の手順を実行します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	interface interface-id	インターフェイス コンフィギュレーション モードを開始し、設定するレイヤ 3 インターフェイスを指定します。
ステップ 3	ip proxy-arp	インターフェイスでプロキシ ARP をイネーブルにします。
ステップ 4	end	特権 EXEC モードに戻ります。
ステップ 5	show ip interface [interface-id]	指定されたインターフェイスまたはすべてのインターフェイスの設定を確認します。
ステップ 6	copy running-config startup-config	(任意) コンフィギュレーション ファイルに設定を保存します。

インターフェイスでプロキシ ARP をディセーブルにするには、**no ip proxy-arp** インターフェイス コンフィギュレーション コマンドを使用します。

IP ルーティングがディセーブルの場合のルーティング支援機能

次のメカニズムを使用することで、スイッチは IP ルーティングがイネーブルでない場合、別のネットワークへのルートを取得できます。

- 「プロキシ ARP」(P.38-13)
- 「デフォルト ゲートウェイ」(P.38-13)
- 「ICMP Router Discovery Protocol (IRDP)」(P.38-14)

プロキシ ARP

プロキシ ARP は、他のルートを取得する場合の最も一般的な方法です。プロキシ ARP を使用すると、ルーティング情報を持たないイーサネット ホストと、他のネットワークまたはサブネット上のホストとの通信が可能になります。このホストでは、すべてのホストが同じローカル イーサネット上にあり、ARP を使用して MAC アドレスを学習すると想定されています。送信元と異なるネットワーク上にあるホストに宛てた ARP 要求を受信したスイッチは、そのホストへの最適なルートがあるかどうかを調べます。最適ルートがある場合、スイッチはスイッチ自身のイーサネット MAC アドレスが格納された ARP 応答パケットを送信します。要求の送信元ホストはパケットをスイッチに送信し、スイッチは目的のホストにパケットを転送します。プロキシ ARP は、すべてのネットワークをローカルな場合と同様に処理し、IP アドレスごとに ARP 処理を実行します。

プロキシ ARP は、デフォルトでイネーブルに設定されています。ディセーブル化されたプロキシ ARP をイネーブルにするには、「[プロキシ ARP のイネーブル化](#)」(P.38-12) を参照してください。プロキシ ARP は、他のルータでサポートされている限り有効です。

デフォルト ゲートウェイ

ルートを特定するもう 1 つの方法は、デフォルト ルータ、つまりデフォルト ゲートウェイを定義する方法です。ローカルでないすべてのパケットはこのルータに送信されます。このルータは適切なルーティングを行う、または IP Control Message Protocol (ICMP) リダイレクト メッセージを返信するという方法で、ホストが使用するローカル ルータを定義します。スイッチはリダイレクト メッセージをキャッシュに格納し、各パケットをできるだけ効率的に転送します。この方法には、デフォルト ルータがダウンした場合、または使用できなくなった場合に、検出が不可能となる制限があります。

IP ルーティングがディセーブルの場合にデフォルト ゲートウェイ (ルータ) を定義するには、特権 EXEC モードで次の手順を実行します。

	コマンド	目的
ステップ 1	<code>configure terminal</code>	グローバル コンフィギュレーション モードを開始します。
ステップ 2	<code>ip default-gateway ip-address</code>	デフォルト ゲートウェイ (ルータ) を設定します。
ステップ 3	<code>end</code>	特権 EXEC モードに戻ります。
ステップ 4	<code>show ip redirects</code>	設定を確認するため、デフォルト ゲートウェイ ルータのアドレスを表示します。
ステップ 5	<code>copy running-config startup-config</code>	(任意) コンフィギュレーション ファイルに設定を保存します。

この機能をディセーブルにするには、`no ip default-gateway` グローバル コンフィギュレーション コマンドを使用します。

ICMP Router Discovery Protocol (IRDP)

ルータ ディスカバリを使用すると、スイッチは IRDP を使用し、他のネットワークへのルートを動的に取得します。ホストは IRDP を使用し、ルータを特定します。クライアントとして動作しているスイッチは、ルータ ディスカバリ パケットを生成します。ホストとして動作しているスイッチは、ルータ ディスカバリ パケットを受信します。スイッチは Routing Information Protocol (RIP) ルーティングの更新を受信し、この情報からルータの場所を推測することもできます。実際のところ、ルーティング デバイスによって送信されたルーティング テーブルは、スイッチに格納されません。どのシステムがデータを送信しているのかが記録されるだけです。IRDP を使用する利点は、プライオリティと、パケットが受信されなくなってからデバイスがダウンしていると思なされるまでの期間をルータごとに両方指定できることです。

検出された各デバイスは、デフォルト ルータの候補となります。現在のデフォルト ルータがダウンしたと宣言された場合、または再送信が多すぎて TCP 接続がタイムアウトになりつつある場合、プライオリティが上位のルータが検出されると、最も高いプライオリティを持つ新しいルータが選択されます。

インターフェイスで IRDP ルーティングを行う場合は、インターフェイスで IRDP 処理をイネーブルにしてください。IRDP 処理をイネーブルにすると、デフォルトのパラメータが適用されます。これらのパラメータを変更することもできます。

インターフェイス上で IRDP をイネーブルにして設定するには、特権 EXEC モードで次の手順を実行します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	interface interface-id	インターフェイス コンフィギュレーション モードを開始し、設定するレイヤ 3 インターフェイスを指定します。
ステップ 3	ip irdp	インターフェイス上で IRDP 処理をイネーブルにします。
ステップ 4	ip irdp multicast	(任意) IP ブロードキャストの代わりとして、マルチキャスト アドレス (224.0.0.1) に IRDP アドバタイズメントを送信します。 (注) このコマンドを使用すると、IRDP パケットをマルチキャストとして送信するサン マイクロシステムズ社の Solaris との互換性を維持できます。実装機能の中には、これらのマルチキャストを受信できないものも多くあります。このコマンドを使用する前に、エンドホストがこの機能に対応していることを確認してください。
ステップ 5	ip irdp holdtime seconds	(任意) アドバタイズメントが有効である IRDP 期間を設定します。デフォルトは maxadvertinterval 値の 3 倍です。 maxadvertinterval 値よりも大きな値 (9000 秒以下) を指定する必要があります。 maxadvertinterval 値を変更すると、この値も変更されます。
ステップ 6	ip irdp maxadvertinterval seconds	(任意) アドバタイズメント間の IRDP の最大インターバルを設定します。デフォルト値は 600 秒です。
ステップ 7	ip irdp minadvertinterval seconds	(任意) アドバタイズメント間の IRDP の最小インターバルを設定します。デフォルトは maxadvertinterval 値の 0.75 倍です。 maxadvertinterval を変更すると、この値も新しいデフォルト値 (maxadvertinterval の 0.75 倍) に変更されます。
ステップ 8	ip irdp preference number	(任意) デバイスの IRDP 初期設定レベルを設定します。指定できる範囲は $-2^{31} \sim 2^{31}$ です。デフォルトは 0 です。大きな値を設定すると、ルータの初期設定レベルも高くなります。

	コマンド	目的
ステップ 9	ip irdp address <i>address</i> [<i>number</i>]	(任意) プロキシアドパタイズメントを行うために必要な IRDP アドレスと初期設定を指定します。
ステップ 10	end	特権 EXEC モードに戻ります。
ステップ 11	show ip irdp	IRDP 値を表示し、設定を確認します。
ステップ 12	copy running-config startup-config	(任意) コンフィギュレーション ファイルに設定を保存します。

maxadvertinterval 値を変更すると、**holdtime** 値および **minadvertinterval** 値も変更されます。最初に **maxadvertinterval** 値を変更し、次に **holdtime** 値または **minadvertinterval** 値のいずれかを手動で変更することが重要です。

IRDP ルーティングをディセーブルにするには、**no ip irdp** インターフェイス コンフィギュレーション コマンドを使用します。

ブロードキャスト パケットの処理方法の設定

IP インターフェイス アドレスを設定した後で、ルーティングをイネーブルにしたり、1 つまたは複数のルーティング プロトコルを設定したり、ネットワーク ブロードキャストへのスイッチの応答方法を設定したりできます。ブロードキャストは、物理ネットワーク上のすべてのホスト宛のデータ パケットです。2 種類のブロードキャストがサポートされています。

- 指定ブロードキャスト パケット：特定のネットワークまたは一連のネットワークに送信されます。指定ブロードキャスト アドレスには、ネットワークまたはサブネット フィールドが含まれます。
- フラッドイング ブロードキャスト パケット：すべてのネットワークに送信されます。



(注)

storm-control インターフェイス コンフィギュレーション コマンドを使用して、トラフィック抑制レベルを設定し、レイヤ 2 インターフェイスでブロードキャスト、ユニキャスト、マルチキャスト トラフィックを制限することもできます。詳細は、第 25 章「ポート単位のトラフィック制御の設定」を参照してください。

ルータはローカル ケーブル長を制限して、ブロードキャスト ストームを防ぎます。ブリッジ（インテリジェントなブリッジを含む）はレイヤ 2 デバイスであるため、ブロードキャストはすべてのネットワーク セグメントに転送され、ブロードキャスト ストームが伝播します。ブロードキャスト ストーム問題を解決する最善の方法は、ネットワーク上で単一のブロードキャスト アドレス方式を使用することです。最新の IP 実装機能ではほとんどの場合、アドレスをブロードキャスト アドレスとして使用するように設定できます。スイッチをはじめ、多数の実装機能では、ブロードキャスト メッセージを転送するためのアドレス方式が複数サポートされています。

これらの方式をイネーブルにするには、次に示す作業を実行します。

- 「指定ブロードキャストから物理ブロードキャストへの変換のイネーブル化」(P.38-15)
- 「UDP ブロードキャスト パケットおよびプロトコルの転送」(P.38-17)
- 「IP ブロードキャスト アドレスの確立」(P.38-17)
- 「IP ブロードキャストのフラッドイング」(P.38-18)

指定ブロードキャストから物理ブロードキャストへの変換のイネーブル化

デフォルトでは、IP 指定ブロードキャストがドロップされるため、転送されることはありません。IP 指定ブロードキャストが廃棄されると、ルータが DoS 攻撃にさらされる危険が少なくなります。

ブロードキャストが物理 (MAC レイヤ) ブロードキャストになるインターフェイスでは、IP 指定ブロードキャストの転送をイネーブルにできます。**ip forward-protocol** グローバル コンフィギュレーション コマンドを使用し、設定されたプロトコルだけを転送できます。

転送するブロードキャストを制御するアクセス リストを指定できます。アクセス リストを指定すると、アクセス リストで許可されている IP パケットだけが、指定ブロードキャストから物理ブロードキャストに変換できるようになります。アクセス リストの詳細については、第 34 章「ACL によるネットワーク セキュリティの設定」を参照してください。

インターフェイス上で IP 指定ブロードキャストの転送をイネーブルにするには、特権 EXEC モードで次の手順を実行します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	interface interface-id	インターフェイス コンフィギュレーション モードを開始し、設定するインターフェイスを指定します。
ステップ 3	ip directed-broadcast [access-list-number]	インターフェイス上で、指定ブロードキャストから物理ブロードキャストへの変換をイネーブルにします。転送するブロードキャストを制御するアクセス リストを指定できます。アクセス リストを指定すると、アクセス リストで許可されている IP パケットだけが変換可能になります。 (注) ip directed-broadcast インターフェイス コンフィギュレーション コマンドは VPN Routing/Forwarding (VRF) インターフェイスで設定でき、こうすると VRF 認識になります。指定ブロードキャスト トラフィックが VRF 内だけでルーティングされます。
ステップ 4	exit	グローバル コンフィギュレーション モードに戻ります。
ステップ 5	ip forward-protocol {udp [port] nd sdns}	ブロードキャスト パケットを転送するとき、ルータによって転送されるプロトコルおよびポートを指定します。 • udp : UDP データグラムを転送します。 <i>port</i> : (任意) 転送される UDP サービスを制御する宛先ポートです。 • nd : ND データグラムを転送します。 • sdns : SDNS データグラムを転送します。
ステップ 6	end	特権 EXEC モードに戻ります。
ステップ 7	show ip interface [interface-id] または show running-config	指定されたインターフェイスまたはすべてのインターフェイスの設定を確認します。
ステップ 8	copy running-config startup-config	(任意) コンフィギュレーション ファイルに設定を保存します。

指定ブロードキャストから物理ブロードキャストへの変換をディセーブルにするには、**no ip directed-broadcast** インターフェイス コンフィギュレーション コマンドを使用します。プロトコルまたはポートを削除するには、**no ip forward-protocol** グローバル コンフィギュレーション コマンドを使用します。

UDP ブロードキャスト パケットおよびプロトコルの転送

UDP は IP のホスト間レイヤ プロトコルで、TCP と同様です。UDP はオーバーヘッドが少ない、コネクションレスのセッションを 2 つのエンドシステム間に提供しますが、受信されたデータグラムの確認応答は行いません。場合に応じてネットワーク ホストは UDP ブロードキャストを使用し、アドレス、コンフィギュレーション、名前に関する情報を検索します。このようなホストが、サーバを含まないネットワーク セグメント上にある場合、通常 UDP ブロードキャストは転送されません。この状況を改善するには、特定のクラスのブロードキャストをヘルパー アドレスに転送するように、ルータのインターフェイスを設定します。インターフェイスごとに、複数のヘルパー アドレスを使用できます。

UDP 宛先ポートを指定し、転送される UDP サービスを制御できます。複数の UDP プロトコルを指定することもできます。旧式のディスクレス Sun ワークステーションおよびネットワーク セキュリティ プロトコル SDNS で使用される Network Disk (ND) プロトコルも指定できます。

ヘルパー アドレスがインターフェイスに定義されている場合、デフォルトでは UDP と ND の両方の転送がイネーブルになっています。『Cisco IOS IP Command Reference, Volume 1 of 3: Addressing and Services, Release 12.2』の **ip forward-protocol** インターフェイス コンフィギュレーション コマンドの説明には、UDP ポートを指定しない場合にデフォルトで転送されるポートがリストされています。

UDP ブロードキャストの転送を設定するときに UDP ポートを指定しないと、ルータは BOOTP 転送 エージェントとして動作するように設定されます。BOOTP パケットは Dynamic Host Configuration Protocol (DHCP) 情報を伝達します。

インターフェイスで UDP ブロードキャスト パケットの転送をイネーブルにし、宛先アドレスを指定するには、特権 EXEC モードで次の手順を実行します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	interface interface-id	インターフェイス コンフィギュレーション モードを開始し、設定するレイヤ 3 インターフェイスを指定します。
ステップ 3	ip helper-address address	転送をイネーブルにし、BOOTP などの UDP ブロードキャスト パケットを転送するための宛先アドレスを指定します。
ステップ 4	exit	グローバル コンフィギュレーション モードに戻ります。
ステップ 5	ip forward-protocol {udp [port] nd sdns}	ブロードキャスト パケットを転送するときに、ルータによって転送されるプロトコルを指定します。
ステップ 6	end	特権 EXEC モードに戻ります。
ステップ 7	show ip interface [interface-id] または show running-config	指定されたインターフェイスまたはすべてのインターフェイスの設定を確認します。
ステップ 8	copy running-config startup-config	(任意) コンフィギュレーション ファイルに設定を保存します。

特定アドレスへのブロードキャスト パケットの転送をディセーブルにするには、**no ip helper-address** インターフェイス コンフィギュレーション コマンドを使用します。プロトコルまたはポートを削除するには、**no ip forward-protocol** グローバル コンフィギュレーション コマンドを使用します。

IP ブロードキャスト アドレスの確立

最も一般的な (デフォルトの) IP ブロードキャスト アドレスは、すべて 1 で構成されているアドレスです (255.255.255.255)。ただし、任意の形式の IP ブロードキャスト アドレスを生成するようにスイッチを設定することもできます。

インターフェイス上で IP ブロードキャスト アドレスを設定するには、特権 EXEC モードで次の手順を実行します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	interface interface-id	インターフェイス コンフィギュレーション モードを開始し、設定するインターフェイスを指定します。
ステップ 3	ip broadcast-address ip-address	デフォルト値と異なるブロードキャスト アドレス (128.1.255.255 など) を入力します。
ステップ 4	end	特権 EXEC モードに戻ります。
ステップ 5	show ip interface [interface-id]	指定されたインターフェイスまたはすべてのインターフェイスのブロードキャスト アドレスを確認します。
ステップ 6	copy running-config startup-config	(任意) コンフィギュレーション ファイルに設定を保存します。

デフォルトの IP ブロードキャスト アドレスに戻すには、**no ip broadcast-address** インターフェイス コンフィギュレーション コマンドを使用します。

IP ブロードキャストのフラッディング

IP ブロードキャストをインターネットワーク全体に、制御可能な方法でフラッディングできるようにするには、ブリッジング STP で作成されたデータベースを使用します。この機能を使用すると、ループを回避することもできます。この機能を使用できるようにするには、フラッディングが行われるインターフェイスごとにブリッジングを設定する必要があります。ブリッジングが設定されていないインターフェイス上でも、ブロードキャストを受信できます。ただし、ブリッジングが設定されていないインターフェイスでは、受信したブロードキャストが転送されません。また、異なるインターフェイスで受信されたブロードキャストを送信する場合、このインターフェイスは使用されません。

IP ヘルパー アドレスのメカニズムを使用して単一のネットワーク アドレスに転送されるパケットを、フラッディングできます。各ネットワーク セグメントには、パケットのコピーが 1 つだけ送信されます。

フラッディングを行う場合、パケットは次の条件を満たす必要があります (これらの条件は、IP ヘルパー アドレスを使用してパケットを転送するときの条件と同じです)。

- パケットは MAC レベルのブロードキャストでなければなりません。
- パケットは IP レベルのブロードキャストでなければなりません。
- パケットは Trivial File Transfer Protocol (TFTP; 簡易ファイル転送プロトコル)、Domain Name System (DNS; ドメイン ネーム システム)、Time、NetBIOS、ND、または BOOTP パケット、または **ip forward-protocol udp** グローバル コンフィギュレーション コマンドで指定された UDP でなければなりません。
- パケットの Time To Live (TTL) 値は 2 以上でなければなりません。

フラッディングされた UDP データグラムには、出力インターフェイスで **ip broadcast-address** インターフェイス コンフィギュレーション コマンドによって指定された宛先アドレスを設定します。宛先アドレスを、任意のアドレスに設定できます。このため、データグラムがネットワーク内を伝播するにつれ、宛先アドレスが変更されることもあります。送信元アドレスは変更されません。TTL 値が減ります。

フラッディングされた UDP データグラムがインターフェイスから送信されると (場合によっては宛先アドレスが変更される)、データグラムは通常の IP 出力ルーチンに渡されます。このため、出力インターフェイスにアクセス リストがある場合、データグラムはその影響を受けます。

ブリッジング スパニング ツリー データベースを使用し、UDP データグラムをフラッディングするには、特権 EXEC モードで次の手順を実行します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	ip forward-protocol spanning-tree	ブリッジング スパニング ツリー データベースを使用し、UDP データグラムをフラッディングします。
ステップ 3	end	特権 EXEC モードに戻ります。
ステップ 4	show running-config	設定を確認します。
ステップ 5	copy running-config startup-config	(任意) コンフィギュレーション ファイルに設定を保存します。

IP ブロードキャストのフラッディングをディセーブルにするには、**no ip forward-protocol spanning-tree** グローバル コンフィギュレーション コマンドを使用します。

スイッチでは、パケットの大部分がハードウェアで転送され、スイッチの CPU を経由しません。CPU に送信されるパケットの場合は、ターボフラッディングを使用し、スパニング ツリーベースの UDP フラッディングを約 4 ～ 5 倍高速化します。この機能は、ARP カプセル化用に設定されたイーサネット インターフェイスでサポートされています。

スパニング ツリーベースのフラッディングを向上させるには、特権 EXEC モードで次の手順を実行します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	ip forward-protocol turbo-flood	スパニング ツリー データベースを使用し、UDP データグラムのフラッディングを高速化します。
ステップ 3	end	特権 EXEC モードに戻ります。
ステップ 4	show running-config	設定を確認します。
ステップ 5	copy running-config startup-config	(任意) コンフィギュレーション ファイルに設定を保存します。

この機能をディセーブルにするには、**no ip forward-protocol turbo-flood** グローバル コンフィギュレーション コマンドを使用します。

IP アドレスのモニタおよびメンテナンス

特定のキャッシュ、テーブル、またはデータベースの内容が無効になった場合、または無効である可能性がある場合は、**clear** 特権 EXEC コマンドを使用し、すべての内容を消去できます。表 38-2 に、内容を消去するために使用するコマンドを示します。

表 38-2 キャッシュ、テーブル、データベースを消去するコマンド

コマンド	目的
clear arp-cache	IP ARP キャッシュおよび高速スイッチング キャッシュを消去します。
clear host {name *}	ホスト名およびアドレス キャッシュから 1 つまたはすべてのエントリを削除します。
clear ip route {network [mask] *}	IP ルーティング テーブルから 1 つまたは複数のルート削除します。

IP ルーティング テーブル、キャッシュ、データベースの内容、ノードへの到達可能性、ネットワーク内のパケットのルーティング パスなど、特定の統計情報を表示できます。表 38-3 に、IP 統計情報を表示するために使用する特権 EXEC コマンドを示します。

表 38-3 キャッシュ、テーブル、データベースを表示するコマンド

コマンド	目的
<code>show arp</code>	ARP テーブルのエントリを表示します。
<code>show hosts</code>	デフォルトのドメイン名、検索サービスの方式、サーバ ホスト名、およびキャッシュに格納されているホスト名とアドレスのリストを表示します。
<code>show ip aliases</code>	TCP ポートにマッピングされた IP アドレスを表示します (エイリアス)。
<code>show ip arp</code>	IP ARP キャッシュを表示します。
<code>show ip interface [interface-id]</code>	インターフェイスの IP ステータスを表示します。
<code>show ip irdp</code>	IRDP 値を表示します。
<code>show ip masks address</code>	ネットワーク アドレスに対して使用されるマスクおよび各マスクを使用するサブネット番号を表示します。
<code>show ip redirects</code>	デフォルト ゲートウェイのアドレスを表示します。
<code>show ip route [address [mask]] [protocol]</code>	ルーティング テーブルの現在のステートを表示します。
<code>show ip route summary</code>	ルーティング テーブルの現在のステートをサマリー形式で表示します。

IP ユニキャスト ルーティングのイネーブル化

デフォルトで、スイッチはレイヤ 2 スイッチング モード、IP ルーティングはディセーブルとなっています。スイッチのレイヤ 3 機能を使用するには、IP ルーティングをイネーブルにする必要があります。IP ルーティングをイネーブルにするには、特権 EXEC モードで次の手順を実行します。

	コマンド	目的
ステップ 1	<code>configure terminal</code>	グローバル コンフィギュレーション モードを開始します。
ステップ 2	<code>ip routing</code>	IP ルーティングをイネーブルにします。
ステップ 3	<code>router ip_routing_protocol</code>	IP ルーティング プロトコルを指定します。このステップでは、他のコマンドを実行することもできます。たとえば、 network (RIP) ルータ コンフィギュレーション コマンドを使用し、ルーティングするネットワークを指定できます。具体的なプロトコルの詳細については、この章の後半および『Cisco IOS IP Configuration Guide, Release 12.2』を参照してください。 (注) IP ベース イメージは、ルーティング プロトコルとして RIP だけをサポートします。
ステップ 4	<code>end</code>	特権 EXEC モードに戻ります。
ステップ 5	<code>show running-config</code>	設定を確認します。
ステップ 6	<code>copy running-config startup-config</code>	(任意) コンフィギュレーション ファイルに設定を保存します。

ルーティングをディセーブルにするには、**no ip routing** グローバル コンフィギュレーション コマンドを使用します。
次に、ルーティング プロトコルとして RIP を使用し、IP ルーティングをイネーブルにする例を示します。

```
Switch# configure terminal
Enter configuration commands, one per line. End with CNTL/Z.
Switch(config)# ip routing
Switch(config)# router rip
Switch(config-router)# network 10.0.0.0
Switch(config-router)# end
```

ここで、選択したルーティング プロトコルのパラメータを設定できます。具体的な手順は次のとおりです。

- 「RIP の設定」(P.38-21)
- 「OSPF の設定」(P.38-27)
- 「EIGRP の設定」(P.38-38)
- 「BGP の設定」(P.38-47)
- 「プロトコル独立機能の設定」(P.38-94) (任意)

RIP の設定

RIP は、小規模な同種ネットワーク間で使用するために作成された Interior Gateway Protocol (IGP; 内部ゲートウェイ プロトコル) です。RIP は、ブロードキャスト UDP データ パケットを使用してルーティング情報を交換するディスタンス ベクタ ルーティング プロトコルです。このプロトコルは RFC 1058 に文書化されています。RIP の詳細については、『*IP Routing Fundamentals*』(Cisco Press 刊)を参照してください。



(注)

RIP は IP ベース イメージでサポートされている唯一のルーティング プロトコルです。その他のルーティング プロトコルを使用する場合は、スタック マスター上で IP サービス イメージを稼働させる必要があります。

スイッチは RIP を使用し、30 秒ごとにルーティング情報アップデート (アドバタイズメント) を送信します。180 秒以上を経過しても別のルータからアップデートがルータに届かない場合、該当するルータから送られたルートは使用不能としてマークされます。240 秒が経過してもアップデートが届かない場合、アップデートを行わないルータに関するすべてのルーティング テーブル エントリは削除されます。

RIP では、各ルートの値を評価するためにホップ カウントが使用されます。ホップ カウントは、ルート内で経由されるルータ数です。直接接続されているネットワークのホップ カウントは 0 です。ホップ カウントが 16 のネットワークに到達することはできません。このように範囲 (0 ~ 15) が狭いため、RIP は大規模ネットワークには適していません。

ルータにデフォルトのネットワーク パスが設定されている場合、RIP はルータを疑似ネットワーク 0.0.0.0 にリンクするルートをアドバタイズします。0.0.0.0 ネットワークは存在しません。RIP はデフォルトのルーティング機能を実行するためのネットワークとして、このネットワークを処理します。デフォルト ネットワークが RIP によって取得された場合、またはルータが最終ゲートウェイで、RIP がデフォルト メトリックによって設定されている場合、スイッチはデフォルト ネットワークをアドバタイズします。RIP は指定されたネットワーク内のインターフェイスにアップデートを送信します。インターフェイスのネットワークを指定しないと、RIP アップデート中にアドバタイズされません。

ここでは、次の設定情報について説明します。

- 「RIP のデフォルト設定」(P.38-22)
- 「基本的な RIP パラメータの設定」(P.38-22)
- 「RIP 認証の設定」(P.38-24)

- ・「サマリー アドレスおよびスプリット ホライズンの設定」(P.38-25)

RIP のデフォルト設定

表 38-4 に、RIP のデフォルト設定を示します。

表 38-4 RIP のデフォルト設定

機能	デフォルト設定
自動サマリー	イネーブル。
デフォルト情報送信元	ディセーブル。
デフォルト メトリック	自動メトリック変換（組み込み）。
IP RIP 認証キーチェーン	認証なし。 認証モード：クリア テキスト。
IP RIP 受信バージョン	version ルータ コンフィギュレーション コマンドに準拠。
IP RIP 送信バージョン	version ルータ コンフィギュレーション コマンドに準拠。
IP RIP の起動	version ルータ コンフィギュレーション コマンドに準拠。
IP スプリット ホライズン	メディアにより異なる。
ネイバー	未定義。
ネットワーク	指定なし。
オフセット リスト	ディセーブル。
出力遅延	0 ミリ秒。
タイマー基準	• update : 30 秒。 • invalid : 180 秒。 • holddown : 180 秒。 • flush : 240 秒。
アップデート送信元の検証	イネーブル。
バージョン	RIP バージョン 1 およびバージョン 2 パケットを受信し、バージョン 1 パケットを送信します。

基本的な RIP パラメータの設定

RIP を設定するには、ネットワークに対して RIP ルーティングをイネーブルにします。他のパラメータを設定することもできます。Catalyst 3750 スイッチでは、ネットワーク番号を設定するまで RIP コンフィギュレーション コマンドは無視されます。

RIP をイネーブルにして設定するには、特権 EXEC モードで次の手順を実行します。

	コマンド	目的
ステップ 1	<code>configure terminal</code>	グローバル コンフィギュレーション モードを開始します。
ステップ 2	<code>ip routing</code>	IP ルーティングをイネーブルにします (IP ルーティングがディセーブルになっている場合に限り、必須です)。

	コマンド	目的
ステップ 3	router rip	RIP ルーティング プロセスをイネーブルにし、ルータ コンフィギュレーション モードを開始します。
ステップ 4	network network number	ネットワークを RIP ルーティング プロセスに関連付けます。複数の network コマンドを指定できます。RIP ルーティング アップデートの送受信は、これらのネットワークのインターフェイスを経由する場合に限り可能です。 (注) RIP コマンドを有効にするためにネットワーク番号を設定する必要があります。
ステップ 5	neighbor ip-address	(任意) ルーティング情報を交換するネイバー ルータを定義します。このステップを使用すると、RIP (通常はブロードキャスト プロトコル) からのルーティング アップデートが非ブロードキャスト ネットワークに到達するようになります。
ステップ 6	offset list [access-list number name] {in out} offset [type number]	(任意) オフセット リストをルーティング メトリックに適用し、RIP によって取得したルートへの着信および発信メトリックを増加します。アクセス リストまたはインターフェイスを使用し、オフセット リストを制限できます。
ステップ 7	timers basic update invalid holddown flush	(任意) ルーティング プロトコル タイマーを調整します。すべてのタイマーの有効範囲は 0 ~ 4294967295 秒です。 update : ルーティング アップデートの送信間隔。デフォルト値は 30 秒です。 invalid : ルートが無効と宣言された後の時間。デフォルト値は 180 秒です。 holddown : ルートがルーティング テーブルから削除されるまでの時間。デフォルト値は 180 秒です。 flush : ルーティング アップデートが延期される時間。デフォルト値は 240 秒です。
ステップ 8	version {1 2}	(任意) RIP バージョン 1 または RIP バージョン 2 のパケットだけを送受信するようにスイッチを設定します。デフォルトの場合、スイッチではバージョン 1 および 2 を受信しますが、バージョン 1 だけを送信します。インターフェイス コマンド ip rip {send receive} version 1 2 1 2 を使用し、インターフェイスでの送受信に使用するバージョンを制御することもできます。
ステップ 9	no auto summary	(任意) 自動サマライズをディセーブルにします。デフォルトでは、クラスフル ネットワーク境界を通過するときにサブプレフィクスがサマライズされます。サマライズをディセーブルにし (RIP バージョン 2 に限る)、クラスフル ネットワーク境界にサブネットおよびホスト ルーティング情報をアドバタイズします。
ステップ 10	no validate-update-source	(任意) 着信 RIP ルーティング アップデートの送信元 IP アドレスの検証をディセーブルにします。デフォルトでは、スイッチが着信 RIP ルーティング アップデートの送信元 IP アドレスを検証します。送信元アドレスが無効な場合は、アップデートが廃棄されます。通常の使用の場合は、この機能をディセーブルにしないでください。ただし、ネットワークに接続されていないルータがあり、そのルータのアップデートを受信する場合は、このコマンドを使用できます。

	コマンド	目的
ステップ 11	output-delay <i>delay</i>	(任意) 送信する RIP アップデートにパケット間遅延を追加します。デフォルトでは、複数のパケットからなる RIP アップデートのパケットに、パケット間遅延が追加されません。パケットを低速なデバイスに送信する場合は、8 ～ 50 ミリ秒のパケット間遅延を追加できます。
ステップ 12	end	特権 EXEC モードに戻ります。
ステップ 13	show ip protocols	設定を確認します。
ステップ 14	copy running-config startup-config	(任意) コンフィギュレーション ファイルに設定を保存します。

RIP ルーティング プロセスをオフにするには、**no router rip** グローバル コンフィギュレーション コマンドを使用します。

アクティブなルーティング プロトコル プロセスのパラメータと現在のステータスを表示するには、**show ip protocols** 特権 EXEC コマンドを使用します。RIP データベースのサマリー アドレス エントリを表示するには、**show ip rip database** 特権 EXEC コマンドを使用します。

RIP 認証の設定

RIP バージョン 1 では、認証がサポートされていません。RIP バージョン 2 のパケットを送受信する場合は、インターフェイスで RIP 認証をイネーブルにできます。インターフェイスで使用できる一連の鍵は、キー チェーンによって決まります。キー チェーンが設定されていないと、デフォルトの場合でも認証は実行されません。「[認証鍵の管理](#)」(P.38-109) に記載されている作業も実行してください。

RIP 認証がイネーブルであるインターフェイスでは、プレーン テキストと MD5 という 2 つの認証モードがサポートされています。デフォルトはプレーン テキストです。

インターフェイスに RIP 認証を設定するには、特権 EXEC モードで次の手順を実行します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	interface <i>interface-id</i>	インターフェイス コンフィギュレーション モードを開始し、設定するインターフェイスを指定します。
ステップ 3	ip rip authentication key-chain <i>name-of-chain</i>	RIP 認証をイネーブルにします。
ステップ 4	ip rip authentication mode [<i>text</i> <i>md5</i>]	プレーン テキスト認証 (デフォルト) または MD5 ダイジェスト認証を使用するように、インターフェイスを設定します。
ステップ 5	end	特権 EXEC モードに戻ります。
ステップ 6	show running-config interface [<i>interface-id</i>]	設定を確認します。
ステップ 7	copy running-config startup-config	(任意) コンフィギュレーション ファイルに設定を保存します。

クリア テキスト認証に戻すには、**no ip rip authentication mode** インターフェイス コンフィギュレーション コマンドを使用します。認証を禁止するには、**no ip rip authentication key-chain** インターフェイス コンフィギュレーション コマンドを使用します。

サマリー アドレスおよびスプリット ホライズンの設定

ブロードキャストタイプの IP ネットワークに接続され、ディスタンス ベクタ ルーティング プロトコルを使用するルータでは、通常ルーティング ループの発生を抑えるために、スプリット ホライズン メカニズムが使用されます。スプリット ホライズンは、ルートに関する情報がその情報の発信元であるインターフェイスで、ルータによってアドバタイズされないようにします。この機能を使用すると、通常の場合は複数のルータ間通信が最適化されます（特にリンクが壊れている場合）。



(注) ルートを適切にアドバタイズするため、スプリット ホライズンをディセーブルにすることがアプリケーションに必要な場合を除き、通常はこの機能をディセーブルにしないでください。

ダイヤルアップ クライアント用のネットワーク アクセス サーバで、サマライズされたローカルな IP アドレス プールをアドバタイズするように、RIP が動作しているインターフェイスを設定する場合は、**ip summary-address rip** インターフェイス コンフィギュレーション コマンドを使用します。



(注) スプリット ホライズンがイネーブルの場合、自動サマリーとインターフェイス IP サマリー アドレスはともにアドバタイズされません。

サマライズされたローカル IP アドレスをアドバタイズし、インターフェイスのスプリット ホライズンをディセーブルにするようにインターフェイスを設定するには、特権 EXEC モードで次の手順を実行します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	interface interface-id	インターフェイス コンフィギュレーション モードを開始し、設定するレイヤ 3 インターフェイスを指定します。
ステップ 3	ip address ip-address subnet-mask	IP アドレスおよび IP サブネットを設定します。
ステップ 4	ip summary-address rip ip address ip-network mask	サマライズする IP アドレスおよび IP ネットワーク マスクを設定します。
ステップ 5	no ip split horizon	インターフェイスでスプリット ホライズンをディセーブルにします。
ステップ 6	end	特権 EXEC モードに戻ります。
ステップ 7	show ip interface interface-id	設定を確認します。
ステップ 8	copy running-config startup-config	(任意) コンフィギュレーション ファイルに設定を保存します。

IP サマライズをディセーブルにするには、**no ip summary-address rip** ルータ コンフィギュレーション コマンドを使用します。

次の例では、主要ネットは 10.0.0.0 です。自動サマリー アドレス 10.0.0.0 はサマリー アドレス 10.2.0.0 によって上書きされるため、10.2.0.0 はインターフェイス ギガビット イーサネット ポート 2 からアドバタイズされますが、10.0.0.0 はアドバタイズされません。次の例では、インターフェイスがまだレイヤ 2 モード（デフォルト）の場合、**no switchport** インターフェイス コンフィギュレーション コマンドを入力してから、**ip address** インターフェイス コンフィギュレーション コマンドを入力する必要があります。



(注)

スプリット ホライズンがイネーブルである場合、(**ip summary-address rip** ルータ コンフィギュレーション コマンドによって設定される) 自動サマリーとインターフェイス サマリー アドレスはともにアドバタイズされません。

```
Switch(config)# router rip
Switch(config-router)# interface gigabitethernet1/0/2
Switch(config-if)# ip address 10.1.5.1 255.255.255.0
Switch(config-if)# ip summary-address rip 10.2.0.0 255.255.0.0
Switch(config-if)# no ip split-horizon
Switch(config-if)# exit
Switch(config)# router rip
Switch(config-router)# network 10.0.0.0
Switch(config-router)# neighbor 2.2.2.2 peer-group mygroup
Switch(config-router)# end
```

スプリット ホライズンの設定

ブロードキャストタイプの IP ネットワークに接続され、ディスタンス ベクタ ルーティング プロトコルを使用するルータでは、通常ルーティング ループの発生を抑えるために、スプリット ホライズン メカニズムが使用されます。スプリット ホライズンは、ルートに関する情報がその情報の発信元であるインターフェイスで、ルータによってアドバタイズされないようにします。この機能を使用すると、複数のルータ間通信が最適化されます (特にリンクが壊れている場合)。



(注)

ルートを適切にアドバタイズするために、スプリット ホライズンをディセーブルにすることがアプリケーションに必要である場合を除き、通常この機能をディセーブルにしないでください。

インターフェイスでスプリット ホライズンをディセーブルにするには、特権 EXEC モードで次の手順を実行します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	interface interface-id	インターフェイス コンフィギュレーション モードを開始し、設定するインターフェイスを指定します。
ステップ 3	ip address ip-address subnet-mask	IP アドレスおよび IP サブネットを設定します。
ステップ 4	no ip split-horizon	インターフェイスでスプリット ホライズンをディセーブルにします。
ステップ 5	end	特権 EXEC モードに戻ります。
ステップ 6	show ip interface interface-id	設定を確認します。
ステップ 7	copy running-config startup-config	(任意) コンフィギュレーション ファイルに設定を保存します。

スプリット ホライズン メカニズムをイネーブルにするには、**ip split-horizon** インターフェイス コンフィギュレーション コマンドを使用します。

OSPF の設定

ここでは、OSPF の設定方法について簡単に説明します。OSPF コマンドの詳細については、『*Cisco IOS IP Command Reference, Volume 2 of 3: Routing Protocols, Release 12.2*』を参照してください。これには、Cisco.com ホームページ ([Documentation] > [Cisco IOS Software] > [12.2 Mainline] > [Command References]) からアクセスできます。



(注)

OSPF では、各メディアがブロードキャスト ネットワーク、非ブロードキャスト ネットワーク、ポイントツーポイント ネットワークに分類されます。スイッチでは、ブロードキャスト ネットワーク (イーサネット、トークン リング、FDDI) およびポイントツーポイント ネットワーク (ポイントツーポイント リンクとして設定されたイーサネット インターフェイス) がサポートされます。

OSPF は IP ネットワーク専用の IGP で、IP サブネット化、および外部から取得したルーティング情報のタグ付けをサポートしています。OSPF を使用するとパケット認証も可能になり、パケットを送受信するときに IP マルチキャストが使用されます。シスコの実装機能では、RFC1253 の OSPF MIB (管理情報ベース) がサポートされています。

シスコの実装機能は、次の主要機能を含む OSPF バージョン 2 仕様に準拠します。

- スタブ エリアの定義がサポートされています。
- 任意の IP ルーティング プロトコルによって取得されたルートは、別の IP ルーティング プロトコルに再配信されます。つまり、ドメイン内レベルで、OSPF は EIGRP および RIP によって取得したルートを取り込むことができます。OSPF ルートを RIP に伝達することもできます。
- エリア内のネイバー ルータ間でのプレーン テキスト認証および MD5 認証がサポートされています。
- 設定可能なルーティング インターフェイス パラメータには、インターフェイス出力コスト、再送信インターバル、インターフェイス送信遅延、ルータ プライオリティ、ルータの dead と hello インターバル、認証鍵などがあります。
- 仮想リンクがサポートされています。
- RFC 1587 に基づく Not-So-Stubby-Area (NSSA) がサポートされています。

通常、OSPF を使用するには、多くの内部ルータ、複数のエリアに接続された *Area Border Router* (ABR; エリア境界ルータ)、および *Autonomous System Boundary Router* (ASBR; 自律システム境界ルータ) 間で調整する必要があります。最小設定では、すべてのデフォルト パラメータ値、エリアに割り当てられたインターフェイスが使用され、認証は行われません。環境をカスタマイズする場合は、すべてのルータの設定を調整する必要があります。

ここでは、次の設定情報について説明します。

- 「OSPF のデフォルト設定」 (P.38-28)
- 「基本的な OSPF パラメータの設定」 (P.38-31)
- 「OSPF インターフェイスの設定」 (P.38-32)
- 「OSPF エリア パラメータの設定」 (P.38-33)
- 「その他の OSPF パラメータの設定」 (P.38-34)
- 「LSA グループ同期設定の変更」 (P.38-36)
- 「ループバック インターフェイスの設定」 (P.38-37)
- 「OSPF のモニタ」 (P.38-37)



(注)

OSPF をイネーブルにするには、スタック マスター上で IP サービス イメージが稼動している必要があります。

OSPF のデフォルト設定

表 38-5 に、OSPF のデフォルト設定を示します。

表 38-5 OSPF のデフォルト設定

機能	デフォルト設定
インターフェイス パラメータ	コスト：デフォルト コストは未定義。 再送信インターバル：5 秒。 送信遅延：1 秒。 プライオリティ：1。 hello インターバル：10 秒。 dead インターバル：hello インターバルの 4 倍。 認証なし。 パスワードの指定なし。 MD5 認証はディセーブル。
エリア	認証タイプ：0（認証なし）。 デフォルト コスト：1。 範囲：ディセーブル。 スタブ：スタブ エリアは未定義。 NSSA：NSSA エリアは未定義。
自動コスト	100 Mbps。
デフォルト情報送信元	ディセーブル。イネーブルの場合、デフォルトのメトリック設定は 10 で、外部ルート タイプのデフォルトはタイプ 2 です。
デフォルト メトリック	各ルーティング プロトコルに適切な、組み込みの自動メトリック変換。
距離 OSPF	dist1（エリア内のすべてのルート）：110。 dist2（エリア間のすべてのルート）：110。 dist3（他のルーティング ドメインからのルート）：110。
OSPF データベース フィルタ	ディセーブル。すべての発信 LSA がインターフェイスにフラッディングされます。
IP OSPF 名検索	ディセーブル。
隣接関係変更ログ	イネーブル
ネイバー	指定なし
ネイバー データベース フィルタ	ディセーブル。すべての発信 LSA はネイバーにフラッディングされます。
ネットワーク エリア	ディセーブル

表 38-5 OSPF のデフォルト設定 (続き)

機能	デフォルト設定
NSF ¹ 認識	IP サービス イメージが稼動しているスイッチの場合、イネーブル。 レイヤ 3 スイッチでは、ハードウェアやソフトウェアの変更中に、ネイバー NSF 対応ルータからのパケットを転送し続けることができます。
NSF 機能	ディセーブル。Catalyst 3750 スイッチは、IPv4 対応の OSPF NSF (NSF 対応ルーティング) をサポートしています。
ルータ ID	OSPF ルーティング プロセスは未定義
サマリー アドレス	ディセーブル
タイマー LSA グループの同期設定	240 秒
タイマー Shortest Path First (SPF)	spf-delay : 5 秒 spf-holdtime : 10 秒
仮想リンク	エリア ID またはルータ ID は未定義 hello インターバル : 10 秒 再送信インターバル : 5 秒 送信遅延 : 1 秒 dead インターバル : 40 秒 認証鍵 : 鍵は未定義 メッセージダイジェスト鍵 (MD5) : 鍵は未定義

1. NSF = Nonstop Forwarding

OSPF for Routed Access

Cisco IOS Release 12.2(55)SE では、IP ベース イメージで OSPF for Routed Access がサポートされます。IP サービス イメージは、ルート制限のない複数の OSPFv2 および OSPFv3 インスタンスが必要な場合に必要になります。また、IP サービス イメージは、マルチ VRF CE 機能をイネーブルにするためにも必要です。

OSPF for Routed Access は、レイヤ 3 ルーティング機能をワイヤリング クローゼットにまで拡張できるように、特に設計されたものです。



(注)

OSPF for Routed Access は、単一の OSPFv2 インスタンスおよび単一の OSPFv3 インスタンスと、ダイナミックに学習された合計 200 個のルートをサポートします。IP ベース イメージは、OSPF for Routed Access を提供します。

ただし、このリリースでは、これらの制限は適用されません。

キャンパス環境では、一般に、ローカルでないすべてのトラフィックをディストリビューション レイヤに転送するディストリビューション スイッチ (ハブ) にワイヤリング クローゼット (スポーク) が接続されたトポロジ (ハブ アンド スポーク) が使用されます。このトポロジでは、ワイヤリング クローゼット スイッチは、完全なルーティング テーブルを保持する必要がありません。ワイヤリング ク

ローゼットで OSPF for Routed Access を使用する場合は、ディストリビューション スイッチがワイヤリング クローゼット スイッチにデフォルト ルートを送信してエリア間および外部のルートに到達するベストプラクティス設計（OSPF スタブまたは完全なスタブ エリア設定）を使用する必要があります。

詳細については、次の URL にアクセスしてください。

<http://www.cisco.com/en/US/partner/docs/solutions/Enterprise/Campus/routed-ex.html>

OSPF Nonstop Forwarding

スイッチ スタックは 2 つのレベルの NSF をサポートしています。

- ・「OSPF NSF 認識」(P.38-30)
- ・「OSPF NSF 機能」(P.38-30)

OSPF NSF 認識

IP サービス イメージは IPv4 の OSPF NSF 認識をサポートしています。ネイバー ルータが NSF 対応で、レイヤ 3 スイッチでは、プライマリ Route Processor (RP) に障害が発生してルータのバックアップ RP によって引き継がれる前に、または処理を中断させずにソフトウェア アップグレードを行うためにプライマリ RP を手動でリロードしている間、ルータからパケットを転送し続けます。

この機能をディセーブルにすることはできません。この機能の詳細については、次の URL にある『OSPF Nonstop Forwarding (NSF) Awareness Feature Guide』を参照してください。

http://www.cisco.com/en/US/products/sw/iosswrel/ps1839/products_white_paper09186a0080153edd.shtml

OSPF NSF 機能

IP サービス イメージは、良好なコンバージェンスおよびスタック マスター変更後のトラフィックの損失を低減させるために、IPv4 の OSPF NSF 対応ルーティングをサポートしています。スタック マスターの変更が OSPF NSF 対応スタックで発生すると、新規スタック マスターでは、OSPF ネイバーとのリンクステート データベースを再同期させるために、次の 2 つのを行う必要があります。

- ・ ネイバー関係をリセットせずにネットワーク上にある、利用可能な OSPF ネイバーを解放する。
- ・ ネットワークのリンクステート データベースの内容を再取得する。

スタック マスターの変更後、新規マスターは OSPF NSF 信号をネイバー NSF 認識デバイスに送信します。デバイスでは、この信号を認識してスタックとのネイバー関係をリセットする必要がないことを把握します。NSF 対応スタック マスターは、ネットワーク上の他のルートから信号を受信すると、ネイバー リストの再構築を開始します。

ネイバー関係が再構築されると、NSF 対応スタック マスターがデータベースと NSF 認識ネイバーを再同期させ、ルーティング情報を OSPF ネイバーと交換します。新規スタック マスターがこのルーティング情報を使用して無効なルートを削除し、Routing Information Database (RIB) をアップデートし、新規情報で Forwarding Information Base (FIB; 転送情報ベース) をアップデートします。これで、OSPF プロトコルが完全にコンバージェンスします。



(注)

OSPF NSF では、すべてのネイバー ネットワーキング デバイスが NSF 認識となっている必要があります。NSF 対応ルータがネットワーク セグメント内で非 NSF 対応ネイバーを検出すると、そのセグメントの NSF 機能がディセーブルになります。すべてのデバイスが NSF 認識または NSF 対応である他のネットワーク セグメントでは、NSF 機能が提供され続けます。

OSPF NSF ルーティングをイネーブルにするには、**nsf** OSPF ルーティング コンフィギュレーション コマンドを使用します。このルーティングがイネーブルであることを確認するには、**show ip ospf** 特権 EXEC コマンドを使用します。

NSF の詳細については、次の URL にある『Cisco Nonstop Forwarding Feature Overview』を参照してください。

http://www.cisco.com/en/US/products/sw/iosswrel/ps1829/products_feature_guide09186a00800ab7fc.html



(注) NSF は、Hot Standby Router Protocol (HSRP; ホットスタンバイ ルータ プロトコル) に設定されたインターフェイスをサポートしていません。

基本的な OSPF パラメータの設定

OSPF をイネーブルにするには、OSPF ルーティング プロセスを作成し、ルーティング プロセスに関連付ける IP アドレスの範囲を指定して、この範囲に関連付けるエリア ID を割り当てる必要があります。

OSPF をイネーブルにするには、特権 EXEC モードで次の手順を実行します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	router ospf process-id	OSPF ルーティングをイネーブルにし、ルータ コンフィギュレーション モードを開始します。プロセス ID はローカルに割り当てられ、内部で使用される識別パラメータで、任意の正の整数を指定できます。各 OSPF ルーティング プロセスには一意の値があります。
ステップ 3	nsf	(任意) OSPF の NSF 動作をイネーブルにします。
ステップ 4	network address wildcard-mask area area-id	OSPF が動作するインターフェイス、およびそのインターフェイスのエリア ID を定義します。単一のコマンドにワイルドカード マスクを指定し、特定の OSPF エリアに関連付けるインターフェイスを 1 つまたは複数定義できます。エリア ID には 10 進数または IP アドレスを指定できます。
ステップ 5	end	特権 EXEC モードに戻ります。
ステップ 6	show ip protocols	設定を確認します。
ステップ 7	copy running-config startup-config	(任意) コンフィギュレーション ファイルに設定を保存します。

OSPF ルーティング プロセスを終了するには、**no router ospf process-id** グローバル コンフィギュレーション コマンドを使用します。

次に、OSPF ルーティング プロセスを設定し、プロセス番号 109 を割り当てる例を示します。

```
Switch(config)# router ospf 109
Switch(config-router)# network 131.108.0.0 255.255.255.0 area 24
```

OSPF インターフェイスの設定

ip ospf インターフェイス コンフィギュレーション コマンドを使用すると、インターフェイス固有の OSPF パラメータを変更できます。これらのパラメータを変更する必要はありませんが、一部のインターフェイス パラメータ (**hello** インターバル、**dead** インターバル、認証鍵など) については、接続されたネットワーク内のすべてのルータで統一性を維持する必要があります。これらのパラメータを変更した場合は、ネットワーク内のすべてのルータの値も同様に更新してください。



(注) **ip ospf** インターフェイス コンフィギュレーション コマンドはすべて任意です。

OSPF インターフェイス パラメータを変更するには、特権 EXEC モードで次の手順を実行します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	interface interface-id	インターフェイス コンフィギュレーション モードを開始し、設定するレイヤ 3 インターフェイスを指定します。
ステップ 3	ip ospf cost	(任意) インターフェイスでパケットを送信するコストを明確に指定します。
ステップ 4	ip ospf retransmit-interval seconds	(任意) リンクステート アドバタイズメント送信間隔を秒数で指定します。範囲は 1 ～ 65535 秒です。デフォルト値は 5 秒です。
ステップ 5	ip ospf transmit-delay seconds	(任意) リンク ステート アップデート パケットを送信するまでの予測待機時間を秒数で設定します。範囲は 1 ～ 65535 秒です。デフォルト値は 1 秒です。
ステップ 6	ip ospf priority number	(任意) ネットワークに対して、OSPF で指定されたルータを検索するときに役立つプライオリティを設定します。指定できる範囲は 0 ～ 255 です。デフォルト値は 1 です。
ステップ 7	ip ospf hello-interval seconds	(任意) OSPF インターフェイスで hello パケットの送信間隔を秒数で設定します。値はネットワークのすべてのノードで同じとします。範囲は 1 ～ 65535 秒です。デフォルト値は 10 秒です。
ステップ 8	ip ospf dead-interval seconds	(任意) 最後のデバイスで hello パケットが確認されてから、OSPF ルータがダウンしていることがネイバーによって宣言されるまでの時間を秒数で設定します。値はネットワークのすべてのノードで同じとします。範囲は 1 ～ 65535 秒です。デフォルト値は hello インターバルの 4 倍です。
ステップ 9	ip ospf authentication-key key	(任意) ネイバー OSPF ルータで使用するパスワードを割り当てます。パスワードには、キーボードから入力した任意の文字列 (最大 8 バイト長) を指定できます。同じネットワーク上のすべてのネイバー ルータには、OSPF 情報を交換するため、同じパスワードを設定する必要があります。
ステップ 10	ip ospf message digest-key keyid md5 key	(任意) MD5 認証をイネーブルにします。 keyid : 1 ～ 255 の ID key : 最大 16 バイトの英数字パスワード
ステップ 11	ip ospf database-filter all out	(任意) インターフェイスへの OSPF LSA パケットのフラッドイングを阻止します。デフォルトでは、LSA が着信するインターフェイスを除き、同じエリア内のすべてのインターフェイスに OSPF は新しい LSA をフラッドイングします。

	コマンド	目的
ステップ 12	end	特権 EXEC モードに戻ります。
ステップ 13	show ip ospf interface <i>[interface-name]</i>	OSPF に関連するインターフェイス情報を表示します。
ステップ 14	show ip ospf neighbor detail	ネイバー スイッチの NSF 認証ステータスを表示します。出力は、次のいずれかに一致します。 <i>Options is 0x52</i> <i>LLS Options is 0x1 (LR)</i> - これらの行の両方が表示される場合、ネイバー スイッチが NSF アウェアです。 <i>Options is 0x42</i> : ネイバー スイッチが NSF アウェアでないことを示します。
ステップ 15	copy running-config startup-config	(任意) コンフィギュレーション ファイルに設定を保存します。

設定されたパラメータ値を削除する場合、またはデフォルト値に戻す場合は、上記コマンドの **no** 形式を使用します。

OSPF エリア パラメータの設定

複数の OSPF エリア パラメータを設定することもできます。設定できるパラメータには、エリア、スタブ エリア、および Not-So-Stubby-Area (NSSA) への無許可アクセスをパスワードによって阻止する認証用パラメータがあります。スタブ エリアに外部ルートに関する情報は送信されませんが、代わりに、Autonomous System (AS; 自律システム) 外の宛先に対するデフォルトの外部ルートが、Area Border Router (ABR) によって生成されます。NSSA ではコアからそのエリアへ向かう LSA の一部がフラッドイングされませんが、再配信することによって、エリア内の AS 外部ルートを取り込むことができます。

ルートのサマライズは、アドバタイズされたアドレスを、他のエリアでアドバタイズされる単一のサマリー ルートに統合することです。ネットワーク番号が連続する場合は、**area range** ルータ コンフィギュレーション コマンドを使用し、範囲内のすべてのネットワークを対象とするサマリー ルートをアドバタイズするように ABR を設定できます。



(注) OSPF **area** ルータ コンフィギュレーション コマンドはすべて任意です。

エリア パラメータを設定するには、特権 EXEC モードで次の手順を実行します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	router ospf <i>process-id</i>	OSPF ルーティングをイネーブルにし、ルータ コンフィギュレーション モードを開始します。
ステップ 3	area <i>area-id</i> authentication	(任意) 特定のエリアへの無許可アクセスに対して、パスワードベースの保護を可能にします。ID には 10 進数または IP アドレスのいずれかを指定できます。
ステップ 4	area <i>area-id</i> authentication message-digest	(任意) エリアに関して MD5 認証をイネーブルにします。

	コマンド	目的
ステップ 5	<code>area area-id stub [no-summary]</code>	(任意) エリアをスタブ エリアとして定義します。 no-summary キーワードを指定すると、ABR はサマリー リンク アドバタイズメントをスタブ エリアに送信できなくなります。
ステップ 6	<code>area area-id nssa [no-redistribution]</code> <code>[default-information-originate]</code> <code>[no-summary]</code>	(任意) エリアを NSSA として定義します。同じエリア内のすべてのルータは、エリアが NSSA であることを認識する必要があります。次のキーワードのいずれかを選択します。 • no-redistribution : ルータが NSSA ABR の場合、redistribute コマンドを使用して、ルートを NSSA でなく通常のエリアに取り込む場合に選択します。 • default-information-originate : タイプ 7 LSA を NSSA に取り込むようにする場合に、ABR で選択します。 • no-redistribution : サマリー LSA を NSSA に送信しない場合に選択します。
ステップ 7	<code>area area-id range address mask</code>	(任意) 単一のルートをアドバタイズするアドレス範囲を指定します。このコマンドは、ABR に対してだけ使用します。
ステップ 8	<code>end</code>	特権 EXEC モードに戻ります。
ステップ 9	<code>show ip ospf [process-id]</code> <code>show ip ospf [process-id [area-id]]</code> <code>database</code>	設定を確認するため、一般的な OSPF ルーティング プロセスまたは特定のプロセス ID に関する情報を表示します。 特定のルータの OSPF データベースに関連する情報のリストを表示します。
ステップ 10	<code>copy running-config startup-config</code>	(任意) コンフィギュレーション ファイルに設定を保存します。

設定されたパラメータ値を削除する場合、またはデフォルト値に戻す場合は、上記コマンドの **no** 形式を使用します。

その他の OSPF パラメータの設定

ルータ コンフィギュレーション モードで、その他の OSPF パラメータを設定することもできます。

- ルート サマライズ : 他のプロトコルからのルートを再配信すると（「[ルート マップによるルーティング情報の再配信](#)」(P.38-99) を参照）、各ルートは外部 LSA 内で個別にアドバタイズされます。OSPF リンク ステート データベースのサイズを小さくするには、**summary-address** ルータ コンフィギュレーション コマンドを使用し、指定されたネットワーク アドレスおよびマスクに含まれる、再配信されたすべてのルートを単一のルータにアドバタイズします。
- 仮想リンク : OSPF では、すべてのエリアがバックボーン エリアに接続されている必要があります。バックボーンが不連続である場合に仮想リンクを確立するには、2 つの ABR を仮想リンクのエンドポイントとして設定します。設定情報には、他の仮想エンドポイント（他の ABR）の ID、および 2 つのルータに共通する非バックボーン リンク（通過エリア）などがあります。仮想リンクをスタブ エリアから設定することはできません。
- デフォルト ルート : OSPF ルーティング ドメイン内へのルート再配信を設定すると、ルータは自動的に自律システム境界ルータ (ASBR) になります。ASBR を設定し、強制的に OSPF ルーティング ドメインにデフォルト ルートを生成できます。
- すべての OSPF **show** 特権 EXEC コマンドで使用される Domain Name Server (DNS) 名を使用すると、ルータ ID やネイバー ID を指定して表示する場合に比べ、ルータを簡単に特定できます。

- デフォルト メトリック : OSPF は、インターフェイスの帯域幅に従ってインターフェイスの OSPF メトリックを計算します。メトリックは、帯域幅で分割された *ref-bw* として計算されます。ここでの *ref* のデフォルト値は 10 で、帯域幅 (*bw*) は **bandwidth** インターフェイス コンフィギュレーション コマンドによって指定されます。大きな帯域幅を持つ複数のリンクの場合は、大きな数値を指定し、これらのリンクのコストを区別できます。
- 管理距離は、ルーティング情報送信元の信頼性を表す数値です。0 ～ 255 の整数を指定でき、値が大きいほど信頼性は低下します。管理距離が 255 の場合はルーティング情報送信元をまったく信頼できないため、無視する必要があります。OSPF では、エリア内のルート (エリア内)、別のエリアへのルート (エリア間)、および再配信によって取得した別のルーティング ドメインからのルート (外部) の 3 つの管理距離が使用されます。どの管理距離の値でも変更できます。
- パッシブ インターフェイス : イーサネット上の 2 つのデバイス間のインターフェイスは 1 つのネットワーク セグメントしか表しません。このため、OSPF が送信側インターフェイスに hello パケットを送信しないようにするには、送信側デバイスをパッシブ インターフェイスに設定する必要があります。両方のデバイスは受信側インターフェイス宛の hello パケットを使用することで、相互の識別を可能にします。
- ルート計算タイマー : OSPF がトポロジ変更を受信してから SPF 計算を開始するまでの遅延時間、および 2 つの SPF 計算の間のホールド タイムを設定できます。
- ネイバー変更ログ : OSPF ネイバー ステートが変更されたときに Syslog メッセージを送信するようにルータを設定し、ルータの変更内容の概要を表示できます。

上記の OSPF パラメータを設定するには、特権 EXEC モードで次の手順を実行します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	router ospf process-id	OSPF ルーティングをイネーブルにし、ルータ コンフィギュレーション モードを開始します。
ステップ 3	summary-address address mask	(任意) 1 つのサマリー ルートだけがアドバタイズされるように、再配信されたルートのアドレスおよび IP サブネット マスクを指定します。
ステップ 4	area area-id virtual-link router-id [hello-interval seconds] [retransmit-interval seconds] [trans] [[authentication-key key] message-digest-key keyid md5 key]]	(任意) 仮想リンクを確立し、パラメータを設定します。パラメータ定義については「 OSPF インターフェイスの設定 」(P.38-32)、仮想リンクのデフォルト設定については表 38-5 (P.38-28) を参照してください。
ステップ 5	default-information originate [always] [metric metric-value] [metric-type type-value] [route-map map-name]	(任意) 強制的に OSPF ルーティング ドメインにデフォルト ルートを生成するように ASBR を設定します。パラメータはすべて任意です。
ステップ 6	ip ospf name-lookup	(任意) DNS 名検索を設定します。デフォルトはディセーブルです。
ステップ 7	ip auto-cost reference-bandwidth ref-bw	(任意) 単一のルートをアドバタイズするアドレス範囲を指定します。このコマンドは、ABR に対してだけ使用します。
ステップ 8	distance ospf {[inter-area dist1] [inter-area dist2] [external dist3]}	(任意) OSPF の距離の値を変更します。各タイプのルートのデフォルト距離は 110 です。指定できる範囲は 1 ～ 255 です。
ステップ 9	passive-interface type number	(任意) 指定されたインターフェイス経由の hello パケットの送信を抑制します。

	コマンド	目的
ステップ 10	timers throttle spf <i>spf-delay spf-holdtime spf-wait</i>	(任意) ルート計算タイマーを設定します。 • <i>spf-delay</i> : SPF 計算の変更を受信する間の遅延。指定できる範囲は 1 ～ 600000 ミリ秒です。 • <i>spf-holdtime</i> : 最初と 2 番目の SPF 計算の間の遅延。指定できる範囲は 1 ～ 600000 ミリ秒です。 • <i>spf-wait</i> : SPF 計算の最大待機時間 (ミリ秒)。指定できる範囲は 1 ～ 600000 ミリ秒です。
ステップ 11	ospf log-adj-changes	(任意) ネイバー ステートが変更されたとき、Syslog メッセージを送信します。
ステップ 12	end	特権 EXEC モードに戻ります。
ステップ 13	show ip ospf [<i>process-id</i> [<i>area-id</i>]] database	特定のルータの OSPF データベースに関連する情報のリストを表示します。キーワード オプションの一部については、「 OSPF のモニタ 」(P.38-37) を参照してください。
ステップ 14	copy running-config startup-config	(任意) コンフィギュレーション ファイルに設定を保存します。

LSA グループ同期設定の変更

OSPF LSA グループ同期設定機能を使用すると、OSPF LSA をグループ化し、リフレッシュ、チェックサム、エージング機能の同期を取って、ルータをより効率的に使用することが可能となります。デフォルトでこの機能はイネーブルとなっています。デフォルトの同期インターバルは 4 分間です。通常は、このパラメータを変更する必要はありません。最適なグループ同期インターバルは、ルータがリフレッシュ、チェックサム、エージングを行う LSA 数に反比例します。たとえば、データベース内に約 10,000 個の LSA が格納されている場合は、同期設定インターバルを短くすると便利です。小さなデータベース (40 ～ 100 LSA) を使用する場合は、同期インターバルを長くし、10 ～ 20 分に設定してください。

OSPF LSA 同期を設定するには、特権 EXEC モードで次の手順を実行します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	router ospf <i>process-id</i>	OSPF ルーティングをイネーブルにし、ルータ コンフィギュレーション モードを開始します。
ステップ 3	timers lsa-group-pacing <i>seconds</i>	LSA のグループ同期を変更します。
ステップ 4	end	特権 EXEC モードに戻ります。
ステップ 5	show running-config	設定を確認します。
ステップ 6	copy running-config startup-config	(任意) コンフィギュレーション ファイルに設定を保存します。

デフォルト値に戻すには、**no timers lsa-group-pacing** ルータ コンフィギュレーション コマンドを使用します。

ループバック インターフェイスの設定

OSPF は、インターフェイスに設定されている最大の IP アドレスをルータ ID として使用します。このインターフェイスがダウンした場合、または削除された場合、OSPF プロセスは新しいルータ ID を再計算し、すべてのルーティング情報をそのルータのインターフェイスから再送信します。ループバック インターフェイスが IP アドレスによって設定されている場合、より大きな IP アドレスが他のインターフェイスにある場合でも、OSPF はこの IP アドレスをルータ ID として使用します。ループバック インターフェイスに障害は発生しないため、安定性は増大します。OSPF は他のインターフェイスよりもループバック インターフェイスを自動的に優先し、すべてのループバック インターフェイスの中で最大の IP アドレスを選択します。

ループバック インターフェイスを設定するには、特権 EXEC モードで次の手順を実行します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	interface loopback 0	ループバック インターフェイスを作成し、インターフェイス コンフィギュレーション モードを開始します。
ステップ 3	ip address address mask	このインターフェイスに IP アドレスを割り当てます。
ステップ 4	end	特権 EXEC モードに戻ります。
ステップ 5	show ip interface	設定を確認します。
ステップ 6	copy running-config startup-config	(任意) コンフィギュレーション ファイルに設定を保存します。

ループバック インターフェイスをディセーブルにするには、**no interface loopback 0** グローバル コンフィギュレーション コマンドを使用します。

OSPF のモニタ

IP ルーティング テーブル、キャッシュ、データベースの内容など、特定の統計情報を表示できます。

表 38-6 に、統計情報を表示するために使用する特権 EXEC コマンドの一部を示します。**show ip ospf database** 特権 EXEC コマンドのオプションおよび表示されるフィールドの詳細については、『*Cisco IOS IP Command Reference, Volume 2 of 3: Routing Protocols, Release 12.2*』を参照してください。

表 38-6 IP OSPF 統計情報の表示コマンド

コマンド	目的
show ip ospf [process-id]	OSPF ルーティング プロセスに関する一般的な情報を表示します。
show ip ospf [process-id] database [router] [link-state-id]	OSPF データベースに関連する情報を表示します。
show ip ospf [process-id] database [router] [self-originate]	
show ip ospf [process-id] database [router] [adv-router [ip-address]]	
show ip ospf [process-id] database [network] [link-state-id]	
show ip ospf [process-id] database [summary] [link-state-id]	
show ip ospf [process-id] database [asbr-summary] [link-state-id]	
show ip ospf [process-id] database [external] [link-state-id]	
show ip ospf [process-id area-id] database [database-summary]	

表 38-6 IP OSPF 統計情報の表示コマンド

コマンド	目的
<code>show ip ospf border-routes</code>	内部の OSPF ルーティング ABR および ASBR テーブル エントリを表示します。
<code>show ip ospf interface [interface-name]</code>	OSPF に関連するインターフェイス情報を表示します。
<code>show ip ospf neighbor [interface-name] [neighbor-id] detail</code>	OSPF インターフェイス ネイバー情報を表示します。
<code>show ip ospf virtual-links</code>	OSPF に関連する仮想リンク情報を表示します。

EIGRP の設定

EIGRP は IGRP のシスコ独自の拡張バージョンです。EIGRP は IGRP と同じディスタンス ベクタ アルゴリズムおよび距離情報を使用しますが、EIGRP では収束性および動作効率が大幅に改善されています。

コンバージェンス技術には、Diffusing Update Algorithm (DUAL) と呼ばれるアルゴリズムが採用されています。DUAL を使用すると、ルート計算の各段階でループが発生しなくなり、トポロジの変更に関連するすべてのデバイスを同時に同期できます。トポロジ変更の影響を受けないルータは、再計算から除外されます。

IP EIGRP を導入すると、ネットワークの幅が広がります。RIP の場合、ネットワークの最大幅は 15 ホップです。EIGRP メトリックは数千ホップをサポートするほど大きいので、ネットワークを拡張するときに問題となるのは、トランスポート レイヤのホップ カウンタだけです。IP パケットが 15 台のルータを経由し、宛先方向のネクスト ホップが EIGRP によって取得されている場合だけ、EIGRP は転送制御フィールドの値を増やします。RIP ルートを宛先へのネクスト ホップとして使用する場合、転送制御フィールドでは、通常どおり値が増加します。

EIGRP には次の機能があります。

- 高速コンバージェンス
- 差分更新：宛先のステートが変更された場合、ルーティング テーブルの内容全体を送信する代わりに差分更新を行い、EIGRP パケットに必要な帯域幅を最小化します。
- 低い CPU 使用率：受信ごとに完全更新パケットを処理する必要がないため、CPU 使用率が低下します。
- プロトコルに依存しないネイバー ディスカバリ メカニズム：このメカニズムを使用しネイバー ルータに関する情報を取得します。
- Variable-Length Subnet Mask (VLSM; 可変長サブネット マスク)
- 任意のルート サマライズ
- 大規模ネットワークへの対応

EIGRP には次に示す 4 つの基本コンポーネントがあります。

- **ネイバー ディスカバリおよび回復**：直接接続されたネットワーク上の他のルータに関する情報を動的に取得するために、ルータで使用されるプロセスです。ネイバーが到達不能になる場合、または操作不能になった場合、ルータもこの情報を検出する必要があります。ネイバー ディスカバリおよび回復は、サイズの小さな hello パケットを定期的を送信することにより、わずかなオーバーヘッドで実現されます。hello パケットが受信されている限り、Cisco IOS ソフトウェアは、ネイバーが有効に機能していると学習します。このように判別された場合、ネイバー ルータはルーティング情報を交換できます。

- **信頼できるトランスポート プロトコル**: EIGRP パケットをすべてのネイバーに確実に、順序どおりに配信します。マルチキャストおよびユニキャスト パケットが混在する送信もサポートされます。EIGRP パケットには確実に送信する必要があるものと、そうでないものがあります。効率を高めるために、必要な場合だけ信頼性が確保されます。たとえば、マルチキャスト機能があるマルチアクセス ネットワーク（イーサネットなど）では、すべてのネイバーにそれぞれ hello パケットを確実に送信する必要はありません。したがって、EIGRP はパケットへの確認応答が不要であることを知らせる、レシーバー宛の情報をパケットに格納し、単一のマルチキャスト hello を送信します。他のタイプのパケット（アップデートなど）の場合は、確認応答（ACK パケット）を要求します。信頼性の高い伝送であれば、ペンディング中の未確認応答パケットがある場合、マルチキャスト パケットを迅速に送信できます。このため、リンク速度が変化する場合でも、コンバージェンス時間を短く保つことができます。
- **DUAL 有限状態マシン**: すべてのルート計算に関する決定プロセスを統合し、すべてのネイバーによってアドバタイズされたすべてのルートをトラッキングします。DUAL は距離情報（メトリックともいう）を使用して、効率的な、ループのないパスを選択し、さらに DUAL は適切な後継ルータに基づいて、ルーティング テーブルに挿入するルートを選択します。後継ルータは、宛先への最小コスト パス（ルーティング ループに関連しないことが保証されている）を持つ、パケット転送に使用されるネイバー ルータです。適切な後継ルータが存在しなくても、宛先にアドバタイズするネイバーが存在する場合は再計算が行われ、この結果、新しい後継ルータが決定されます。ルートの再計算に要する時間によって、コンバージェンス時間が変わります。再計算はプロセスに負荷がかかるため、必要でない場合は、再計算しないようにしてください。トポロジが変更されると、DUAL は適切な後継ルータの有無を調べます。適切な後継ルータが存在する場合は、それらを探して使用し、不要な再計算を回避します。
- **プロトコル依存モジュール**: ネットワーク レイヤ プロトコル特有の作業を行います。たとえば、IP EIGRP モジュールは、IP でカプセル化された EIGRP パケットを送受信します。このモジュールは、EIGRP パケットを解析し、受信した新しい情報を DUAL に通知する作業を行います。EIGRP は DUAL にルーティング決定を行うように要求しますが、結果は IP ルーティング テーブルに格納されます。EIGRP は、他の IP ルーティング プロトコルによって取得したルートの再配信も行います。

ここでは、次の設定情報について説明します。

- 「EIGRP のデフォルト設定」(P.38-39)
- 「基本的な EIGRP パラメータの設定」(P.38-42)
- 「EIGRP インターフェイスの設定」(P.38-43)
- 「EIGRP ルート認証の設定」(P.38-44)
- 「EIGRP スタブ ルーティングの設定」(P.38-45)
- 「EIGRP のモニタおよびメンテナンス」(P.38-46)



(注) EIGRP をイネーブルにするには、スタック マスター上で IP サービス イメージが稼動している必要があります。

EIGRP のデフォルト設定

表 38-7 に、EIGRP のデフォルト設定を示します。

表 38-7 EIGRP のデフォルト設定

機能	デフォルト設定
自動サマリー	イネーブル。クラスフル ネットワーク境界を通過するとき、この境界にサブプレフィクスがサマライズされます。
デフォルト情報	再配信中は外部ルートが許可され、EIGRP プロセス間でデフォルト情報が渡されます。
デフォルト メトリック	デフォルト メトリックなしで再配信できるのは、接続されたルートおよびインターフェイスのスタティック ルートだけです。デフォルト メトリックは次のとおりです。 - 帯域幅：0 kbps 以上 - 遅延（10 マイクロ秒）：0 または 39.1 ナノ秒の倍数である任意の正の数値 - 信頼性：0 ～ 255 の任意の数値（255 の場合は信頼性が 100%） - 負荷：0 ～ 255 の数値で表される有効帯域幅（255 の場合は 100% の負荷） - Maximum Transmission Unit（MTU; 最大伝送ユニット）：バイトで表されたルートの MTU サイズ（0 または任意の正の整数）
距離	内部距離：90 外部距離：170
EIGRP のネイバー関係変更ログ	ディセーブル。隣接関係の変更はロギングされません。
IP 認証キーチェーン	認証なし
IP 認証モード	認証なし
IP 帯域幅比率	50%
IP hello 間隔	低速の Nonbroadcast Multiaccess（NBMA; 非ブロードキャスト マルチアクセス）ネットワークの場合：60 秒、それ以外のネットワークの場合：5 秒
IP ホールド タイム	低速の NBMA ネットワークの場合：180 秒、それ以外のネットワークの場合：15 秒
IP スプリットホライズン	イネーブル
IP サマリー アドレス	サマリー集約アドレスは未定義
メトリック ウェイト	tos：0。k1 および k3：1。k2、k4、および k5：0。
ネットワーク	指定なし
NSF ¹ 認識	IP サービス イメージが稼動しているスイッチの場合、イネーブル。 レイヤ 3 スイッチでは、ハードウェアやソフトウェアの変更中に、ネイバー NSF 対応ルータからのパケットを転送し続けることができます。
NSF 機能	ディセーブル (注) Catalyst 3750 スイッチは IPv4 の EIGRP NSF 対応ルーティングをサポートしています。
オフセットリスト	ディセーブル
ルータ EIGRP	ディセーブル

表 38-7 EIGRP のデフォルト設定 (続き)

機能	デフォルト設定
メトリック設定	ルート マップにはメトリック設定なし
トラフィック共有	メトリックの比率に応じて配分
差異	1 (等価コスト ロード バランシング)

1. NSF = Nonstop Forwarding

EIGRP ルーティング プロセスを作成するには、EIGRP をイネーブルにし、ネットワークを関連付ける必要があります。EIGRP は指定されたネットワーク内のインターフェイスにアップデートを送信します。インターフェイス ネットワークを指定しないと、どの EIGRP アップデートでもアドバタイズされません。



(注)

ネットワーク上に IGRP 用に設定されているルータがあり、この設定を EIGRP に変更する場合は、IGRP と EIGRP の両方が設定された移行ルータを指定する必要があります。この場合は、この次のセクションに記載されているステップ 1 ~ 3 を実行してください ([「スプリット ホライズンの設定」\(P.38-26\)](#) も参照)。ルートを自動的に再配信するには、同じ AS 番号を使用する必要があります。

EIGRP Nonstop Forwarding

スイッチ スタックは 2 つのレベルの EIGRP をサポートしています。

- 「EIGRP NSF 認識」(P.38-41)
- 「EIGRP NSF 機能」(P.38-41)

EIGRP NSF 認識

EIGRP NSF 認識機能が IP サービス イメージの IPv4 でサポートされています。ネイバー ルータが NSF 対応である場合、レイヤ 3 スイッチでは、ルータに障害が発生してプライマリ RP がバックアップ RP によって引き継がれる間、または処理を中断せずにソフトウェア アップグレードを行うためにプライマリ RP を手動でリロードしている間、ネイバー ルータからパケットを転送し続けます。

この機能をディセーブルにすることはできません。この機能の詳細については、次の URL にある

『*EIGRP Nonstop Forwarding (NSF) Awareness Feature Guide*』を参照してください。

http://www.cisco.com/en/US/products/sw/iosswrel/ps1839/products_feature_guide09186a0080160010.html

EIGRP NSF 機能

良好なコンバージェンスと、スタック マスター変更後のトラフィック損失を低減させるために、IP サービス イメージで IPv4 の EIGRP NSF 対応ルーティングをサポートしています。EIGRP NSF 対応スタック マスターがリスタートする際、または新規スタック マスターが起動して NSF が再起動する際、スイッチにはネイバーがなく、トポロジテーブルは空です。スイッチでは、スイッチ スタックに向かうトラフィックを中断せずに、インターフェイスを始動させ、ネイバーを再取得して、トポロジとルーティング テーブルを再構築する必要があります。EIGRP ピア ルータでは、新規スタック マスターから学習したルートを維持し、NSF リスタート プロセスを通じてトラフィックを転送し続けます。

ネイバーによる隣接リセットを避けるために、新規スタック マスターは新規リスタート (RS) ビットを使用して EIGRP パケット ヘッダーにリスタートを表示します。ネイバーがこれを受信すると、ピアリスト内のスタックを同期させてスタックとの隣接関係を維持します。次にネイバーは、NSF 認識を示し新規スタック マスターを援助するように設定された RS ビットとともにトポロジテーブルをスタック マスターに送信します。

少なくともスタック ピア ネイバーの 1 つが NSF 認識である場合、スタック マスターはアップデートを受信してそのデータベースを再構築します。各 NSF 認識ネイバーは、最後のアップデート パケットで End of Table (EOT) マーカーを送信して、テーブル内容の終了を示します。スタック マスターで EOT マーカーを受信すると、コンバージェンスを認識し、アップデートの送信を開始します。スタック マスターがすべての EOT マーカーをネイバーから受信するか、あるいは NSF コンバージ タイマーが切れると、EIGRP によって RIB がコンバージェンスが通知され、トポロジテーブルがすべての NSF 認識ピアにフラッディングされます。



(注) NSF は、Hot Standby Router Protocol (HSRP; ホットスタンバイ ルータ プロトコル) に設定されたインターフェイスをサポートしていません。

EIGRP NSF ルーティングをイネーブルにするには、**nsf** EIGRP ルーティング コンフィギュレーション コマンドを使用します。NFS がデバイスでイネーブルであることを確認するには、**show ip protocol** 特権 EXEC コマンドを使用します。**nsf** コマンドの詳細については、このリリースに対応するコマンド リファレンスを参照してください。

基本的な EIGRP パラメータの設定

EIGRP を設定するには、特権 EXEC モードで次の手順を実行します。ルーティング プロセスの設定は必須ですが、それ以外のステップは任意です。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	router eigrp <i>autonomous-system number</i>	EIGRP ルーティング プロセスをイネーブルにし、ルータ コンフィギュレーション モードを開始します。AS 番号によって他の EIGRP ルータへのルートを特定し、ルーティング情報をタグ付けします。
ステップ 3	nsf	(任意) EIGRP NSF をイネーブルにします。スタック マスターとすべてのピアにこのコマンドを入力します。
ステップ 4	network <i>network-number</i>	ネットワークを EIGRP ルーティング プロセスに関連付けます。EIGRP は指定されたネットワーク内のインターフェイスにアップデートを送信します。
ステップ 5	eigrp log-neighbor-changes	(任意) EIGRP ネイバー関係変更のログギングをイネーブルにし、ルーティング システムの安定性をモニタします。
ステップ 6	metric weights <i>tos k1 k2 k3 k4 k5</i>	(任意) EIGRP メトリックを調整します。デフォルト値はほとんどのネットワークで適切に動作するよう入念に設定されていますが、調整することも可能です。 注意 メトリックを設定する作業は複雑です。熟練したネットワーク設計者の指導がない場合は、行わないでください。

	コマンド	目的
ステップ 7	offset list [<i>access-list number</i> <i>name</i>] { in out } <i>offset</i> [<i>type number</i>]	(任意) オフセット リストをルーティング メトリックに適用し、EIGRP によって取得したルートへの着信および発信メトリックを増やします。アクセス リストまたはインターフェイスを使用し、オフセット リストを制限できます。
ステップ 8	no auto-summary	(任意) ネットワークレベル ルートへのサブネット ルートの自動サマライズをディセーブルにします。
ステップ 9	ip summary-address eigrp <i>autonomous-system-number address mask</i>	(任意) サマリー集約を設定します。
ステップ 10	end	特権 EXEC モードに戻ります。
ステップ 11	show ip protocols	設定を確認します。
ステップ 12	show ip protocols	設定を確認します。 NSF 認識の場合、出力に次のように表示されます。 *** IP Routing is NSF aware *** EIGRP NSF enabled
ステップ 13	copy running-config startup-config	(任意) コンフィギュレーション ファイルに設定を保存します。

機能をディセーブルにする場合、または設定をデフォルト値に戻す場合は、上記コマンドの **no** 形式を使用します。

EIGRP インターフェイスの設定

インターフェイスごとに、他の EIGRP パラメータを任意で設定できます。

EIGRP インターフェイスを設定するには、特権 EXEC モードで次の手順を実行します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	interface <i>interface-id</i>	インターフェイス コンフィギュレーション モードを開始し、設定するレイヤ 3 インターフェイスを指定します。
ステップ 3	ip bandwidth-percent eigrp <i>percent</i>	(任意) インターフェイスで EIGRP が使用できる帯域幅の割合を設定します。デフォルト値は 50% です。
ステップ 4	ip summary-address eigrp <i>autonomous-system-number address mask</i>	(任意) 指定されたインターフェイスのサマリー集約アドレスを設定します (auto-summary がイネーブルの場合は、通常設定する必要はありません)。
ステップ 5	ip hello-interval eigrp <i>autonomous-system-number seconds</i>	(任意) EIGRP ルーティング プロセスの hello タイム インターバルを変更します。範囲は 1 ~ 65535 秒です。低速 NBMA ネットワークの場合のデフォルトは 60 秒、その他のすべてのネットワークでは 5 秒です。

EIGRP の設定

	コマンド	目的
ステップ 6	ip hold-time eigrp autonomous-system-number seconds	(任意) EIGRP ルーティング プロセスのホールド タイム インターバルを変更します。範囲は 1 ～ 65535 秒です。低速 NBMA ネットワークの場合のデフォルトは 180 秒、その他のすべてのネットワークでは 15 秒です。  注意 ホールド タイムを調整する前に、シスコのテクニカル サポートにお問い合わせください。
ステップ 7	no ip split-horizon eigrp autonomous-system-number	(任意) スプリット ホライズンをディセーブルにし、ルート情報が情報元インターフェイスからルータによってアドバタイズされるようにします。
ステップ 8	end	特権 EXEC モードに戻ります。
ステップ 9	show ip eigrp interface	EIGRP がアクティブであるインターフェイス、およびそれらのインターフェイスに関連する EIGRP の情報を表示します。
ステップ 10	copy running-config startup-config	(任意) コンフィギュレーション ファイルに設定を保存します。

機能をディセーブルにする場合、または設定をデフォルト値に戻す場合は、上記コマンドの **no** 形式を使用します。

EIGRP ルート認証の設定

EIGRP ルート認証を行うと、EIGRP ルーティング プロトコルからのルーティング アップデートに関する MD5 認証が可能になり、承認されていない送信元から無許可または問題のあるルーティング メッセージを受け取ることがなくなります。

認証をイネーブルにするには、特権 EXEC モードで次の手順を実行します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	interface interface-id	インターフェイス コンフィギュレーション モードを開始し、設定するレイヤ 3 インターフェイスを指定します。
ステップ 3	ip authentication mode eigrp autonomous-system md5	IP EIGRP パケットの MD5 認証をイネーブルにします。
ステップ 4	ip authentication key-chain eigrp autonomous-system key-chain	IP EIGRP パケットの認証をイネーブルにします。
ステップ 5	exit	グローバル コンフィギュレーション モードに戻ります。
ステップ 6	key chain name-of-chain	キー チェーンを識別し、キーチェーン コンフィギュレーション モードを開始します。ステップ 4 で設定した名前を指定します。
ステップ 7	key number	キーチェーン コンフィギュレーション モードで、鍵番号を識別します。
ステップ 8	key-string text	キーチェーン コンフィギュレーション モードで、キー スtring を識別します。

	コマンド	目的
ステップ 9	accept-lifetime <i>start-time</i> { infinite <i>end-time</i> duration <i>seconds</i> }	(任意) 鍵を受信する期間を指定します。 <i>start-time</i> および <i>end-time</i> 構文には、 <i>hh:mm:ss Month date year</i> または <i>hh:mm:ss date Month year</i> のいずれかを使用できます。デフォルトはデフォルトの <i>start-time</i> 以降、無制限です。指定できる最初の日付は 1993 年 1 月 1 日です。デフォルトの <i>end-time</i> および duration は infinite です。
ステップ 10	send-lifetime <i>start-time</i> { infinite <i>end-time</i> duration <i>seconds</i> }	(任意) 鍵を送信する期間を指定します。 <i>start-time</i> および <i>end-time</i> 構文には、 <i>hh:mm:ss Month date year</i> または <i>hh:mm:ss date Month year</i> のいずれかを使用できます。デフォルトはデフォルトの <i>start-time</i> 以降、無制限です。指定できる最初の日付は 1993 年 1 月 1 日です。デフォルトの <i>end-time</i> および duration は infinite です。
ステップ 11	end	特権 EXEC モードに戻ります。
ステップ 12	show key chain	認証鍵情報を表示します。
ステップ 13	copy running-config startup-config	(任意) コンフィギュレーション ファイルに設定を保存します。

機能をディセーブルにする場合、または設定をデフォルト値に戻す場合は、上記コマンドの **no** 形式を使用します。

EIGRP スタブ ルーティングの設定

EIGRP スタブ ルーティング機能は、すべてのイメージで使用することができ、エンド ユーザの近くにルーテッド トラフィックを移動することでリソースの利用率を低減させます。



(注)

IP ベース イメージに含まれているのは EIGRP スタブ ルーティング機能だけです。この機能は、ルーティング テーブルからネットワークの他のスイッチに接続ルートまたは集約ルートをアドバタイズするだけです。スイッチはアクセス レイヤで EIGRP スタブ ルーティングを使用するため、その他の種類のルーティング アドバタイズを使用する必要がなくなります。拡張機能および完全な EIGRP ルーティングのために、スイッチは IP サービス イメージを実行している必要があります。IP ベース イメージが稼動しているスイッチで、マルチ VRF CE と EIGRP スタブ ルーティングを同時に設定しようとする場合、この設定は許可されません。

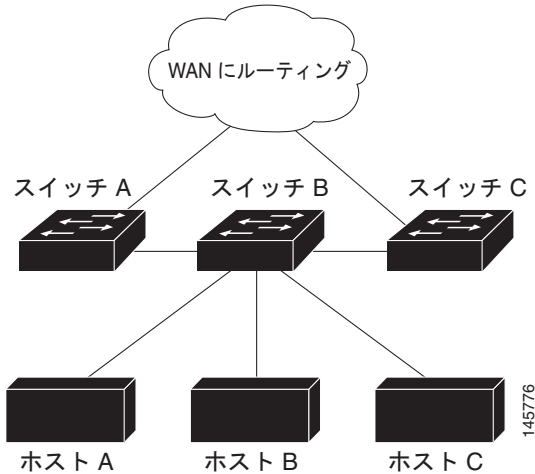
EIGRP スタブ ルーティングを使用するネットワークでは、ユーザへの IP トラフィックの許可ルートだけが EIGRP スタブ ルーティングを設定しているスイッチを通過します。スイッチは、ユーザ インターフェイスとして設定されているインターフェイスまたは他のデバイスに接続されているインターフェイスにルーテッド トラフィックを送信します。

EIGRP スタブ ルーティングを使用しているときは、EIGRP を使用してスイッチだけをスタブとして設定するように、分散ルータおよびリモート ルータを設定する必要があります。指定したルートだけがスイッチから伝播されます。スイッチは、サマリー、接続ルート、およびルーティング アップデートに対するすべてのクエリーに応答します。

スタブ ステータスを通知するパケットを受信するネイバーは、スタブ ルータのクエリーを実行せず、スタブ ピアを有するルータはそのピアのクエリーを実行しません。スタブ ルータは、分散ルータに依存してすべてのピアに適切なアップデートを送信します。

図 38-4 では、スイッチ B が EIGRP スタブ ルータとして設定されています。スイッチ A および C は残りの WAN に接続されています。スイッチ B は、接続ルート、スタティック ルート、再配信ルート、およびサマリー ルートをスイッチ A および C にアドバタイズします。スイッチ B は、スイッチ A から学習したルートをアドバタイズしません（その逆も同様）。

図 38-4 EIGRP スタブ ルータ設定



(注)

eigrp stub ルータ コンフィギュレーション コマンドを入力した後、**eigrp stub connected summary** コマンドだけが有効です。Command-Line Interface (CLI; コマンドライン インターフェイス) のヘルプが **receive-only** および **static** キーワードを表示し、これらのキーワードを入力したとしても、IP ベース イメージが稼動しているスイッチは、常に **connected** および **summary** キーワードが設定されているように動作します。

EIGRP スタブ ルーティングの詳細については、『Cisco IOS IP Configuration Guide, Volume 2 of 3: Routing Protocols, Release 12.2』の「Configuring EIGRP Stub Routing」を参照してください。このマニュアルには、Cisco.com の [Documentation] > [Cisco IOS Software] > [12.2 Mainline] > [Configuration Guides] からアクセスできます。

EIGRP のモニタおよびメンテナンス

ネイバー テーブルからネイバーを削除できます。さらに、各種 EIGRP ルーティング統計情報を表示することもできます。表 38-8 に、ネイバー削除および統計情報表示用の特権 EXEC コマンドを示しています。表示されるフィールドの詳細については、『Cisco IOS IP Command Reference, Volume 2 of 3: Routing Protocols, Release 12.2』を参照してください。このマニュアルには、Cisco.com の [Documentation] > [Cisco IOS Software] > [12.2 Mainline] > [Command References] からアクセスできます。

表 38-8 IP EIGRP の clear および show コマンド

コマンド	目的
<code>clear ip eigrp neighbors [if-address interface]</code>	ネイバー テーブルからネイバーを削除します。
<code>show ip eigrp interface [interface] [as number]</code>	EIGRP 用に設定されたインターフェイスの情報を表示します。
<code>show ip eigrp neighbors [type-number]</code>	EIGRP によって検出されたネイバーを表示します。

表 38-8 IP EIGRP の clear および show コマンド (続き)

コマンド	目的
<code>show ip eigrp topology [autonomous-system-number] [[ip-address] mask]</code>	指定されたプロセスの EIGRP トポロジテーブルを表示します。
<code>show ip eigrp traffic [autonomous-system-number]</code>	すべてまたは特定の EIGRP プロセスの送受信パケット数を表示します。

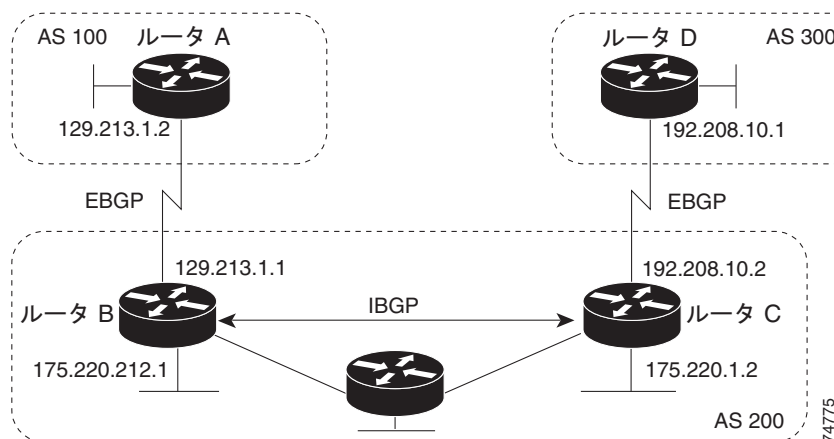
BGP の設定

Border Gateway Protocol (BGP; ボーダー ゲートウェイ プロトコル) は、Exterior Gateway Protocol (EGP; 外部ゲートウェイ プロトコル) です。Autonomous System (AS; 自律システム) 間で、ループの発生しないルーティング情報交換を保証するドメイン間ルーティング システムを設定するために使用されます。AS は、同じ管理下で動作して RIP や OSPF などの IGP を境界内で実行し、EGP を使用して相互接続されるルータで構成されます。BGP バージョン 4 は、インターネット内でドメイン間ルーティングを行うための標準 EGP です。このプロトコルは、RFC 1163、1267、および 1771 で定義されています。RIP の詳細については、Cisco Press 発行の『*Internet Routing Architectures*』および、『*Cisco IP and IP Routing Configuration Guide*』の「Configuring BGP」の章を参照してください。この資料には、Cisco.com の [Documentation] > [Cisco IOS Software] > [12.2 Mainline] > [Configuration Guides] からアクセスできます。

BGP コマンドおよびキーワードの詳細については、『*Cisco IOS IP Command Reference, Volume 2 of 3: Routing Protocols, Release 12.2*』の「IP Routing Protocols」を参照してください。このマニュアルには、Cisco.com の [Documentation] > [Cisco IOS Software] > [12.2 Mainline] > [Command References] からアクセスできます。表示されているにもかかわらずスイッチでサポートされない BGP コマンドについては、付録 D「Cisco IOS Release 12.2(55)SE でサポートされていないコマンド」を参照してください。

BGP アップデートを交換する場合、同じ AS に属するルータは *Internal BGP* (IBGP) を実行し、異なる AS に属するルータは *External BGP* (EBGP) を実行します。大部分のコンフィギュレーション コマンドは、EBGP と IBGP で同じですが、ルーティング アップデートが AS 間で交換されるか (EBGP)、または AS 内で交換されるか (IBGP) という点で異なります。図 38-5 に、EBGP と IBGP の両方が稼動するネットワークを示します。

図 38-5 EBGp、IBGP、および複数の AS



外部 AS と情報を交換する前に、BGP は AS 内のルータ間で内部 BGP ピアリングを定義し、IGRP や OSPF など AS 内で稼動する IGP に BGP ルーティング情報を再配信して、AS 内のネットワークに到達することを確認します。

BGP ルーティング プロセスを実行するルータは、通常 BGP スピーカーと呼ばれます。BGP はトランスポート プロトコルとして TCP を使用します (特にポート 179)。ルーティング情報を交換するため相互に TCP 接続された 2 つの BGP スピーカーを、ピアまたはネイバーと呼びます。図 38-5 では、ルータ A と B、ルータ B と C、およびルータ C と D がそれぞれ BGP ピアです。ルーティング情報は、宛先ネットワークへの完全パスを示す一連の AS 番号です。BGP はこの情報を使用し、ループのない AS マップを作成します。

このネットワークの特徴は次のとおりです。

- ルータ A および B では EBGp が、ルータ B および C では IBGP が稼動しています。EBGP ピアは直接接続されていますが、IBGP ピアは直接接続されていないことに注意してください。IGP が稼動し、2 つのネイバーが相互に到達する限り、IBGP ピアを直接接続する必要はありません。
- AS 内のすべての BGP スピーカーは、相互にピア関係を確立する必要があります。つまり、AS 内の BGP スピーカーは、論理的な完全メッシュ型に接続する必要があります。BGP4 は、論理的な完全メッシュに関する要求を軽減する 2 つの技術 (連合およびルート リフレクタ) を提供します。
- AS 200 は AS 100 および AS 300 の中継 AS です。つまり、AS 200 は AS 100 と AS 300 間でパケットを転送するために使用されます。

BGP ピアは完全な BGP ルーティング テーブルを最初に交換し、差分更新だけを送信します。BGP ピアはキープアライブ メッセージ (接続が有効であることを確認)、および通知メッセージ (エラーまたは特殊条件に回答) を交換することもできます。

BGP の場合、各ルートはネットワーク番号、情報が通過した AS のリスト (AS パス)、および他のパス アトリビュート リストで構成されます。BGP システムの主な機能は、AS パスのリストに関する情報など、ネットワークの到達可能性情報を他の BGP システムと交換することです。この情報は、AS が接続されているかどうかを判別したり、ルーティング ループをブリーニングしたり、AS レベル ポリシー判断を行うために使用できます。

Cisco IOS が稼動しているルータまたはスイッチが IBGP ルートを選択または使用するのには、ネクスト ホップ ルータで使用可能なルートがあり、IGP から同期信号を受信している (IGP 同期がディセーブルの場合は除く) 場合です。複数のルートが使用可能な場合、BGP はアトリビュート値に基づいてパスを選択します。BGP アトリビュートの詳細については、「BGP 判断アトリビュートの設定」(P.38-56) を参照してください。

BGP バージョン 4 では Classless Interdomain Routing (CIDR) がサポートされているため、集約ルートを作成してスーパーネットを構築し、ルーティング テーブルのサイズを削減できます。CIDR は、BGP 内部のネットワーク クラスの概念をエミュレートし、IP プレフィックスのアドバタイズメントをサポートします。

ここでは、次の設定情報について説明します。

- 「BGP のデフォルト設定」(P.38-49)
- 「BGP ルーティングのイネーブル化」(P.38-52)
- 「ルーティング ポリシー変更の管理」(P.38-54)
- 「BGP 判断アトリビュートの設定」(P.38-56)
- 「ルート マップによる BGP フィルタリングの設定」(P.38-58)
- 「ネイバーによる BGP フィルタリングの設定」(P.38-59)
- 「BGP フィルタリング用のプレフィックス リストの設定」(P.38-60)
- 「BGP コミュニティ フィルタリングの設定」(P.38-61)
- 「BGP ネイバーおよびピア グループの設定」(P.38-63)

- 「集約アドレスの設定」(P.38-65)
- 「ルーティング ドメイン連合の設定」(P.38-65)
- 「BGP ルート リフレクタの設定」(P.38-66)
- 「ルート ダンピング化の設定」(P.38-67)
- 「BGP のモニタおよびメンテナンス」(P.38-68)

BGP 設定の詳細については、『Cisco IOS IP Configuration Guide, Release 12.2』の「IP Routing Protocols」の「Configuring BGP」を参照してください。特定コマンドの詳細については、『Cisco IOS IP Command Reference, Volume 2 of 3: Routing Protocol, Release 12.2』を参照してください。これらのマニュアルは、Cisco.com ホームページ ([Documentation] > [Cisco IOS Software] > [12.2 Mainline] > [Configuration Guides] または [Command References]) にあります。

表示されているにもかかわらずスイッチでサポートされない BGP コマンドについては、付録 D 「Cisco IOS Release 12.2(55)SE でサポートされていないコマンド」を参照してください。

BGP のデフォルト設定

表 38-9 に、BGP の基本的なデフォルト設定を示します。すべての特性の詳細については、『Cisco IOS IP Command Reference, Volume 2 of 3: Routing Protocols, Release 12.2』の特定のコマンドを参照してください。

表 38-9 BGP のデフォルト設定

機能	デフォルト設定
集約アドレス	ディセーブル：未定義
AS パス アクセス リスト	未定義
自動サマリー	イネーブル
最適パス	- ルータはルートを選択する場合に AS パスを考慮します。外部 BGP ピアからの類似ルートは比較されません。 - ルータ ID の比較：ディセーブル
BGP コミュニティ リスト	- 番号：未定義 コミュニティ番号を示す特定の値を許可すると、許可されていないその他すべてのコミュニティ番号は、暗黙の拒否にデフォルト設定されます。 - フォーマット：シスコ デフォルト フォーマット (32 ビット番号)
BGP 連合 ID/ピア	- ID：未設定 - ピア：識別なし
BGP 高速外部フォールオーバー	イネーブル
BGP ローカル初期設定	100。指定できる範囲は 0 ～ 4294967295 です（大きな値を推奨）。
BGP ネットワーク	指定なし。バックドア ルートのアドバタイズメントなし。
BGP ルート ダンピング化	デフォルトでディセーブル イネーブルの場合は、次のようになります。 - 半減期は 15 分 - 再使用は 750（10 秒増分） - 抑制は 2000（10 秒増分） - 最大抑制時間は半減期の 4 倍（60 分）

表 38-9 BGP のデフォルト設定 (続き)

機能	デフォルト設定
BGP ルータ ID	ループバック インターフェイスに IP アドレスが設定されている場合は、ループバック インターフェイスの IP アドレス、またはルータの物理インターフェイスに対して設定された最大の IP アドレス
デフォルトの情報送信元 (プロトコルまたはネットワーク再配信)	ディセーブル
デフォルト メトリック	自動メトリック変換 (組み込み)
距離	- 外部ルート管理距離 : 20 (有効値は 1 ~ 255) - 内部ルート管理距離 : 200 (有効値は 1 ~ 255) - ローカル ルート管理距離 : 200 (有効値は 1 ~ 255)
ディストリビュート リスト	- 入力 (アップデート中に受信されたネットワークをフィルタリング) : ディセーブル - 出力 (アップデート中のネットワークのアドバタイズを抑制) : ディセーブル
内部ルート再配信	ディセーブル
IP プレフィクス リスト	未定義
Multi Exit Discriminator (MED)	- 常に比較 : ディセーブル 異なる AS 内のネイバーからのパスに対して、MED を比較しません。 - 最適パスの比較 : ディセーブル - 最悪パスである MED の除外 : ディセーブル - 決定的な MED 比較 : ディセーブル

表 38-9 BGP のデフォルト設定 (続き)

機能	デフォルト設定
ネイバー	• アドバタイズメント インターバル：外部ピアの場合は 30 秒、内部ピアの場合は 5 秒 • ロギング変更：イネーブル • 条件付きアドバタイズメント：ディセーブル • デフォルト送信元：ネイバーに送信されるデフォルト ルートはなし • 説明：なし • ディストリビュート リスト：未定義 • 外部 BGP マルチホップ：直接接続されたネイバーだけを許可 • フィルタ リスト：使用しない • 受信したプレフィックスの最大数：制限なし • ネクスト ホップ (BGP ネイバーのネクスト ホップとなるルータ)：ディセーブル • パスワード：ディセーブル • ピア グループ：定義なし。割り当てメンバーなし • プレフィックス リスト：指定なし • リモート AS (ネイバー BGP テーブルへのエントリ追加)：ピア定義なし • プライベート AS 番号の削除：ディセーブル • ルート マップ：ピアへの適用なし • コミュニティ アトリビュート送信：ネイバーへの送信なし • シャットダウンまたはソフト再設定：ディセーブル • タイマー：キープアライブ：60 秒。ホールドタイム：180 秒。 • アップデート送信元：最適ローカル アドレス • バージョン：BGP バージョン 4 • ウェイト：BGP ピアによって学習されたルート：0。ローカル ルータから取得されたルート：32768。
NSF ¹ 認識	ディセーブル レイヤ 3 スイッチでは、ハードウェアやソフトウェアの変更中に、ネイバー NSF 対応ルータからのパケットを転送し続けることができます。 (注) NSF 認識は、グレースフル リスタートをイネーブルすることにより、IP サービス イメージが稼動しているスイッチの IPv4 に対してイネーブルにできます。
ルート リフレクタ	未設定
同期化 (BGP および IGP)	イネーブル
テーブル マップ アップデート	ディセーブル
タイマー	キープアライブ：60 秒。ホールドタイム：180 秒。

1. NSF = Nonstop Forwarding

NSF 認識

BGP NSF 認識機能が IP サービス イメージの IPv4 でサポートされます。BGP ルーティングでこの機能をイネーブルにするには、グレースフル リスタートをイネーブルにする必要があります。ネイバー ルータが NSF 対応で、この機能がイネーブルである場合、レイヤ 3 スイッチでは、ルータに障害が発

生してプライマリ RP がバックアップ RP によって引き継がれる間、または処理を中断させずにソフトウェア アップグレードを行うためにプライマリ RP を手動でリロードしている間、ネイバー ルータからパケットを転送し続けます。

詳細については、次の URL にある『*BGP Nonstop Forwarding (NSF) Awareness Feature Guide*』を参照してください。
http://www.cisco.com/en/US/products/sw/iosswrel/ps1839/products_feature_guide09186a008015fed.html

BGP ルーティングのイネーブル化

BGP ルーティングをイネーブルにするには、BGP ルーティング プロセスを確立し、ローカル ネットワークを定義します。BGP はネイバーとの関係を完全に認識する必要があるため、BGP ネイバーも指定する必要があります。

BGP は、内部および外部の 2 種類のネイバーをサポートします。内部ネイバーは同じ AS 内に、外部ネイバーは異なる AS 内にあります。通常の場合、外部ネイバーは相互に隣接し、1 つのサブネットを共有しますが、内部ネイバーは同じ AS 内の任意の場所に存在します。

スイッチではプライベート AS 番号を使用できます。プライベート AS 番号は通常サービス プロバイダーによって割り当てられ、ルートが外部ネイバーにアダプタイズされないシステムに設定されます。プライベート AS 番号の範囲は 64512 ~ 65535 です。AS パスからプライベート AS 番号を削除するように外部ネイバーを設定するには、neighbor remove-private-as ルータ コンフィギュレーション コマンドを使用します。この結果、外部ネイバーにアップデートを渡すとき、AS パス内にプライベート AS 番号が含まれている場合は、これらの番号が削除されます。

AS が別の AS からさらに別の AS にトラフィックを渡す場合は、アダプタイズメント対象のルートに矛盾が存在しないことが重要です。BGP がルートをアダプタイズしてから、ネットワーク内のすべてのルータが IGP を通してルートを学習した場合、AS は一部のルータがルーティングできなかったトラフィックを受信することがあります。このような事態を避けるため、BGP は IGP が AS に情報を伝播し、BGP が IGP と同期化されるまで、待機する必要があります。同期化は、デフォルトでイネーブルに設定されています。AS が特定の AS から別の AS にトラフィックを渡さない場合、または AS 内のすべてのルータで BGP が稼動している場合は、同期化をディセーブルにし、IGP 内で伝送されるルート数を少なくして、BGP がより短時間で収束するようにします。



(注) BGP をイネーブルにするには、スタック マスター上で IP サービス イメージが稼動している必要があります。

BGP ルーティングをイネーブルにして BGP ルーティング プロセスを確立し、ネイバーを指定するには、特権 EXEC モードで次の手順を実行します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	ip routing	IP ルーティングをイネーブルにします (IP ルーティングがディセーブルになっている場合に限り必須)。
ステップ 3	router bgp autonomous-system	BGP ルーティング プロセスをイネーブルにして AS 番号を割り当て、ルータ コンフィギュレーション モードを開始します。指定できる AS 番号は 1 ~ 65535 です。64512 ~ 65535 は、プライベート AS 番号専用です。
ステップ 4	network network-number [mask network-mask] [route-map route-map-name]	この AS に対してローカルとなるようにネットワークを設定し、BGP テーブルにネットワークを格納します。

	コマンド	目的
ステップ 5	neighbor { <i>ip-address</i> <i>peer-group-name</i> } remote-as <i>number</i>	BGP ネイバー テーブルに設定を追加し、IP アドレスによって識別されるネイバーが、指定された AS に属することを示します。 EBGP の場合、通常ネイバーは直接接続されており、IP アドレスは接続のもう一方の端におけるインターフェイスのアドレスです。 IBGP の場合、IP アドレスにはルータ インターフェイス内の任意のアドレスを指定できます。
ステップ 6	neighbor { <i>ip-address</i> <i>peer-group-name</i> } remove-private-as	(任意) 発信ルーティング アップデート内の AS パスからプライベート AS 番号を削除します。
ステップ 7	no synchronization	(任意) BGP と IGP の同期化をディセーブルにします。
ステップ 8	no auto-summary	(任意) 自動ネットワーク サマライズをディセーブルにします。デフォルトでは、IGP から BGP にサブネットが再配信された場合、ネットワーク ルートだけが BGP テーブルに追加されます。
ステップ 9	bgp fast-external-fallover	(任意) 外部ネイバー間のリンクが切断された場合、BGP セッションを自動的にリセットします。デフォルトで、セッションは即座にリセットされません。
ステップ 10	bgp graceful-restart	(任意) NSF 認識をスイッチでイネーブルにします。NSF 認識はデフォルトではディセーブルです。
ステップ 11	end	特権 EXEC モードに戻ります。
ステップ 12	show ip bgp network <i>network-number</i> または show ip bgp neighbor	設定を確認します。 NSF 認識（グレースフル リスタート）がネイバーでイネーブルにされていることを確認します。 スイッチおよびネイバーで NSF 認識がイネーブルである場合は、次のメッセージが表示されます。 <i>Graceful Restart Capability: advertised and received</i> スイッチで NSF 認識がイネーブルであり、ネイバーでディセーブルである場合は、次のメッセージが表示されます。 <i>Graceful Restart Capability: advertised</i>
ステップ 13	copy running-config startup-config	(任意) コンフィギュレーション ファイルに設定を保存します。

BGP AS を削除するには、**no router bgp autonomous-system** グローバル コンフィギュレーション コマンドを使用します。BGP テーブルからネットワークを削除するには、**no network network-number** ルータ コンフィギュレーション コマンドを使用します。ネイバーを削除するには、**no neighbor {ip-address | peer-group-name} remote-as number** ルータ コンフィギュレーション コマンドを使用します。ネイバーにアップデート内のプライベート AS 番号を追加するには、**no neighbor {ip-address | peer-group-name} remove-private-as** ルータ コンフィギュレーション コマンドを使用します。同期化を再度イネーブルにするには、**synchronization** ルータ コンフィギュレーション コマンドを使用します。

次に、図 38-5 に示されたルータ上で BGP を設定する例を示します。

ルータ A :

```
Switch(config)# router bgp 100
Switch(config-router)# neighbor 129.213.1.1 remote-as 200
```

ルータ B :

```
Switch(config)# router bgp 200
Switch(config-router)# neighbor 129.213.1.2 remote-as 100
Switch(config-router)# neighbor 175.220.1.2 remote-as 200
```

ルータ C :

```
Switch(config)# router bgp 200
Switch(config-router)# neighbor 175.220.212.1 remote-as 200
Switch(config-router)# neighbor 192.208.10.1 remote-as 300
```

ルータ D :

```
Switch(config)# router bgp 300
Switch(config-router)# neighbor 192.208.10.2 remote-as 200
```

BGP ピアが稼動していることを確認するには、**show ip bgp neighbors** 特権 EXEC コマンドを使用します。次に、ルータ A にこのコマンドを実行した場合の出力例を示します。

```
Switch# show ip bgp neighbors
```

```
BGP neighbor is 129.213.1.1, remote AS 200, external link
  BGP version 4, remote router ID 175.220.212.1
  BGP state = established, table version = 3, up for 0:10:59
  Last read 0:00:29, hold time is 180, keepalive interval is 60 seconds
  Minimum time between advertisement runs is 30 seconds
  Received 2828 messages, 0 notifications, 0 in queue
  Sent 2826 messages, 0 notifications, 0 in queue
  Connections established 11; dropped 10
```

state = established 以外の情報が出力された場合、ピアは稼動していません。リモート ルータ ID は、ルータ（または最大のループバック インターフェイス）上の最大の IP アドレスです。テーブルが新規情報でアップデートされるたびに、テーブルのバージョン番号は増加します。継続的にテーブルバージョン番号が増加している場合は、ルートがフラッピングし、ルーティング アップデートが絶えず発生しています。

外部プロトコルの場合、**network** ルータ コンフィギュレーション コマンドから IP ネットワークへの参照によって制御されるのは、アドバタイズされるネットワークだけです。これは、**network** コマンドを使用してアップデートの送信先を指定する IGP（EIGRP など）と対照的です。

BGP 設定の詳細については、『*Cisco IOS IP Configuration Guide, Release 12.2*』の「IP Routing Protocols」を参照してください。特定コマンドの詳細については、『*Cisco IOS IP Command Reference, Volume 2 of 3: Routing Protocol, Release 12.2*』を参照してください。表示されているにもかかわらずスイッチでサポートされない BGP コマンドについては、付録 D「Cisco IOS Release 12.2(55)SE でサポートされていないコマンド」を参照してください。

ルーティング ポリシー変更の管理

ピアのルーティング ポリシーには、着信または発信ルーティング テーブル アップデートに影響する可能性があるすべての設定が含まれます。BGP ネイバーとして定義された 2 台のルータは、BGP 接続を形成し、ルーティング情報を交換します。この後で BGP フィルタ、ウェイト、距離、バージョン、またはタイマーを変更する場合、または同様の設定変更を行う場合は、BGP セッションをリセットし、設定の変更を有効にする必要があります。

リセットには、ハードリセットとソフトリセットの 2 つのタイプがあります。事前に設定を行わなくても、ソフトリセットを使用できます。事前設定なしにソフトリセットを使用するには、両方の BGP ピアでソフトルートリフレッシュ機能がサポートされていなければなりません。この機能は、ピアによって TCP セッションが確立されたときに送信される OPEN メッセージに格納されてアドバタイズされます。ソフトリセットを使用すると、BGP ルータ間でルートリフレッシュ要求およびルーティング情報を動的に交換したり、それぞれの発信ルーティングテーブルを後で再アドバタイズできます。

- ソフトリセットによってネイバーから着信アップデートが生成された場合、このリセットは **ダイナミック着信ソフトリセット** といいます。
- ソフトリセットによってネイバーに一連のアップデートが送信された場合、このリセットは **発信ソフトリセット** といいます。

ソフト着信リセットが発生すると、新規着信ポリシーが有効になります。ソフト発信リセットが発生すると、BGP セッションがリセットされずに、新規ローカル発信ポリシーが有効になります。発信ポリシーのリセット中に新しい一連のアップデートが送信されると、新規着信ポリシーも有効になる場合があります。

表 38-10 に、ハードリセットとソフトリセットの利点および欠点を示します。

表 38-10 ハードリセットとソフトリセットの利点および欠点

リセットタイプ	利点	欠点
ハードリセット	メモリ オーバーヘッドが発生しません。	ネイバーから提供された BGP、IP、および Forwarding Information Base (FIB; 転送情報ベース) テーブルのプレフィックスが失われます。推奨しません。
発信ソフトリセット	ルーティングテーブルアップデートが設定、保管されません。	着信ルーティングテーブルアップデートがリセットされません。
ダイナミック着信ソフトリセット	BGP セッションおよびキャッシュがクリアされません。 ルーティングテーブルアップデートを保管する必要がなく、メモリ オーバーヘッドが発生しません。	両方の BGP ルータでルートリフレッシュ機能をサポートする必要があります。

BGP ピアがルートリフレッシュ機能をサポートするかどうかを学習して、BGP セッションをリセットするには、特権 EXEC モードで次の手順を実行します。

	コマンド	目的
ステップ 1	show ip bgp neighbors	ネイバーがルートリフレッシュ機能をサポートするかどうかを表示します。サポートされている場合は、ルータに関する次のメッセージが表示されます。 <i>Received route refresh capability from peer.</i>
ステップ 2	clear ip bgp { * address peer-group-name }	指定された接続上でルーティングテーブルをリセットします。 • すべての接続をリセットする場合は、アスタリスク (*) を入力します。 • 特定の接続をリセットする場合は、IP アドレスを入力します。 • ピアグループをリセットする場合は、ピアグループ名を入力します。

コマンド	目的
ステップ 3 clear ip bgp {* <i>address</i> <i>peer-group-name</i> } soft out	(任意) 指定された接続上で着信ルーティング テーブルをリセットするには、発信ソフト リセットを実行します。このコマンドは、ルート リフレッシュがサポートされている場合に使用してください。 - すべての接続をリセットする場合は、アスタリスク (*) を入力します。 - 特定の接続をリセットする場合は、IP アドレスを入力します。 - ピア グループをリセットする場合は、ピア グループ名を入力します。
ステップ 4 show ip bgp show ip bgp neighbors	ルーティング テーブルおよび BGP ネイバーに関する情報をチェックし、リセットされたことを確認します。

BGP 判断アトリビュートの設定

BGP スピーカーが複数の AS から受信したアップデートが、同じ宛先に対して異なるパスを示している場合、BGP スピーカーはその宛先に到達する最適パスを 1 つ選択する必要があります。選択されたパスは BGP ルーティング テーブルに格納され、ネイバーに伝播されます。この判断は、アップデートに格納されているアトリビュート値、および BGP で設定可能な他の要因に基づいて行われます。

BGP ピアはネイバー AS からプレフィクスに対する 2 つの EBGП パスを学習するとき、最適パスを選択して IP ルーティング テーブルに挿入します。BGP マルチパス サポートがイネーブルで、同じネイバー AS から複数の EBGП パスを学習する場合、単一の最適パスの代わりに、複数のパスが IP ルーティング テーブルに格納されます。その後、パケット スイッチング中に、複数のパス間でパケット単位または宛先単位のロード バランシングが実行されます。**maximum-paths** ルータ コンフィギュレーション コマンドは、許可されるパス数を制御します。

これらの要因により、BGP が最適パスを選択するためにアトリビュートを評価する順序が決まります。

1. パスで指定されているネクスト ホップが到達不能な場合、このアップデートは削除されます。BGP のネクスト ホップのアトリビュート (ソフトウェアによって自動判別される) は、宛先に到達するために使用されるネクスト ホップの IP アドレスです。EBGP の場合、通常このアドレスは **neighbor remote-as** ルータ コンフィギュレーション コマンドで指定されたネイバーの IP アドレスです。ネクスト ホップの処理をディセーブルにするには、ルート マップまたは **neighbor next-hop-self** ルータ コンフィギュレーション コマンドを使用します。
2. 最大ウェイトのパスを推奨します (シスコ独自のパラメータ)。ウェイト アトリビュートはルータにローカルであるため、ルーティング アップデートで伝播されません。デフォルトでは、ルータ送信元のパスに関するウェイト アトリビュートは 32768 で、それ以外のパスのウェイト アトリビュートは 0 です。最大ウェイトのルートを推奨します。ウェイトを設定するには、アクセス リスト、ルート マップ、または **neighbor weight** ルータ コンフィギュレーション コマンドを使用します。
3. ローカル初期設定値が最大のルートを推奨します。ローカル初期設定はルーティング アップデートに含まれ、同じ AS 内のルータ間で交換されます。ローカル初期設定アトリビュートのデフォルト値は 100 です。ローカル初期設定を設定するには、**bgp default local-preference** ルータ コンフィギュレーション コマンドまたはルート マップを使用します。
4. ローカル ルータ上で稼動する BGP から送信されたルートを推奨します。
5. AS パスが最短のルートを推奨します。
6. 送信元タイプが最小のルートを推奨します。内部ルートまたは IGP は、EGP によって学習されたルートよりも小さく、EGP で学習されたルートは、未知の送信元のルートまたは別の方法で学習されたルートよりも小さくなります。

7. 想定されるすべてのルートについてネイバー AS が同じである場合は、MED メトリック アトリビュートが最小のルートを推奨します。MED を設定するには、ルート マップまたは **default-metric** ルータ コンフィギュレーション コマンドを使用します。IBGP ピアに送信されるアップデートには、MED が含まれます。
8. 内部 (IBGP) パスより、外部 (EBGP) パスを推奨します。
9. 最も近い IGP ネイバー (最小の IGP メトリック) を通って到達できるルートを推奨します。ルータは、AS 内の最短の内部パス (BGP のネクスト ホップへの最短パス) を使用し、宛先に到達するためです。
10. 次の条件にすべて該当する場合は、このパスのルートを IP ルーティング テーブルに挿入してください。
 - 最適ルートと目的のルートがともに外部ルートである
 - 最適ルートと目的のルートの両方が、同じネイバー AS からのルートである
 - maximum-paths がイネーブルである
11. マルチパスがイネーブルでない場合は、BGP ルータ ID の IP アドレスが最小であるルートを推奨します。通常、ルータ ID はルータ上の最大の IP アドレスまたはループバック (仮想) アドレスですが、実装に依存することがあります。

同じ判断アトリビュートを設定するには、特権 EXEC モードで次の手順を実行します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	router bgp autonomous-system	BGP ルーティング プロセスをイネーブルにして AS 番号を割り当て、ルータ コンフィギュレーション モードを開始します。
ステップ 3	bgp best-path as-path ignore	(任意) ルート選択中に AS パス長を無視するようにルータを設定します。
ステップ 4	neighbor {ip-address peer-group-name} next-hop-self	(任意) ネクスト ホップ アドレスの代わりに使用される特定の IP アドレスを入力し、ネイバーへの BGP アップデートに関するネクスト ホップの処理をディセーブルにします。
ステップ 5	neighbor {ip-address peer-group-name} weight weight	(任意) ネイバー接続にウェイトを割り当てます。指定できる値は 0 ～ 65535 です。最大ウェイトのルートを推奨します。別の BGP ピアから学習されたルートのデフォルト ウェイトは 0 です。ローカル ルータから送信されたルートのデフォルト ウェイトは 32768 です。
ステップ 6	default-metric number	(任意) 推奨パスを外部ネイバーに設定するように MED メトリックを設定します。MED を持たないすべてのルータも、この値に設定されます。指定できる範囲は 1 ～ 4294967295 です。最小値を推奨します。
ステップ 7	bgp bestpath med missing-as-worst	(任意) MED がない場合は無限の値が指定されていると見なし、MED 値を持たないパスが最も望ましくないパスになるように、スイッチを設定します。
ステップ 8	bgp always-compare med	(任意) 異なる AS 内のネイバーからのパスに対して、MED を比較するようにスイッチを設定します。デフォルトでは、MED は同じ AS 内のパス間だけで比較されます。
ステップ 9	bgp bestpath med confed	(任意) 連合内の異なるサブ AS によってアドバタイズされたパスから特定のパスを選択する場合に、MED を考慮するようにスイッチを設定します。

	コマンド	目的
ステップ 10	bgp deterministic med	(任意) 同じ AS 内の異なるピアによってアドパタイズされたルートから選択する場合に、MED 変数を考慮するようにスイッチを設定します。
ステップ 11	bgp default local-preference value	(任意) デフォルトのローカル初期設定値を変更します。指定できる範囲は 0 ～ 4294967295 です。デフォルトは 100 です。最大のローカル初期設定値を推奨します。
ステップ 12	maximum-paths number	(任意) IP ルーティング テーブルに追加するパスの数を設定します。デフォルトでは、最適パスだけがルーティング テーブルに追加されます。指定できる範囲は 1 ～ 16 です。複数の値を指定すると、パス間のロード バランシングが可能になります (スイッチ ソフトウェア では最大 32 の等価コスト ルーティングが許可されていますが、スイッチ ハードウェアはルートあたり 17 パス以上は使用しません)。
ステップ 13	end	特権 EXEC モードに戻ります。
ステップ 14	show ip bgp show ip bgp neighbors	ルーティング テーブルおよび BGP ネイバーに関する情報をチェックし、リセットされたことを確認します。
ステップ 15	copy running-config startup-config	(任意) コンフィギュレーション ファイルに設定を保存します。

デフォルト ステートに戻すには、このコマンドの **no** 形式を使用します。

ルート マップによる BGP フィルタリングの設定

BGP 内でルート マップを使用すると、ルーティング情報を制御、変更したり、ルーティング ドメイン間でルートを再配信する条件を定義できます。ルート マップの詳細については、「[ルート マップによるルーティング情報の再配信](#)」(P.38-99) を参照してください。各ルート マップには、ルート マップを識別する名前 (マップ タグ) およびオプションのシーケンス番号が付いています。

ルート マップを使用してネクスト ホップ処理をディセーブルにするには、特権 EXEC モードで次の手順を実行します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	route-map map-tag [[permit deny] sequence-number]]	ルート マップを作成し、ルート マップ コンフィギュレーション モードを開始します。
ステップ 3	set ip next-hop ip-address [...ip-address] [peer-address]	(任意) ネクスト ホップ処理をディセーブルにするようにルート マップを設定します。 - 着信ルート マップの場合は、一致するルートのネクスト ホップをネイバー ピア アドレスに設定し、サードパーティのネクスト ホップを上書きします。 - BGP ピアの発信ルート マップの場合は、ネクスト ホップをローカル ルータのピア アドレスに設定して、ネクスト ホップ計算をディセーブルにします。
ステップ 4	end	特権 EXEC モードに戻ります。

	コマンド	目的
ステップ 5	show route-map <i>[map-name]</i>	設定を確認するため、設定されたすべてのルート マップ、または指定されたルート マップだけを表示します。
ステップ 6	copy running-config startup-config	(任意) コンフィギュレーション ファイルに設定を保存します。

ルート マップを削除するには、**no route-map map-tag** コマンドを使用します。ネクスト ホップ処理を再びイネーブルにするには、**no set ip next-hop ip-address** コマンドを使用します。

ネイバーによる BGP フィルタリングの設定

BGP アドバタイズメントをフィルタリングするには、**as-path access-list** グローバル コンフィギュレーション コマンドや **neighbor filter-list** ルータ コンフィギュレーション コマンドなどの AS パス フィルタを使用します。**neighbor distribute-list** ルータ コンフィギュレーション コマンドとアクセス リストを併用することもできます。distribute-list フィルタはネットワーク番号に適用されます。**distribute-list** コマンドの詳細については、「[ルーティング アップデートのアドバタイズメントおよび処理の制御](#)」(P.38-107) を参照してください。

ネイバー単位でルート マップを使用すると、アップデートをフィルタリングしたり、各アトリビュートを変更できます。ルート マップは、着信アップデートまたは発信アップデートのいずれかに適用できます。ルート マップを渡すルートだけが、アップデート内で送信または許可されます。着信および発信の両方のアップデートで、AS パス、コミュニティ、およびネットワーク番号に基づくマッチングがサポートされています。AS パスのマッチングには **match as-path access-list** ルート マップ コマンド、コミュニティに基づくマッチングには **match community-list** ルート マップ コマンド、ネットワークに基づくマッチングには **ip access-list** グローバル コンフィギュレーション コマンドが必要です。

ネイバー単位のルート マップを適用するには、特権 EXEC モードで次の手順を実行します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	router bgp <i>autonomous-system</i>	BGP ルーティング プロセスをイネーブルにして AS 番号を割り当て、ルータ コンフィギュレーション モードを開始します。
ステップ 3	neighbor <i>{ip-address peer-group name}</i> distribute-list <i>{access-list-number name}</i> <i>{in out}</i>	(任意) アクセス リストの指定に従って、ネイバーに対して送受信される BGP ルーティング アップデートをフィルタリングします。 (注) neighbor prefix-list ルータ コンフィギュレーション コマンドを使用して、アップデートをフィルタリングすることもできますが、両方のコマンドを使用して同じ BGP ピアを設定することはできません。
ステップ 4	neighbor <i>{ip-address peer-group name}</i> route-map <i>map-tag</i> <i>{in out}</i>	(任意) ルート マップを適用し、着信または発信ルートをフィルタリングします。
ステップ 5	end	特権 EXEC モードに戻ります。
ステップ 6	show ip bgp neighbors	設定を確認します。
ステップ 7	copy running-config startup-config	(任意) コンフィギュレーション ファイルに設定を保存します。

ネイバーからアクセス リストを削除するには、**no neighbor distribute-list** コマンドを使用します。ネイバーからルート マップを削除するには、**no neighbor route-map map-tag** ルータ コンフィギュレーション コマンドを使用します。

BGP AS パスに基づいて着信および発信の両方のアップデートにアクセス リスト フィルタを指定して、フィルタリングすることもできます。各フィルタは、正規表現に基づくアクセス リストです（正規表現の作成方法については、『Cisco IOS Dial Technologies Command Reference, Release 12.2』の付録「Regular Expressions」を参照してください）。この方法を使用するには、AS パスのアクセス リストを定義し、特定のネイバーに対して送受信されるアップデートに適用します。

BGP パス フィルタリングを設定するには、特権 EXEC モードで次の手順を実行します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	ip as-path access-list access-list-number {permit deny} as-regular-expressions	BGP 関連アクセス リストを定義します。
ステップ 3	router bgp autonomous-system	BGP ルータ コンフィギュレーション モードを開始します。
ステップ 4	neighbor {ip-address peer-group name} filter-list {access-list-number name} {in out weight weight}	アクセス リストに基づいて、BGP フィルタを確立します。
ステップ 5	end	特権 EXEC モードに戻ります。
ステップ 6	show ip bgp neighbors [paths regular-expression]	設定を確認します。
ステップ 7	copy running-config startup-config	(任意) コンフィギュレーション ファイルに設定を保存します。

BGP フィルタリング用のプレフィクス リストの設定

neighbor distribute-list ルータ コンフィギュレーション コマンドを含む多数の BGP ルート フィルタリング コマンドでは、アクセス リストの代わりにプレフィクス リストを使用できます。プレフィクス リストを使用すると、大規模リストのロードおよび検索パフォーマンスが改善し、差分更新がサポートされ、CLI（コマンドライン インターフェイス）設定が簡素化され、柔軟性が増すなどの利点が生じます。

プレフィクス リストによるフィルタリングでは、アクセス リストの照合の場合と同様に、プレフィクス リストに記載されたプレフィクスとルートとのプレフィクスが照合されます。一致が存在する場合は、一致したルートが使用されます。プレフィクスが許可されるか、または拒否されるかは、次に示す規則に基づいて決定されます。

- 空のプレフィクス リストはすべてのプレフィクスを許可します。
- 指定されたプレフィクスがプレフィクス リスト内のどのエン트리とも一致しない場合は、暗黙の拒否が使用されます。
- 指定されたプレフィクスと一致するエントリがプレフィクス リスト内に複数存在する場合は、シーケンス番号が最小であるプレフィクス リスト エントリが識別されます。

デフォルトでは、シーケンス番号は自動生成され、5 ずつ増分します。シーケンス番号の自動生成をディセーブルにした場合は、エントリごとにシーケンス番号を指定する必要があります。シーケンス番号を指定する場合の増分値に制限はありません。増分値が 1 の場合は、このリストに追加エントリを挿入できません。増分値が大きい場合は、値がなくなることがあります。

コンフィギュレーション エントリを削除する場合は、シーケンス番号を指定する必要はありません。**show** コマンドの出力には、シーケンス番号が含まれます。

コマンド内でプレフィクス リストを使用する場合は、あらかじめプレフィクス リストを設定しておく必要があります。プレフィクス リストを作成したり、プレフィクス リストにエントリを追加するには、特権 EXEC モードで次の手順を実行します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	ip prefix-list list-name [seq seq-value] deny permit network/len [ge ge-value] [le le-value]	一致条件のために、アクセスを拒否 (deny) または許可 (permit) するプレフィクス リストを作成します。シーケンス番号を指定することもできます。少なくとも 1 つの permit コマンドまたは deny コマンドを入力する必要があります。 network/len は、ネットワーク番号およびネットワーク マスク の長さ (ビット単位) です。 (任意) ge および le の値は、照合するプレフィクス長の範囲を指定します。指定された ge-value および le-value は、次の条件を満たす必要があります。$len < ge-value < le-value < 32$
ステップ 3	ip prefix-list list-name seq seq-value deny permit network/len [ge ge-value] [le le-value]	(任意) プレフィクス リストにエントリを追加し、そのエントリにシーケンス番号を割り当てます。
ステップ 4	end	特権 EXEC モードに戻ります。
ステップ 5	show ip prefix list [detail summary] name [network/len] [seq seq-num] [longer] [first-match]	プレフィクス リストまたはプレフィクス リスト エントリに関する情報を表示して、設定を確認します。
ステップ 6	copy running-config startup-config	(任意) コンフィギュレーション ファイルに設定を保存します。

プレフィクス リストまたはそのエントリをすべて削除する場合は、**no ip prefix-list list-name** グローバル コンフィギュレーション コマンドを使用します。プレフィクス リストから特定のエントリを削除する場合は、**no ip prefix-list seq seq-value** グローバル コンフィギュレーション コマンドを使用します。シーケンス番号の自動生成をディセーブルにするには **no ip prefix-list sequence number** コマンドを、自動生成を再びイネーブルにするには **ip prefix-list sequence number** コマンドを使用します。プレフィクス リスト エントリのヒット数テーブルをクリアするには、**clear ip prefix-list** 特権 EXEC コマンドを使用します。

BGP コミュニティ フィルタリングの設定

BGP コミュニティ フィルタリングは、COMMUNITIES アトリビュートの値に基づいてルーティング情報の配信を制御する BGP の方法の 1 つです。このアトリビュートによって、宛先はコミュニティにグループ化され、コミュニティに基づいてルーティング判断が適用されます。この方法を使用すると、ルーティング情報の配信制御を目的とする BGP スピーカーの設定が簡単になります。

コミュニティは、共通するいくつかのアトリビュートを共有する宛先のグループです。各宛先は複数のコミュニティに属します。AS 管理者は、宛先が属するコミュニティを定義できます。デフォルトでは、すべての宛先が一般的なインターネット コミュニティに属します。コミュニティは、過渡的でグローバルな、オプションの COMMUNITIES アトリビュート (1 ~ 4294967200) によって識別されません。事前に定義された既知のコミュニティの一部を、次に示します。

- **internet** : このルートをインターネット コミュニティにアドバタイズします。すべてのルータが所属します。
- **no-export** : EBGp ピアにこのルートをアドバタイズしません。
- **no-advertise** : いずれのピア (内部または外部) にもこのルートをアドバタイズしません。
- **local-as** : ローカルな AS 外部のピアにこのルートをアドバタイズしません。

コミュニティに基づき、他のネイバーに許可、送信、配信するルーティング情報を制御できます。BGP スピーカーは、ルートを学習、アドバタイズ、または再配信するときに、ルートのコミュニティを設定、追加、または変更します。ルートを集約すると、作成された集約内の COMMUNITIES アトリビュートに、すべての初期ルートの全コミュニティが含まれます。

コミュニティ リストを使用すると、ルート マップの **match** ステートメントで使用されるコミュニティ グループを作成できます。さらに、アクセス リストの場合と同様、一連のコミュニティ リストを作成することもできます。ステートメントは一致が見つかるまでチェックされ、1 つのステートメントが満たされると、テストは終了します。

コミュニティに基づいて COMMUNITIES アトリビュートおよび **match** コマンドを設定するには、「[ルート マップによるルーティング情報の再配信](#)」(P.38-99) に記載されている **match community-list** および **set community** ルート マップ コンフィギュレーション コマンドを参照してください。

デフォルトでは、COMMUNITIES アトリビュートはネイバーに送信されません。COMMUNITIES アトリビュートが特定の IP アドレスのネイバーに送信されるように指定するには、**neighbor send-community** ルータ コンフィギュレーション コマンドを使用します。

コミュニティ リストを作成、適用するには、特権 EXEC モードで次の手順を実行します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	ip community-list <i>community-list-number {permit deny}</i> <i>community-number</i>	コミュニティ リストを作成し、番号を割り当てます。 <i>community-list-number</i> は 1 ～ 99 の整数です。この値は、コミュニティの許可または拒否グループを 1 つまたは複数識別します。 <i>community-number</i> は、set community ルートマップ コンフィギュレーション コマンドで設定される番号です。
ステップ 3	router bgp autonomous-system	BGP ルータ コンフィギュレーション モードを開始します。
ステップ 4	neighbor {ip-address peer-group name} send-community	この IP アドレスのネイバーに送信する COMMUNITIES アトリビュートを指定します。
ステップ 5	set comm-list list-num delete	(任意) ルート マップで指定された標準または拡張コミュニティ リストと一致する着信または発信アップデートのコミュニティ アトリビュートから、コミュニティを削除します。
ステップ 6	exit	グローバル コンフィギュレーション モードに戻ります。
ステップ 7	ip bgp-community new-format	(任意) AA:NN のフォーマットで、BGP コミュニティを表示、解析します。 BGP コミュニティは、2 つの部分からなる 2 バイト長フォーマットで表示されます。シスコのデフォルトのコミュニティ フォーマットは NNAA です。BGP に関する最新の RFC では、コミュニティは AA:NN の形式を取ります。最初の部分は AS 番号で、その次の部分は 2 バイトの数値です。
ステップ 8	end	特権 EXEC モードに戻ります。
ステップ 9	show ip bgp community	設定を確認します。
ステップ 10	copy running-config startup-config	(任意) コンフィギュレーション ファイルに設定を保存します。

BGP ネイバーおよびピア グループの設定

通常、BGP ネイバーの多くは同じアップデート ポリシー（同じ発信ルート マップ、配信リスト、フィルタ リスト、アップデート送信元など）を使用して設定されます。アップデート ポリシーが同じネイバーをピア グループにまとめると設定が簡単になり、アップデートの効率が高まります。多数のピアを設定した場合は、この方法を推奨します。

BGP ピア グループを設定するには、ピア グループを作成し、そこにオプションを割り当てて、ピア グループ メンバーとしてネイバーを追加します。ピア グループを設定するには、**neighbor** ルータ コンフィギュレーション コマンドを使用します。デフォルトでは、ピア グループ メンバーは **remote-as**（設定されている場合）、**version**、**update-source**、**out-route-map**、**out-filter-list**、**out-dist-list**、**minimum-advertisement-interval**、**next-hop-self** など、ピア グループの設定オプションをすべて継承します。すべてのピア グループ メンバーは、ピア グループに対する変更を継承します。また、発信アップデートに影響しないオプションを無効にするように、メンバーを設定することもできます。

各ネイバーに設定オプションを割り当てるには、ネイバーの IP アドレスを使用し、次に示すルータ コンフィギュレーション コマンドのいずれかを指定します。ピア グループにオプションを割り当てるには、ピア グループ名を使用し、いずれかのコマンドを指定します。**neighbor shutdown** ルータ コンフィギュレーション コマンドを使用すると、すべての設定情報を削除せずに、BGP ピアまたはピア グループをディセーブルにできます。

BGP ピアを設定するには、特権 EXEC モードで次のコマンドを使用します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	router bgp <i>autonomous-system</i>	BGP ルータ コンフィギュレーション モードを開始します。
ステップ 3	neighbor <i>peer-group-name</i> peer-group	BGP ピア グループを作成します。
ステップ 4	neighbor <i>ip-address</i> peer-group <i>peer-group-name</i>	BGP ネイバーをピア グループのメンバーにします。
ステップ 5	neighbor {<i>ip-address</i> <i>peer-group-name</i>} remote-as <i>number</i>	BGP ネイバーを指定します。 remote-as <i>number</i> を使用してピア グループが設定されていない場合は、このコマンドを使用し、EBGP ネイバーを含むピア グループを作成します。指定できる範囲は 1 ～ 65535 です。
ステップ 6	neighbor {<i>ip-address</i> <i>peer-group-name</i>} description <i>text</i>	(任意) ネイバーに記述子を関連付けます。
ステップ 7	neighbor {<i>ip-address</i> <i>peer-group-name</i>} default-originate [<i>route-map map-name</i>]	(任意) BGP スピーカー（ローカル ルータ）にネイバーへのデフォルト ルート 0.0.0.0 の送信を許可して、このルートがデフォルト ルートとして使用されるようにします。
ステップ 8	neighbor {<i>ip-address</i> <i>peer-group-name</i>} send-community	(任意) この IP アドレスのネイバーに送信する COMMUNITIES アトリビュートを指定します。
ステップ 9	neighbor {<i>ip-address</i> <i>peer-group-name</i>} update-source <i>interface</i>	(任意) IBGP セッションに、TCP 接続に関するすべての操作インターフェイスの使用を許可します。
ステップ 10	neighbor {<i>ip-address</i> <i>peer-group-name</i>} ebgp-multihop	(任意) ネイバーがセグメントに直接接続されていない場合でも、BGP セッションを使用可能にします。マルチホップ ピア アドレスへの唯一のルートがデフォルト ルート (0.0.0.0) の場合、マルチホップ セッションは確立されません。
ステップ 11	neighbor {<i>ip-address</i> <i>peer-group-name</i>} local-as <i>number</i>	(任意) ローカル AS として使用する AS 番号を指定します。指定できる範囲は 1 ～ 65535 です。
ステップ 12	neighbor {<i>ip-address</i> <i>peer-group-name</i>} advertisement-interval <i>seconds</i>	(任意) BGP ルーティング アップデートを送信する最小インターバルを設定します。

■ BGP の設定

	コマンド	目的
ステップ 13	neighbor { <i>ip-address</i> <i>peer-group-name</i> } maximum-prefix <i>maximum</i> [<i>threshold</i>]	(任意) ネイバーから受信できるプレフィクス数を制御します。指定できる範囲は 1 ～ 4294967295 です。 <i>threshold</i> (任意) は、警告メッセージが生成される基準となる最大値 (パーセント) です。デフォルト値は 75% です。
ステップ 14	neighbor { <i>ip-address</i> <i>peer-group-name</i> } next-hop-self	(任意) ネイバー宛の BGP アップデートに関して、ネクストホップでの処理をディセーブルにします。
ステップ 15	neighbor { <i>ip-address</i> <i>peer-group-name</i> } password <i>string</i>	(任意) TCP 接続での MD5 認証を BGP ピアに設定します。両方の BGP ピアに同じパスワードを設定する必要があります。そうしないと、BGP ピア間に接続が作成されません。
ステップ 16	neighbor { <i>ip-address</i> <i>peer-group-name</i> } route-map <i>map-name</i> { in out }	(任意) 着信または発信ルートにルート マップを適用します。
ステップ 17	neighbor { <i>ip-address</i> <i>peer-group-name</i> } send-community	(任意) この IP アドレスのネイバーに送信する COMMUNITIES アトリビュートを指定します。
ステップ 18	neighbor { <i>ip-address</i> <i>peer-group-name</i> } timers <i>keepalive</i> <i>holdtime</i>	(任意) ネイバーまたはピア グループ用のタイマーを設定します。 <i>keepalive</i> インターバルは、キープアライブ メッセージがピアに送信される間隔です。指定できる範囲は 1 ～ 4294967295 秒です。デフォルトは 60 秒です。 <i>holdtime</i> は、キープアライブ メッセージを受信しなかった場合、ピアが非アクティブと宣言されるまでのインターバルです。指定できる範囲は 1 ～ 4294967295 秒です。デフォルトは 180 秒です。
ステップ 19	neighbor { <i>ip-address</i> <i>peer-group-name</i> } weight <i>weight</i>	(任意) ネイバーからのすべてのルートに関するウェイトを指定します。
ステップ 20	neighbor { <i>ip-address</i> <i>peer-group-name</i> } distribute-list { <i>access-list-number</i> <i>name</i> } { in out }	(任意) アクセス リストの指定に従って、ネイバーに対して送受信される BGP ルーティング アップデートをフィルタリングします。
ステップ 21	neighbor { <i>ip-address</i> <i>peer-group-name</i> } filter-list <i>access-list-number</i> { in out weight <i>weight</i> }	(任意) BGP フィルタを確立します。
ステップ 22	neighbor { <i>ip-address</i> <i>peer-group-name</i> } version <i>value</i>	(任意) ネイバーと通信するとき使用する BGP バージョンを指定します。
ステップ 23	neighbor { <i>ip-address</i> <i>peer-group-name</i> } soft-reconfiguration inbound	(任意) 受信したアップデートの保管を開始するようにソフトウェアを設定します。
ステップ 24	end	特権 EXEC モードに戻ります。
ステップ 25	show ip bgp neighbors	設定を確認します。
ステップ 26	copy running-config startup-config	(任意) コンフィギュレーション ファイルに設定を保存します。

既存の BGP ネイバーまたはネイバー ピア グループをディセーブルにするには、**neighbor shutdown** ルータ コンフィギュレーション コマンドを使用します。ディセーブル化されている既存のネイバーまたはネイバー ピア グループをイネーブルにするには、**no neighbor shutdown** ルータ コンフィギュレーション コマンドを使用します。

集約アドレスの設定

CIDR を使用すると、集約ルート（またはスーパーネット）を作成して、ルーティング テーブルのサイズを最小化できます。BGP 内に集約ルートを設定するには、集約ルートを BGP に再配信するか、または BGP ルーティング テーブル内に集約エントリを作成します。BGP テーブル内に特定のエントリがさらに 1 つまたは複数存在する場合は、BGP テーブルに集約アドレスが追加されます。

ルーティング テーブル内に集約アドレスを作成するには、特権 EXEC モードで次のコマンドを使用します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	router bgp <i>autonomous-system</i>	BGP ルータ コンフィギュレーション モードを開始します。
ステップ 3	aggregate-address <i>address mask</i>	BGP ルーティング テーブル内に集約エントリを作成します。集約ルートは AS からのルートとしてアドバタイズされます。情報が失われた可能性があることを示すため、アトミック集約アトリビュートが設定されます。
ステップ 4	aggregate-address <i>address mask as-set</i>	(任意) AS 設定パス情報を生成します。このコマンドは、この前のコマンドと同じ規則に従う集約エントリを作成します。ただし、アドバタイズされるパスは、すべてのパスに含まれる全要素で構成される AS_SET です。多くのパスを集約するときは、このキーワードを使用しないでください。このルートは絶えず取り消され、更新されます。
ステップ 5	aggregate-address <i>address-mask summary-only</i>	(任意) サマリー アドレスだけをアドバタイズします。
ステップ 6	aggregate-address <i>address mask suppress-map map-name</i>	(任意) 選択された、より具体的なルートを抑制します。
ステップ 7	aggregate-address <i>address mask advertise-map map-name</i>	(任意) ルート マップによって指定された設定に基づいて、集約を生成します。
ステップ 8	aggregate-address <i>address mask attribute-map map-name</i>	(任意) ルート マップで指定されたアトリビュートを持つ集約を生成します。
ステップ 9	end	特権 EXEC モードに戻ります。
ステップ 10	show ip bgp neighbors [<i>advertised-routes</i>]	設定を確認します。
ステップ 11	copy running-config startup-config	(任意) コンフィギュレーション ファイルに設定を保存します。

集約エントリを削除するには、**no aggregate-address *address mask*** ルータ コンフィギュレーション コマンドを使用します。オプションをデフォルト値に戻すには、キーワードを指定してコマンドを使用します。

ルーティング ドメイン連合の設定

IBGP メッシュを削減する方法の 1 つは、AS を複数のサブ AS に分割して、単一の AS として認識される単一の連合にグループ化することです。各 AS は内部で完全にメッシュ化されていて、同じ連合内の他の AS との間には数本の接続があります。異なる AS 内にあるピアでは EBGP セッションが使用されますが、ルーティング情報は IBGP ピアと同様な方法で交換されます。特に、ネクスト ホップ、MED、およびローカル初期設定情報が維持されるため、すべての AS で単一の IGP を使用できます。

BGP 連合を設定するには、AS システム グループの AS 番号として機能する連合 ID を指定する必要があります。

BGP 連合を設定するには、特権 EXEC モードで次のコマンドを使用します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	router bgp <i>autonomous-system</i>	BGP ルータ コンフィギュレーション モードを開始します。
ステップ 3	bgp confederation identifier <i>autonomous-system</i>	BGP 連合 ID を設定します。
ステップ 4	bgp confederation peers <i>autonomous-system</i> [<i>autonomous-system ...</i>]	連合に属する AS、および特殊な EBGp ピアとして処理する AS を指定します。
ステップ 5	end	特権 EXEC モードに戻ります。
ステップ 6	show ip bgp neighbor show ip bgp network	設定を確認します。
ステップ 7	copy running-config startup-config	(任意) コンフィギュレーション ファイルに設定を保存します。

BGP ルート リフレクタの設定

BGP では、すべての IBGP スピーカーを完全メッシュ構造にする必要があります。外部ネイバーからルートを受信したルータは、そのルートをすべての内部ネイバーにアドバタイズする必要があります。ルーティング情報のループを防ぐには、すべての IBGP スピーカーを接続する必要があります。内部ネイバーは、内部ネイバーから取得されたルートを他の内部ネイバーに送信しません。

ルート リフレクタを使用すると、取得されたルートをネイバーに渡す場合に他の方法が使用されるため、すべての IBGP スピーカーを完全メッシュ構造にする必要はありません。IBGP ピアをルート リフレクタに設定すると、その IBGP ピアは IBGP によって取得されたルートを一連の IBGP ネイバーに送信するようになります。ルート リフレクタの内部ピアには、クライアント ピアと非クライアント ピア (AS 内の他のすべてのルータ) の 2 つのグループがあります。ルート リフレクタは、これらの 2 つのグループ間でルートを反映させます。ルート リフレクタおよびそのクライアント ピアは、クラスタを形成します。非クライアント ピアは相互に完全メッシュ構造にする必要がありますが、クライアント ピアはその必要はありません。クラスタ内のクライアントは、そのクラスタ外の IBGP スピーカーと通信しません。

アドバタイズされたルートを受信したルート リフレクタは、ネイバーに応じて、次のいずれかのアクションを実行します。

- EBGp スピーカーからのルートをすべてのクライアントおよび非クライアント ピアにアドバタイズします。
- 非クライアント ピアからのルートをすべてのクライアントにアドバタイズします。
- クライアントからのルートをすべてのクライアントおよび非クライアント ピアにアドバタイズします。したがって、クライアントを完全メッシュ構造にする必要はありません。

通常、クライアントのクラスタにはルート リフレクタが 1 つあり、クラスタはルート リフレクタのルータ ID で識別されます。冗長性を高めて、シングル ポイントでの障害を回避するには、クラスタに複数のルート リフレクタを設定する必要があります。このように設定した場合は、ルート リフレクタが同じクラスタ内のルート リフレクタからのアップデートを認識できるように、クラスタ内のすべてのルート リフレクタに同じクラスタ ID (4 バイト) を設定する必要があります。クラスタを処理するすべてのルート リフレクタは完全メッシュ構造にし、一連の同一なクライアント ピアおよび非クライアント ピアを設定する必要があります。

ルート リフレクタおよびクライアントを設定するには、特権 EXEC モードで次のコマンドを使用します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	router bgp <i>autonomous-system</i>	BGP ルータ コンフィギュレーション モードを開始します。
ステップ 3	neighbor <i>ip-address</i> <i>peer-group-name</i> route-reflector-client	ローカル ルータを BGP ルート リフレクタに、指定されたネイバーをクライアントに設定します。
ステップ 4	bgp cluster-id <i>cluster-id</i>	(任意) クラスタに複数のルート リフレクタが存在する場合、クラスタ ID を設定します。
ステップ 5	no bgp client-to-client reflection	(任意) クライアント間のルート反映をディセーブルにします。デフォルトでは、ルート リフレクタ クライアントからのルートは、他のクライアントに反映されます。ただし、クライアントが完全メッシュ構造の場合、ルート リフレクタはルートをクライアントに反映させる必要がありません。
ステップ 6	end	特権 EXEC モードに戻ります。
ステップ 7	show ip bgp	設定を確認します。送信元の ID およびクラスタリスト アトリビュートを表示します。
ステップ 8	copy running-config startup-config	(任意) コンフィギュレーション ファイルに設定を保存します。

ルート ダンピング化の設定

ルート フラップ ダンピング化は、インターネットワーク内でフラッピング ルートの伝播を最小化するための BGP 機能です。ルートがフラッピングと見なされるのは、ルートが使用可能、使用不可能、使用可能、使用不可能のように、状態が継続的に変化する場合があります。ルート ダンピング化がイネーブルの場合は、フラッピングしているルートに *penalty* 値が割り当てられます。ルートの累積ペナルティが設定された制限値に到達すると、ルートが稼動している場合であっても、BGP はルートのアドバタイズメントを抑制します。再使用限度は、ペナルティと比較される設定可能な値です。ペナルティが再使用限度より小さくなると、起動中の抑制されたルートのアドバタイズメントが再開されます。

IBGP によって取得されたルートには、ダンピング化が適用されません。このポリシーにより、IBGP ピアのペナルティが AS 外部のルートよりも大きくなることはありません。

BGP ルート ダンピング化を設定するには、特権 EXEC モードで次のコマンドを使用します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	router bgp <i>autonomous-system</i>	BGP ルータ コンフィギュレーション モードを開始します。
ステップ 3	bgp dampening	BGP ルート ダンピング化をイネーブルにします。
ステップ 4	bgp dampening <i>half-life</i> <i>reuse suppress</i> <i>max-suppress</i> [<i>route-map map</i>]	(任意) ルート ダンピング化係数のデフォルト値を変更します。
ステップ 5	end	特権 EXEC モードに戻ります。
ステップ 6	show ip bgp flap-statistics [{<i>regexp regexp</i>} {<i>filter-list list</i>} {<i>address mask</i> [<i>longer-prefix</i>]}]	(任意) フラッピングしているすべてのパスのフラップをモニタします。ルートの抑制が終了し、安定状態になると、統計情報が削除されます。
ステップ 7	show ip bgp dampened-paths	(任意) 抑制されるまでの時間を含めて、ダンピングされたルートを表示します。

	コマンド	目的
ステップ 8	<code>clear ip bgp flap-statistics [{regex regexp} {filter-list list} {address mask [longer-prefix]}</code>	(任意) BGP フラップ統計情報を消去して、ルートがダンピング化される可能性を小さくします。
ステップ 9	<code>clear ip bgp dampening</code>	(任意) ルート ダンピング情報を消去して、ルートの抑制を解除します。
ステップ 10	<code>copy running-config startup-config</code>	(任意) コンフィギュレーション ファイルに設定を保存します。

フラップ ダンピング化をディセーブルにするには、キーワードを指定しないで **no bgp dampening** ルータ コンフィギュレーション コマンドを使用します。ダンピング係数をデフォルト値に戻すには、値を指定して **no bgp dampening** ルータ コンフィギュレーション コマンドを使用します。

BGP のモニタおよびメンテナンス

特定のキャッシュ、テーブル、またはデータベースのすべての内容を削除できます。この作業は、特定の構造の内容が無効になる場合、または無効である疑いがある場合に必要となります。

BGP ルーティング テーブル、キャッシュ、データベースの内容など、特定の統計情報を表示できます。さらに、リソースの利用率を取得したり、ネットワーク問題を解決するための情報を使用することもできます。さらに、ノードの到達可能性に関する情報を表示し、デバイスのパケットが経由するネットワーク内のルーティング パスを検出することもできます。

表 38-8 に、BGP を消去および表示するために使用する特権 EXEC コマンドを示しています。画面に表示されるフィールドについては、『Cisco IOS IP Command Reference, Volume 2 of 3: Routing Protocols, Release 12.2』を参照してください。このマニュアルは、Cisco.com の [Documentation] > [Cisco IOS Software] > [12.2 Mainline] > [Command References] からアクセス可能です。

表 38-11 IP BGP の clear および show コマンド

コマンド	目的
<code>clear ip bgp address</code>	特定の BGP 接続をリセットします。
<code>clear ip bgp *</code>	すべての BGP 接続をリセットします。
<code>clear ip bgp peer-group tag</code>	BGP ピア グループのすべてのメンバーを削除します。
<code>show ip bgp prefix</code>	プレフィクスがアドバタイズされるピア グループ、またはピア グループに含まれないピアを表示します。ネクスト ホップやローカル プレフィクスなどのプレフィクス アトリビュートも表示されます。
<code>show ip bgp cidr-only</code>	サブネットおよびスーパーネット ネットワーク マスクを含むすべての BGP ルートを表示します。
<code>show ip bgp community [community-number] [exact]</code>	指定されたコミュニティに属するルートを表示します。
<code>show ip bgp community-list community-list-number [exact-match]</code>	コミュニティ リストで許可されたルートを表示します。
<code>show ip bgp filter-list access-list-number</code>	指定された AS パス アクセス リストによって照合されたルートを表示します。
<code>show ip bgp inconsistent-as</code>	送信元の AS と矛盾するルートを表示します。
<code>show ip bgp regexp regular-expression</code>	コマンドラインに入力された特定の正規表現と一致する AS パスを持つルートを表示します。
<code>show ip bgp</code>	BGP ルーティング テーブルの内容を表示します。

表 38-11 IP BGP の clear および show コマンド (続き)

コマンド	目的
<code>show ip bgp neighbors [address]</code>	各ネイバーとの BGP 接続および TCP 接続に関する詳細情報を表示します。
<code>show ip bgp neighbors [address] [advertised-routes dampened-routes flap-statistics paths regular-expression received-routes routes]</code>	特定の BGP ネイバーから取得されたルートを表示します。
<code>show ip bgp paths</code>	データベース内のすべての BGP パスを表示します。
<code>show ip bgp peer-group [tag] [summary]</code>	BGP ピア グループに関する情報を表示します。
<code>show ip bgp summary</code>	すべての BGP 接続のステータスを表示します。

また、**bgp log-neighbor changes** ルータ コンフィギュレーション コマンドを使用し、BGP ネイバーをリセット、起動、またはダウンさせるときに生成されるメッセージのログギングをイネーブルにすることもできます。

ISO CLNS ルーティングの設定

International Organization for Standardization (ISO; 国際標準化機構) の Connectionless Network Service (CLNS; コネクションレス型ネットワーク サービス) プロトコルは、Open System Interconnection (OSI; オープン システム インターコネクション) モデルのネットワーク層の標準です。ISO ネットワーク アーキテクチャでのアドレスは、Network Service Access Point (NSAP; ネットワーク サービス アクセス ポイント) および Network Entity Title (NET) と呼ばれます。OSI ネットワーク内の各ノードには、1 つ以上の NET が設定されています。さらに、各ノードには多数の NSAP アドレスがあります。

スイッチで **clns routing** グローバル コンフィギュレーション コマンドを使用してコネクションレス ルーティングをイネーブルにすると、スイッチは転送決定だけを行い、ルーティング関連の機能は実行しません。ダイナミック ルーティングの場合は、ルーティング プロトコルもイネーブルにする必要があります。スイッチは、ISO CLNS ネットワーク用の OSI ルーティング プロトコルに基づく Intermediate System-to-Intermediate System (IS-IS) ダイナミック ルーティング プロトコルをサポートします。

ダイナミックにルーティングを行うときは、IS-IS を使用します。このルーティング プロトコルは、エリアの概念をサポートします。エリア内では、すべてのルータはすべてのシステム ID への到達方法を認識しています。エリア間では、ルータは適切なエリアへの到達方法を認識しています。IS-IS は、ステーション ルーティング (エリア内) とエリアルーティング (エリア間) という 2 つのレベルのルーティングをサポートします。

ISO IGRP および IS-IS NSAP のアドレス指定方式の重要な違いは、エリア アドレスの定義に関する部分です。どちらもレベル 1 のルーティング (エリア内のルーティング) にはシステム ID を使用します。一方、エリア ルーティングでのアドレスの指定方法は異なります。ISO IGRP NSAP のアドレスには、ドメイン、エリア、およびシステム ID という 3 つの異なるフィールドがルーティング用に含まれます。IS-IS のアドレスには、単一の連続したエリアフィールド (ドメイン フィールドとエリア フィールドを含みます) およびシステム ID という 2 つのフィールドが含まれます。



(注)

ISO CLNS の詳細については、『*Cisco IOS Apollo Domain, Banyan VINES, DECnet, ISO CLNS and XNS Configuration Guide, Release 12.2*』を参照してください。この章で使用するコマンドの構文および使用方法の詳細については、『*Cisco IOS Apollo Domain, Banyan VINES, DECnet, ISO CLNS and XNS Command Reference, Release 12.2*』を参照するか、IOS コマンド リファレンス マスター インデックスを使用するか、オンラインで検索してください。

IS-IS ダイナミック ルーティングの設定

IS-IS は、ISO のダイナミック ルーティング プロトコルです (ISO 105890 で説明されています)。他のルーティング プロトコルとは異なり、IS-IS をイネーブルにするには、IS-IS ルーティング プロセスを作成し、それをネットワークではなく特定のインターフェイスに割り当てる必要があります。マルチエリアの IS-IS 設定構文を使用して、レイヤ 3 のスイッチまたはルータごとに複数の IS-IS ルーティング プロセスを指定できます。その場合、IS-IS ルーティング プロセスのインスタンスごとにパラメータを設定します。

小規模の IS-IS ネットワークは、ネットワーク内のすべてのルータを含む単一のエリアとして構築します。ネットワークの規模が拡大したら、通常、すべてのエリアのすべてのレベル 2 ルータを接続して構成されるバックボーン エリアとしてネットワークを再編成し、これをローカル エリアに接続します。ローカル エリア内では、ルータはすべてのシステム ID への到達方法を認識しています。エリア間では、ルータはバックボーンへの到達方法を認識し、バックボーン ルータは他のエリアへの到達方法を認識しています。

ルータは、ローカル エリア内のルーティング (ステーション ルーティング) を実行するために、レベル 1 の隣接関係を確立します。ルータは、レベル 1 エリア間のルーティング (エリア ルーティング) を実行するために、レベル 2 の隣接関係を確立します。

1 台の Cisco ルータで、最大 29 エリアのルーティングに参加し、バックボーン内のレベル 2 ルーティングを実行できます。一般に、各ルーティング プロセスはエリアに対応します。デフォルトでは、最初に設定されるルーティング プロセスのインスタンスが、レベル 1 とレベル 2 の両方のルーティングを実行します。追加のルータ インスタンスを設定でき、追加インスタンスは自動的にレベル 1 エリアとして扱われます。IS-IS ルーティング プロセスのインスタンスごとに個別にパラメータを設定する必要があります。

IS-IS マルチエリア ルーティングの場合、レベル 2 ルーティングを実行するように設定できるプロセスは 1 つだけですが、各シスコ ユニットは最大で 29 のレベル 1 エリアを定義できます。いずれかのプロセスでレベル 2 ルーティングを設定した場合、すべての追加プロセスは自動的にレベル 1 として設定されます。レベル 2 ルーティングを設定したプロセスは、同時にレベル 1 ルーティングを実行するように設定できます。レベル 2 ルーティングがルータ インスタンスに必要な場合は、**is-type** グローバル コンフィギュレーション コマンドを使用してレベル 2 機能を削除します。別のルータ インスタンスをレベル 2 ルータとして設定する場合にも、**is-type** コマンドを使用します。



(注)

IS-IS の詳細については、『*Cisco IOS IP Configuration Guide, Release 12.2*』の「IP Routing Protocols」を参照してください。ここで使用するコマンドの構文および使用方法の詳細については、『*Cisco IOS IP Command Reference, Release 12.2*』を参照してください。

ここでは、IS-IS ルーティングの設定方法について簡単に説明します。次の情報が含まれます。

- 「IS-IS のデフォルト設定」(P.38-71)
- 「IS-IS ルーティングのイネーブル化」(P.38-72)
- 「IS-IS グローバル パラメータの設定」(P.38-74)
- 「IS-IS インターフェイス パラメータの設定」(P.38-76)

IS-IS のデフォルト設定

表 38-12 に、IS-IS のデフォルト設定を示します。

表 38-12 IS-IS のデフォルト設定

機能	デフォルト設定
Link-State PDU (LSP; リンクステート PDU) エラーの無視	イネーブル。
IS-IS タイプ	従来型の IS-IS : ルータはレベル 1 (ステーション) およびレベル 2 (エリア) の両方のルータとして動作します。 マルチエリア IS-IS : IS-IS ルーティング プロセスの最初のインスタンスは、レベル 1-2 ルータです。それ以外のインスタンスはレベル 1 ルータです。
デフォルト情報送信元	ディセーブル。
IS-IS 隣接関係ステート変更ログ	ディセーブル。
LSP 生成スロットリング タイマー	連続する 2 つの発生間の最大間隔 : 5 秒。 初期 LSP 生成遅延 : 50 ミリ秒。 1 番めと 2 番めの LSP 生成の間のホールド タイム : 5000 ミリ秒。
LSP 最大存続時間 (リフレッシュなし)	LSP パケットが削除されるまでに 1200 秒 (20 分)。
LSP リフレッシュ インターバル	900 秒 (15 分) ごとに LSP リフレッシュを送信。
最大 LSP パケット サイズ	1497 バイト。
NSF 認識 ¹ (Cisco IOS Release 12.2(25)SEG 以降)	イネーブル レイヤ 3 スイッチでは、ハードウェアやソフトウェアの変更中に、ネイバー NSF 対応ルータからのパケットを転送し続けることができます。
Partial Route Computation (PRC) スロットリング タイマー	最大 PRC 待機インターバル : 5 秒。 トポロジ変更後の最初の PRC 計算遅延 : 2000 ミリ秒。 1 番めと 2 番めの PRC 生成の間のホールド タイム : 5000 ミリ秒。
パーティション回避	ディセーブル。
パスワード	エリアまたはドメインのパスワードは定義されておらず、認証はディセーブルです。
set-overload-bit	ディセーブル。イネーブルにして、引数を入力しない場合は、過負荷ビットがすぐに設定されて、 no set-overload-bit コマンドを入力するまで設定状態を維持します。
Shortest Path First (SPF) スロットリング タイマー	連続する SPF 間の最大インターバル : 10 秒。 トポロジ変更後の最初の SPF 計算 : 5500 ミリ秒。 1 番めと 2 番めの SPF 生成の間のホールド タイム : 5500 ミリ秒。
サマリー アドレス	ディセーブル。

1. NSF = Nonstop Forwarding

NSF 認識

IPv4 に対して統合された IS-IS NSF 認識機能がサポートされています。この機能を使用することにより、NSF 認識である Customer Premises Equipment (CPE; 宅内装置) のルータは、NSF 対応のルータによるパケットの NSF を補助できます。ローカル ルータは NSF を実行しない場合もありますが、NSF を認識することで、ネイバー NSF 対応ルータのルーティング データベースおよびリンクステート データベースの整合性と正確さを、スイッチオーバー プロセスの間に維持できます。

この機能は自動的にイネーブルになり、設定を行う必要はありません。この機能の詳細については、次の URL にある『*Integrated IS-IS Nonstop Forwarding (NSF) Awareness Feature Guide*』を参照してください。
http://www.cisco.com/en/US/products/sw/iosswrel/ps1839/products_white_paper09186a00801541c7.shtml

IS-IS ルーティングのイネーブル化

IS-IS をイネーブルにするには、ルーティング プロセスごとに名前と NET を指定します。その後、インターフェイスで IS-IS ルーティングをイネーブルにし、ルーティング プロセスの各インスタンスに対してエリアを指定します。

IS-IS をイネーブルにして、IS-IS ルーティング プロセスの各インスタンスにエリアを指定するには、特権 EXEC モードで次の手順を実行します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	clns routing	スイッチ上で ISO コネクションレス ルーティングをイネーブルに設定します。
ステップ 3	router isis [area tag]	指定したルーティング プロセスの IS-IS ルーティングをイネーブルにし、IS-IS ルーティングのコンフィギュレーション モードを開始します。 (任意) <i>area tag</i> 引数を使用して、IS-IS ルータの割り当て先のエリアを指定します。複数の IS-IS エリアを設定している場合は、値を入力する必要があります。 最初に設定した IS-IS インスタンスは、デフォルトでレベル 1-2 になります。以降のインスタンスは自動的にレベル 1 になります。 is-type グローバル コンフィギュレーション コマンドを使用して、ルーティングのレベルを変更できます。
ステップ 4	net network-entity-title	ルーティング プロセスに対して NET を設定します。マルチエリア IS-IS を設定している場合は、ルーティング プロセスごとに NET を指定します。NET およびアドレスの名前を指定できます。
ステップ 5	is-type {level-1 level-1-2 level-2-only}	(任意) ルータは、レベル 1 (ステーション) ルータとして、マルチエリア ルーティングのレベル 2 (エリア) ルータとして、またはその両方 (デフォルト) として動作するように設定できます。 - レベル 1: ステーション ルータとしてのみ動作します。 - レベル 1-2: ステーション ルータおよびエリア ルータの両方として動作します。 - レベル 2: エリア ルータとしてのみ動作します。
ステップ 6	exit	グローバル コンフィギュレーション モードに戻ります。
ステップ 7	interface interface-id	IS-IS をルーティングするインターフェイスを指定し、インターフェイス コンフィギュレーション モードを開始します。インターフェイスがレイヤ 3 インターフェイスとしてまだ設定されていない場合は、 no switchport コマンドを入力してレイヤ 3 モードにします。
ステップ 8	ip router isis [area tag]	インターフェイスで ISO CLNS に対する IS-IS ルーティング プロセスを設定し、エリア デジグネータをルーティング プロセスに結合します。
ステップ 9	clns router isis [area tag]	インターフェイスで ISO CLNS をイネーブルにします。

	コマンド	目的
ステップ 10	ip address <i>ip-address-mask</i>	インターフェイスの IP アドレスを定義します。いずれか 1 つのインターフェイスでも IS-IS ルーティング用に設定する場合は、IS-IS をイネーブルにするエリア内のすべてのインターフェイスで IP アドレスが必要です。
ステップ 11	end	特権 EXEC モードに戻ります。
ステップ 12	show isis [<i>area tag</i>] database detail	設定を確認します。
ステップ 13	copy running-config startup-config	(任意) コンフィギュレーション ファイルに設定を保存します。

IS-IS ルーティングをディセーブルにするには、**no router isis *area-tag*** ルータ コンフィギュレーション コマンドを使用します。

次に、IP ルーティング プロトコルとして従来型の IS-IS を実行するように 3 つのルータを設定する例を示します。従来型の IS-IS では、すべてのルータが（デフォルトで）レベル 1 およびレベル 2 のルータとして動作します。

ルータ A

```
Switch(config)# clns routing
Switch(config)# router isis
Switch(config-router)# net 49.0001.0000.0000.000a.00
Switch(config-router)# exit
Switch(config)# interface gigabitethernet1/0/1
Switch(config-if)# ip router isis
Switch(config-if)# clns router isis
Switch(config)# interface gigabitethernet1/0/2
Switch(config-if)# ip router isis
Switch(config-if)# clns router isis
Switch(config-router)# exit
```

ルータ B

```
Switch(config)# clns routing
Switch(config)# router isis
Switch(config-router)# net 49.0001.0000.0000.000b.00
Switch(config-router)# exit
Switch(config)# interface gigabitethernet1/0/1
Switch(config-if)# ip router isis
Switch(config-if)# clns router isis
Switch(config)# interface gigabitethernet1/0/2
Switch(config-if)# ip router isis
Switch(config-if)# clns router isis
Switch(config-router)# exit
```

ルータ C

```
Switch(config)# clns routing
Switch(config)# router isis
Switch(config-router)# net 49.0001.0000.0000.000c.00
Switch(config-router)# exit
Switch(config)# interface gigabitethernet1/0/1
Switch(config-if)# ip router isis
Switch(config-if)# clns router isis
Switch(config)# interface gigabitethernet1/0/2
Switch(config-if)# ip router isis
Switch(config-if)# clns router isis
Switch(config-router)# exit
```

IS-IS グローバル パラメータの設定

必要に応じて設定できる IS-IS グローバル パラメータがいくつかあります。

- ルート マップで制御されるデフォルト ルートを設定することで、IS-IS ルーティング ドメインにデフォルト ルートを強制的に適用できます。ルート マップで設定可能な他のフィルタリング オプションを指定することもできます。
- 受信した IS-IS LSP に内部チェックサム エラーがある場合は無視するように、または壊れている LSP を削除するように、ルータを設定できます。このようにすると、LSP の発信側は LSP を再生成します。
- エリアおよびドメインにパスワードを割り当てることができます。
- ルーティング テーブルでサマリー アドレスによって表される集約アドレスを作成できます（ルート サマライズ）。他のルーティング プロトコルから学習されたルートもサマライズできます。サマリーのアドバタイズに使用されるメトリックは、特定のルート全体の中で最小のメトリックです。
- 過負荷ビットを設定できます。
- LSP リフレッシュ インターバル、および LSP がリフレッシュされないでルータのデータベースに残っていることのできる最大時間を設定できます。
- LSP 生成のスロットリング タイマー、SPF 計算、および PRC を設定できます。
- IS-IS の隣接関係のステートが（アップまたはダウンに）変化したときにログ メッセージを生成するように、スイッチを設定できます。
- ネットワーク内のリンクの Maximum Transmission Unit (MTU; 最大伝送ユニット) のサイズが 1500 バイト未満の場合、そのような状況でもルーティングが発生するように LSP MTU を小さくすることができます。
- partition avoidance ルータ コンフィギュレーション コマンドは、レベル 1-2 境界ルータ、隣接するレベル 1 ルータ、およびエンド ホストの間の完全な接続が失われたときに、エリアが分割されるのを防ぎます。

IS-IS パラメータを設定するには、特権 EXEC モードで次の手順を実行します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	clns routing	スイッチ上で ISO コネクションレス ルーティングをイネーブルに設定します。
ステップ 3	router isis	IS-IS ルーティング プロトコルを指定し、ルータ コンフィギュレーション モードを開始します。
ステップ 4	default-information originate [<i>route-map map-name</i>]	(任意) 強制的に IS-IS ルーティング ドメインにデフォルト ルートを適用します。 route-map map-name を入力すると、ルート マップが満たされている場合、ルーティング プロセスはデフォルト ルートを生成します。
ステップ 5	ignore-lsp-errors	(任意) 内部チェックサム エラーのある LSP を、削除するのではなく無視するようにルータを設定します。このコマンドは、デフォルトでイネーブルになります（壊れている LSP はドロップされます）。壊れている LSP を削除するには、 no ignore-lsp-errors ルータ コンフィギュレーション コマンドを入力します。
ステップ 6	area-password password	(任意) エリア認証パスワードを設定します。パスワードはレベル 1（ステーション ルータ レベル）の LSP に挿入されます。
ステップ 7	domain-password password	(任意) ルーティング ドメイン認証パスワードを設定します。パスワードはレベル 2（エリア ルータ レベル）の LSP に挿入されます。

	コマンド	目的
ステップ 8	summary-address <i>address mask</i> [<i>level-1</i> <i>level-1-2</i> <i>level-2</i>]	(任意) 特定のレベルのアドレスのサマリーを作成します。
ステップ 9	set-overload-bit [on-startup { <i>seconds</i> wait-for-bgp }]	(任意) 過負荷ビット (hippity ビット) を設定し、問題のあるルータを他のルータが SPF 計算で無視できるようにします。 • (任意) on-startup : 起動時にのみ過負荷ビットを設定します。 on-startup を指定しないと、過負荷ビットがすぐに設定されて、no set-overload-bit コマンドを入力するまで設定状態を維持します。 on-startup を指定する場合は、秒数または wait-for-bgp を入力する必要があります。 • <i>seconds</i> : on-startup キーワードを設定すると、過負荷ビットはシステムの起動時に設定されて、この秒数だけ設定状態を維持します。指定できる範囲は 5 ~ 86400 秒です。 • wait-for-bgp : on-startup キーワードを設定すると、過負荷ビットはシステムの起動時に設定されて、BGP がコンバージするまで設定状態を維持します。BGP がコンバージする信号 IS-IS ではない場合は、IS-IS は 10 分後に過負荷ビットをオフにします。
ステップ 10	lsp-refresh-interval <i>seconds</i>	(任意) LSP リフレッシュ インターバルを秒単位で設定します。指定できる範囲は 1 ~ 65535 秒です。デフォルトでは、900 秒 (15 分) ごとに LSP リフレッシュが送信されます。
ステップ 11	max-lsp-lifetime <i>seconds</i>	(任意) LSP パケットをリフレッシュしないでルータ データベースにとどめておく最大時間を設定します。指定できる範囲は 1 ~ 65535 秒です。デフォルトは 1200 秒 (20 分) です。指定したタイム インターバルの後、LSP パケットは削除されます。
ステップ 12	lsp-gen-interval [<i>level-1</i> <i>level-2</i>] <i>lsp-max-wait</i> [<i>lsp-initial-wait</i> <i>lsp-second-wait</i>]	(任意) IS-IS LSP 生成スロットリング タイマーを設定します。 • <i>lsp-max-wait</i> : 連続して生成される 2 つの LSP 間の最大インターバル (秒) です。指定できる範囲は 1 ~ 120 です。デフォルトは 5 です。 • <i>lsp-initial-wait</i> : 初期 LSP 生成の遅延 (ミリ秒) です。指定できる範囲は 1 ~ 10000 です。デフォルトは 50 です。 • <i>lsp-second-wait</i> : 1 番めと 2 番めの LSP 生成の間のホールド タイム (ミリ秒) です。指定できる範囲は 1 ~ 10000 です。デフォルトは 5000 です。
ステップ 13	spf-interval [<i>level-1</i> <i>level-2</i>] <i>spf-max-wait</i> [<i>spf-initial-wait</i> <i>spf-second-wait</i>]	(任意) IS-IS の SPF スロットリング タイマーを設定します。 • <i>spf-max-wait</i> : 連続する SFP の間の最大インターバル (秒) です。指定できる範囲は 1 ~ 120 です。デフォルトは 10 です。 • <i>spf-initial-wait</i> : トポロジ変更後の最初の SFP 計算 (ミリ秒) です。指定できる範囲は 1 ~ 10000 です。デフォルトは 5500 です。 • <i>spf-second-wait</i> : 1 番めと 2 番めの SFP 計算の間のホールド タイム (ミリ秒) です。指定できる範囲は 1 ~ 10000 です。デフォルトは 5500 です。

	コマンド	目的
ステップ 14	prc-interval <i>prc-max-wait</i> [<i>prc-initial-wait prc-second-wait</i>]	(任意) IS-IS の PRC スロットリング タイマーを設定します。 <i>prc-max-wait</i> : 連続する 2 つの PRC 計算の間の最大インターバル (秒) です。指定できる範囲は 1 ~ 120 です。デフォルトは 5 です。 <i>prc-initial-wait</i> : トポロジ変更後の最初の PRC 計算の遅延 (ミリ秒) です。指定できる範囲は 1 ~ 10,000 です。デフォルトは 2000 です。 <i>prc-second-wait</i> : 1 番めと 2 番めの PRC 計算の間のホールド タイム (ミリ秒) です。指定できる範囲は 1 ~ 10,000 です。デフォルトは 5000 です。
ステップ 15	log-adjacency-changes [all]	(任意) IS-IS 隣接関係ステートの変化をログに記録するようにルータを設定します。 all を入力すると、End System-to-Intermediate System PDU や Link State Packet (LSP; リンク ステート パケット) など、Intermediate System-to-Intermediate System Hellos に関係のないイベントによって生成されるすべての変更が含まれます。
ステップ 16	p-mtu size	(任意) 最大 LSP パケット サイズをバイト単位で指定します。指定できる範囲は 128 ~ 4352 です。デフォルトは 1497 バイトです。 (注) ネットワーク内のいずれかのリンクで小さい MTU サイズが設定されている場合は、ネットワーク内のすべてのルータで LSP MTU サイズを変更する必要があります。
ステップ 17	partition avoidance	(任意) 境界ルータ、隣接するすべてのレベル 1 ルータ、およびエンド ホストの間の完全な接続が失われたとき、IS-IS レベル 1-2 境界ルータはレベル 2 バックボーンへのレベル 1 エリア プレフィックスのアドバタイズを停止します。
ステップ 18	end	特権 EXEC モードに戻ります。
ステップ 19	show clns	設定を確認します。
ステップ 20	copy running-config startup-config	(任意) コンフィギュレーション ファイルに設定を保存します。

デフォルト ルートの生成をディセーブルにするには、**no default-information originate** ルータ コンフィギュレーション コマンドを使用します。パスワードをディセーブルにするには、**no area-password** または **no domain-password** ルータ コンフィギュレーション コマンドを使用します。LSP MTU の設定をディセーブルにするには、**no lsp mtu** ルータ コンフィギュレーション コマンドを使用します。サマリー アドレス指定、LSP リフレッシュ インターバル、LSP 存続時間、LSP タイマー、SFP タイマー、PRC タイマーをデフォルトの設定に戻すには、コマンドの **no** 形式を使用します。出力フォーマットをディセーブルにするには、**no partition avoidance** ルータ コンフィギュレーション コマンドを使用します。

IS-IS インターフェイス パラメータの設定

必要に応じて、接続されている他のルータとは別個に、特定のインターフェイス固有の IS-IS パラメータを設定できます。ただし、係数やタイム インターバルなどの一部の値をデフォルトから変更する場合は、複数のルータおよびインターフェイスでも変更しないと意味がありません。ほとんどのインターフェイス パラメータは、レベル 1、レベル 2、または両方に対して設定できます。

設定できるインターフェイス レベル パラメータの一部を次に示します。

- インターフェイスのデフォルト メトリック。IS-IS メトリックの値として使用され、Quality of Service (QoS) ルーティングが実行されていないときに割り当てられます。

- hello インターバル (インターフェイスで送信される hello パケット間の時間の長さ)、または IS-IS hello パケットで送信されるホールドタイムを決定するためにインターフェイスで使用されるデフォルトの hello パケット係数。ホールドタイムにより、ネイバーがネイバー ダウンを宣言する前に別の hello パケットを待機する時間の長さが決定されます。これにより、障害リンクまたはネイバーが検出されてルートを再計算できるまでの時間が決まります。hello パケットが頻繁に失われて、IS-IS 隣接関係の障害が必要以上に発生する環境では、hello 係数を変更してください。hello 係数を大きくし、それに対応して hello インターバルを小さくすることで、リンク障害の検出に必要な時間を延ばすことなく、hello プロトコルの信頼性を高めることができます。
- 他には、次のようなタイム インターバルがあります。
 - Complete Sequence Number PDU (CSNP) インターバル。CSNP は、データベースの同期を維持するために指定されているルータによって送信されます。
 - 再送信インターバル。ポイントツーポイント リンクに対する IS-IS LSP の再送信間の時間です。
 - IS-IS LSP 再送信スロットル インターバル。ポイントツーポイント リンクで IS-IS LSP が再送信される最大速度 (パケット間のミリ秒数) です。このインターバルは、同じ LSP の連続する再送信の間の時間である再送信インターバルとは異なります。
- 代表ルータ選出プライオリティ。マルチアクセス ネットワークで必要な隣接関係の数を減らすことができます。その結果、ルーティング プロトコル トラフィックの量およびトポロジ データベースのサイズが減ります。
- インターフェイス回線タイプ。指定したインターフェイス上のネイバーに適した隣接関係のタイプです。
- インターフェイスのパスワード認証。

IS-IS インターフェイス パラメータを設定するには、特権 EXEC モードで次の手順を実行します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	interface interface-id	設定するインターフェイスを指定し、インターフェイス コンフィギュレーション モードを開始します。インターフェイスがレイヤ 3 インターフェイスとしてまだ設定されていない場合は、 no switchport コマンドを入力してレイヤ 3 モードにします。
ステップ 3	isis metric default-metric [level-1 level-2]	(任意) 指定したインターフェイスのメトリック (コスト) を設定します。指定できる範囲は 0 ～ 63 です。デフォルト値は 10 です。レベルを入力しないと、デフォルトとしてレベル 1 およびレベル 2 の両方のルータに適用されます。
ステップ 4	isis hello-interval {seconds minimal} [level-1 level-2]	(任意) スイッチによって送信される hello パケット間の時間の長さを指定します。デフォルトでは、hello インターバルの <i>seconds</i> の 3 倍の値が、送信される hello パケットの <i>holdtime</i> としてアドバタイズされます。hello インターバルを小さくすると、トポロジ変化の検出は早くなりますが、ルーティング トラフィックが増加します。 • minimal : システムは、結果のホールド タイムが 1 秒になるように、hello 係数に基づいて hello インターバルを計算します。 • seconds : 指定できる範囲は 1 ～ 65535 です。デフォルト値は 10 秒です。

ISO CLNS ルーティングの設定

	コマンド	目的
ステップ 5	isis hello-multiplier <i>multiplier</i> [level-1 level-2]	(任意) ルータが隣接関係をダウンとして宣言するためにネイバーが失う必要のある、IS-IS hello パケットの数を指定します。指定できる範囲は 3 ～ 1000 です。デフォルト値は 3 です。hello-multiplier を小さくするほどコンバージェンスは速くなりますが、ルーティングが不安定になる可能性があります。
ステップ 6	isis csnp-interval <i>seconds</i> [level-1 level-2]	(任意) インターフェイスに対する IS-IS の CSNP インターバルを設定します。指定できる範囲は 0 ～ 65535 です。デフォルト値は 10 秒です。
ステップ 7	isis retransmit-interval <i>seconds</i>	(任意) ポイントツーポイント リンクに対する IS-IS LSP の再送信間の秒数を設定します。ネットワーク上の任意の 2 ルータ間で予想される往復遅延より大きい整数を指定する必要があります。指定できる範囲は 0 ～ 65535 です。デフォルト値は 5 秒です。
ステップ 8	isis retransmit-throttle-interval <i>milliseconds</i>	(任意) IS-IS LSP 再送信スロットル インターバルを設定します。これは、IS-IS LSP がポイントツーポイント リンクで再送信される最大速度 (パケット間のミリ秒数) です。指定できる範囲は 0 ～ 65535 です。デフォルトは、 isis lsp-interval コマンドによって決定されます。
ステップ 9	isis priority <i>value</i> [level-1 level-2]	(任意) 代表ルータ選出に使用するプライオリティを設定します。指定できる範囲は 0 ～ 127 です。デフォルト値は 64 です。
ステップ 10	isis circuit-type {level-1 level-1-2 level-2-only}	(任意) 指定したインターフェイス (インターフェイスの回線タイプを指定) のネイバーに適した隣接関係のタイプを設定します。 • level-1 : このノードとそのネイバーの両方に共通するエリア アドレスが 1 つでもある場合、レベル 1 の隣接関係が確立されます。 • level-1-2 : ネイバーもレベル 1 およびレベル 2 の両方として設定されていて、共通のエリアが 1 つでもある場合、レベル 1 および 2 の隣接関係が確立されます。共通のエリアがない場合は、レベル 2 の隣接関係が確立されます。これがデフォルトです。 • level 2 : レベル 2 の隣接関係が確立されます。ネイバー ルータがレベル 1 ルータの場合、隣接関係は確立されません。
ステップ 11	isis password <i>password</i> [level-1 level-2]	(任意) インターフェイスの認証パスワードを設定します。認証は、デフォルトではディセーブルに設定されています。レベル 1 またはレベル 2 を指定すると、それぞれ、レベル 1 またはレベル 2 のルーティングに対してだけパスワードがイネーブルになります。レベルを指定しないと、デフォルトでレベル 1 およびレベル 2 になります。
ステップ 12	end	特権 EXEC モードに戻ります。
ステップ 13	show clns interface <i>interface-id</i>	設定を確認します。
ステップ 14	copy running-config startup-config	(任意) コンフィギュレーション ファイルに設定を保存します。

デフォルトの設定に戻すには、コマンドの **no** 形式を使用します。

ISO IGRP と IS-IS のモニタおよびメンテナンス

CLNS キャッシュのすべての内容を削除したり、特定のネイバーまたはルートの情報を削除したりできます。ルーティング テーブル、キャッシュ、データベースの内容など、CLNS または IS-IS の特定の統計情報を表示できます。また、特定のインターフェイス、フィルタ、またはネイバーについての情報も表示できます。

表 38-13 に、ISO CLNS と IS-IS ルーティングを消去および表示するために使用する特権 EXEC コマンドを示します。表示フィールドの詳細については、『Cisco IOS Apollo Domain, Banyan VINES, DECnet, ISO CLNS and XNS Command Reference, Release 12.2』を参照するか、Cisco IOS コマンド リファレンス マスター インデックスを使用するか、オンラインで検索してください。

表 38-13 ISO CLNS と IS-IS の clear および show コマンド

コマンド	目的
clear clns cache	CLNS ルーティング キャッシュを消去して再初期化します。
clear clns es-neighbors	隣接関係データベースからエンドシステム (ES) のネイバー情報を削除します。
clear clns is-neighbors	隣接関係データベースから中継システム (IS) のネイバー情報を削除します。
clear clns neighbors	隣接関係データベースから CLNS のネイバー情報を削除します。
clear clns route	動的に抽出される CLNS ルーティング情報を削除します。
show clns	CLNS ネットワークに関する情報を表示します。
show clns cache	CLNS ルーティング キャッシュのエントリを表示します。
show clns es-neighbors	ES ネイバー エントリを表示します。関連付けられたエリアを含みます。
show clns filter-expr	フィルタ式を表示します。
show clns filter-set	フィルタ セットを表示します。
show clns interface [interface-id]	各インターフェイスに関する CLNS 固有の情報または ES-IS の情報を表示します。
show clns neighbor	IS-IS ネイバーに関する情報を表示します。
show clns protocol	このルータの各 IS-IS または ISO IGRP ルーティングプロセスに対するプロトコル固有の情報を一覧表示します。
show clns route	このルータが CLNS パケットのルーティング方法を知っているすべての宛先を表示します。
show clns traffic	このルータが検出した CLNS パケットに関する情報を表示します。
show ip route isis	ISIS IP ルーティング テーブルの現在のステートを表示します。
show isis database	IS-IS リンクステート データベースを表示します。
show isis routes	IS-IS レベル 1 ルーティング テーブルを表示します。
show isis spf-log	IS-IS に対する SPF 計算の履歴を表示します。
show isis topology	全エリアで接続されているすべてのルータの一覧を表示します。
show route-map	設定されたすべてのルート マップ、または指定されたルート マップだけを表示します。
trace clns destination	ネットワーク内のパケットによって、指定した宛先に対して取得されたパスを検出します。
which-route {nsap-address clns-name}	指定した CLNS の宛先が含まれるルーティング テーブルを表示します。

マルチ VRF CE の設定

Virtual Private Network (VPN; バーチャル プライベート ネットワーク) は、ISP バックボーン ネットワーク上でお客様にセキュアな帯域幅共有を提供します。VPN は、共通ルーティング テーブルを共有するサイトの集合です。カスタマー サイトは、1 つまたは複数のインターフェイスでサービス プロバイダー ネットワークに接続され、サービス プロバイダーは、VPN Routing/Forwarding (VRF) テーブルと呼ばれる VPN ルーティング テーブルと各インターフェイスを関連付けます。

Catalyst 3750 スイッチは、スイッチで IP サービス イメージが稼動中の場合に、Customer Edge (CE; カスタマー エッジ) デバイスの multiple VPN Routing/Forwarding (multi-VRF; マルチ VRF) インスタンスをサポートします (マルチ VRF CE)。サービス プロバイダーは、マルチ VRF CE により、重複する IP アドレスで複数の VPN をサポートできます。IP ベース イメージが稼動しているスイッチでこれを設定しようとする、エラー メッセージが表示されます。IP ベース イメージが稼動しているスイッチで、マルチ VRF CE と EIGRP スタブ ルーティングを同時に設定することは許可されていません。



(注)

スイッチでは、VPN のサポートのために Multiprotocol Label Switching (MPLS; マルチプロトコル ラベル スイッチング) が使用されません。MPLS VRF に関する詳細については、『Cisco IOS Switching Services Configuration Guide, Release 12.2』を参照してください。これには、Cisco.com のホームページ ([Documentation] > [Cisco IOS Software] > [12.2 Mainline] > [Command References]) からアクセス可能です。

- 「マルチ VRF CE の概要」 (P.38-80)
- 「マルチ VRF CE のデフォルト設定」 (P.38-82)
- 「マルチ VRF CE の設定時の注意事項」 (P.38-82)
- 「VRF の設定」 (P.38-83)
- 「VRF 認識サービスの設定」 (P.38-85)
- 「VPN ルーティング セッションの設定」 (P.38-89)
- 「BGP PE/CE ルーティング セッションの設定」 (P.38-89)
- 「マルチ VRF CE の設定例」 (P.38-90)
- 「マルチ VRF CE ステータスの表示」 (P.38-94)

マルチ VRF CE の概要

マルチ VRF CE は、サービス プロバイダーが複数の VPN をサポートし、VPN 間で IP アドレスを重複して使用できるようにする機能です。マルチ VRF CE は入力インターフェイスを使用して、さまざまな VPN のルートを区別し、1 つまたは複数のレイヤ 3 インターフェイスと各 VRF を関連付けて仮想パケット転送テーブルを形成します。VRF 内のインターフェイスは、イーサネット ポートのように物理的なもの、または VLAN SVI のように論理的なものにもできますが、複数の VRF に属することはできません。



(注)

マルチ VRF CE インターフェイスは、レイヤ 3 インターフェイスである必要があります。

マルチ VRF CE には、次のデバイスが含まれます。

- お客様は、CE デバイスにより、1 つまたは複数の Provider Edge (PE; プロバイダー エッジ) ルータへのデータ リンクを介してサービス プロバイダー ネットワークにアクセスできます。CE デバイスは、サイトのローカル ルートをルータにアドバタイズし、リモート VPN ルートをそこから学習します。Catalyst 3750 スイッチは、CE にすることができます。
- PE ルータは、スタティック ルーティング、または BGP、RIPv2、OSPF、EIGRP などのルーティング プロトコルを使用して、CE デバイスとルーティング情報を交換します。PE では、直接接続している VPN の VPN ルートだけを維持すればよく、すべてのサービス プロバイダー VPN ルートを維持する必要はありません。各 PE ルータは、直接接続しているサイトごとに VRF を維持します。すべてのサイトが同じ VPN に存在する場合は、PE ルータの複数のインターフェイスを 1 つの

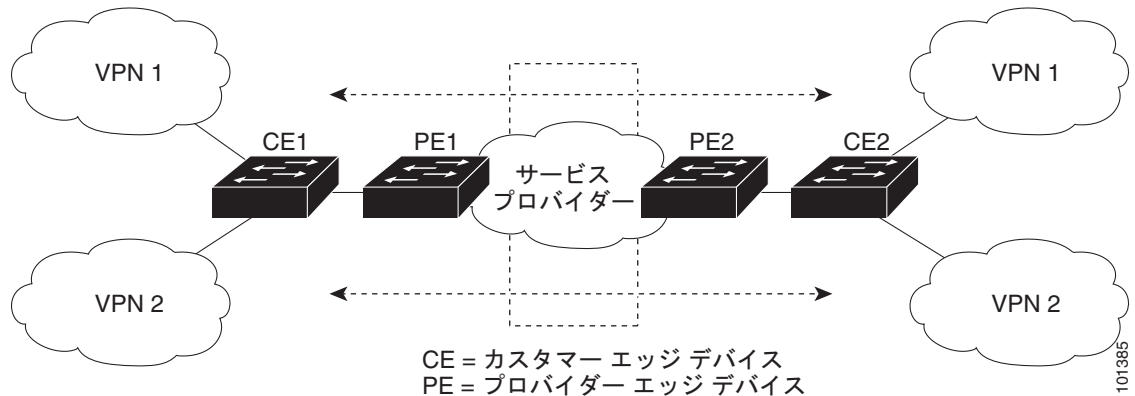
VRF に関連付けることができます。各 VPN は、指定された VRF にマッピングされます。PE ルータは、ローカル VPN ルートを CE から学習した後で、IBGP を使用して別の PE ルータと VPN ルーティング情報を交換します。

- CE デバイスに接続していないサービス プロバイダー ネットワークのルータは、プロバイダー ルータやコア ルータになります。

マルチ VRF CE では、複数のお客様が 1 つの CE を共有でき、CE と PE の間で 1 つの物理リンクだけが使用されます。共有 CE は、お客様ごとに別々の VRF テーブルを維持し、独自のルーティング テーブルに基づいて、お客様ごとにパケットをスイッチングまたはルーティングします。マルチ VRF CE は、制限付きの PE 機能を CE デバイスに拡張して、別々の VRF テーブルを維持し、VPN のプライバシーおよびセキュリティを支店に拡張します。

図 38-6 は、Catalyst 3750 スイッチを複数の仮想 CE として使用した設定を示しています。このシナリオは、中小企業など、VPN サービスの帯域幅要件の低いお客様に適しています。そのような場合、Catalyst 3750 スイッチではマルチ VRF CE のサポートが必要です。マルチ VRF CE はレイヤ 3 機能なので、VRF のそれぞれのインターフェイスはレイヤ 3 インターフェイスである必要があります。

図 38-6 複数の仮想 CE として機能するスイッチ



CE スイッチは、レイヤ 3 インターフェイスを VRF に追加するコマンドを受信すると、マルチ VRF CE 関連のデータ構造で VLAN ID と Policy Label (PL; ポリシー ラベル) の間に適切なマッピングを設定し、VLAN ID と PL を VLAN データベースに追加します。

マルチ VRF CE を設定すると、レイヤ 3 転送テーブルは、次の 2 つのセクションに概念的に分割されます。

- マルチ VRF CE ルーティング セクションには、さまざまな VPN からのルートが含まれます。
- グローバル ルーティング セクションには、インターネットなど、VPN 以外のネットワークへのルートが含まれます。

さまざまな VRF の VLAN ID はさまざまなポリシー ラベルにマッピングされ、処理中に VRF を区別するために使用されます。レイヤ 3 設定機能では、学習した新しい VPN ルートごとに、入力ポートの VLAN ID を使用してポリシー ラベルを取得し、マルチ VRF CE ルーティング セクションにポリシー ラベルおよび新しいルートを挿入します。ルーテッド ポートからパケットを受信した場合は、ポート内部 VLAN ID 番号が使用されます。SVI からパケットを受信した場合は、VLAN 番号が使用されます。

マルチ VRF CE 対応ネットワークのパケット転送処理は次のとおりです。

- スイッチは、VPN からパケットを受信すると、入力ポリシー ラベル番号に基づいてルーティング テーブルを検索します。ルートが見つかったら、スイッチはパケットを PE に転送します。
- 入力 PE は、CE からパケットを受信すると、VRF 検索を実行します。ルートが見つかったら、ルータは対応する MPLS ラベルをパケットに追加し、MPLS ネットワークに送信します。

- 出力 PE は、ネットワークからパケットを受信すると、ラベルを除去してそのラベルを使用し、正しい VPN ルーティング テーブルを識別します。次に、通常のルート検索を実行します。ルートが見つかり、パケットを正しい隣接デバイスに転送します。
- CE は、出力 PE からパケットを受信すると、入力ポリシー ラベルを使用して正しい VPN ルーティング テーブルを検索します。ルートが見つかり、パケットを VPN 内で転送します。

VRF を設定するには、VRF テーブルを作成し、VRF に関連するレイヤ 3 インターフェイスを指定します。次に、VPN、および CE と PE 間でルーティング プロトコルを設定します。プロバイダーのバックボーンで VPN ルーティング情報を配信する場合は、BGP が望ましいルーティング プロトコルです。マルチ VRF CE ネットワークには、次の 3 つの主要コンポーネントがあります。

- VPN ルート ターゲット コミュニティ : VPN コミュニティのその他すべてのメンバーのリスト。VPN コミュニティ メンバーごとに VPN ルート ターゲットを設定する必要があります。
- VPN コミュニティ PE ルータのマルチプロトコル BGP ピアリング : VPN コミュニティのすべてのメンバーに VRF 到達可能性情報を伝播します。VPN コミュニティのすべての PE ルータで BGP ピアリングを設定する必要があります。
- VPN 転送 : VPN サービス プロバイダー ネットワークを介し、全 VPN コミュニティ メンバー間で、全トラフィックを伝送します。

マルチ VRF CE のデフォルト設定

表 38-14 に、VRF のデフォルト設定を示します。

表 38-14 VRF のデフォルト設定

機能	デフォルト設定
VRF	ディセーブル。VRF は定義されていません。
マップ	インポート マップ、エクスポート マップ、ルート マップは定義されていません。
VRF 最大ルート数	ファスト イーサネット スイッチ : 8000。 ギガビット イーサネット スイッチ : 12000。
転送テーブル	インターフェイスのデフォルトは、グローバル ルーティング テーブルです。

マルチ VRF CE の設定時の注意事項



(注)

マルチ VRF CE を使用するには、IP サービス イメージをスイッチにインストールする必要があります。

ネットワークに VRF を設定する場合は、次のことに注意してください。

- マルチ VRF CE を含むスイッチは複数のお客様によって共有され、各お客様には独自のルーティング テーブルがあります。
- お客様は別々の VRF テーブルを使用するので、同じ IP アドレスを再利用できます。別々の VPN では IP アドレスの重複が許可されます。
- マルチ VRF CE では、複数のお客様が、PE と CE の間で同じ物理リンクを共有できます。複数の VLAN を持つトランク ポートでは、パケットがお客様間で分離されます。それぞれのお客様には独自の VLAN があります。

- マルチ VRF CE ではサポートされない MPLS-VRF 機能があります。ラベル交換、LDP 隣接関係、ラベル付きパケットはサポートされません。
- PE ルータの場合、マルチ VRF CE の使用と複数の CE の使用に違いはありません。図 38-6 では、複数の仮想レイヤ 3 インターフェイスがマルチ VRF CE デバイスに接続されています。
- スイッチでは、物理ポートか VLAN SVI、またはその両方の組み合わせを使用して、VRF を設定できます。SVI は、アクセス ポートまたはトランク ポートで接続できます。
- お客様は、別のお客様と重複しない限り、複数の VLAN を使用できます。お客様の VLAN は、スイッチに保存されている適切なルーティング テーブルの識別に使用される特定のルーティング テーブル ID にマッピングされます。
- Catalyst 3750 スイッチは、1 つのグローバル ネットワークおよび最大 26 の VRF をサポートします。
- CE と PE の間では、ほとんどのルーティング プロトコル (BGP、OSPF、RIP、およびスタティック ルーティング) を使用できます。ただし、次の理由から External BGP (EBGP) を使用することを推奨します。
 - BGP では、複数の CE とのやり取りに複数のアルゴリズムを必要としません。
 - BGP は、さまざまな管理者によって稼動するシステム間でルーティング情報を渡すように設計されています。
 - BGP では、ルートのアトリビュートを CE に簡単に渡すことができます。
- マルチ VRF CE は、パケットのスイッチング レートに影響しません。
- VPN マルチキャストはサポートされません。
- マルチ VRF CE 内のラインレート マルチキャスト転送をサポートしています。
- マルチキャスト VRF は、同一インターフェイス上でプライベート VLAN と共存することができません。
- 最大 1000 のマルチキャスト ルータがサポートされていて、すべての VRF で共有可能です。
- VRF を設定しない場合は、105 のポリシーを設定できます。
- VRF を 1 つでも設定する場合は、41 のポリシーを設定できます。
- 41 より多いポリシーを設定する場合は、VRF を設定できません。
- VRF とプライベート VLAN は相互に排他的です。プライベート VLAN では VRF をイネーブルにすることはできません。同じように、VLAN インターフェイスで VRF が設定されている VLAN では、プライベート VLAN をイネーブルにはできません。
- VRF と Policy-Based Routing (PBR; ポリシーベース ルーティング) は、スイッチ インターフェイス上で相互に排他的です。PBR がインターフェイスでイネーブルになっているときは、VRF をイネーブルにすることはできません。同じように、インターフェイスで VRF がイネーブルになっているときは、PBR をイネーブルにはできません。
- VRF と Web Cache Communication Protocol (WCCP) は、スイッチ インターフェイス上で相互に排他的です。インターフェイスで WCCP がイネーブルになっているときは、VRF をイネーブルにすることはできません。同じように、インターフェイスで VRF がイネーブルになっているときは、WCCP をイネーブルにはできません。

VRF の設定

1 つまたは複数の VRF を設定するには、特権 EXEC モードで次の手順を実行します。コマンドの完全な構文と使用方法については、このリリースのスイッチのコマンドリファレンス、および『Cisco IOS Switching Services Command Reference, Release 12.2』を参照してください。

■ マルチ VRF CE の設定

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	ip routing	IP ルーティングをイネーブルにします
ステップ 3	ip vrf vrf-name	VRF に名前を付けて VRF コンフィギュレーション モードを開始します。
ステップ 4	rd route-distinguisher	ルート識別子を指定し、VRF テーブルを作成します。AS 番号と任意の番号 (xxx:y) または IP アドレスと任意の番号 (A.B.C.D:y) を入力します。
ステップ 5	route-target {export import both} route-target-ext-community	指定した VRF のインポート コミュニティ、エクスポート コミュニティ、またはインポートとエクスポートのルート ターゲット コミュニティのリストを作成します。AS システム番号と任意の番号 (xxx:y) または IP アドレスと任意の番号 (A.B.C.D:y) を入力します。 <i>route-target-ext-community</i> は、ステップ 4 で入力した <i>route-distinguisher</i> と同一にする必要があります。
ステップ 6	import map route-map	(任意) ルート マップを VRF に関連付けます。
ステップ 7	interface interface-id	VRF に関連付けるレイヤ 3 インターフェイスを指定し、インターフェイス コンフィギュレーション モードを開始します。インターフェイスはルーテッド ポートまたは SVI に設定できます。
ステップ 8	ip vrf forwarding vrf-name	VRF をレイヤ 3 インターフェイスに関連付けます。
ステップ 9	end	特権 EXEC モードに戻ります。
ステップ 10	show ip vrf [brief detail interfaces] [vrf-name]	設定を確認します。設定した VRF に関する情報を表示します。
ステップ 11	copy running-config startup-config	(任意) コンフィギュレーション ファイルに設定を保存します。

VRF を削除してすべてのインターフェイスを削除するには、**no ip vrf vrf-name** グローバル コンフィギュレーション コマンドを使用します。VRF からあるインターフェイスを削除するには、**no ip vrf forwarding** インターフェイス コンフィギュレーション コマンドを使用します。

マルチキャスト VRF の設定

VRF テーブル内にマルチキャストを設定するには、特権 EXEC モードで次の手順を実行します。コマンドの構文と使用方法の詳細については、このリリースのスイッチのコマンドリファレンス、および『Cisco IOS Switching Services Command Reference, Release 12.2』を参照してください。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	ip routing	IP ルーティング モードをイネーブルにします。
ステップ 3	ip vrf vrf-name	VRF に名前を付けて VRF コンフィギュレーション モードを開始します。
ステップ 4	rd route-distinguisher	ルート識別子を指定し、VRF テーブルを作成します。AS 番号と任意の番号 (xxx:y) または IP アドレスと任意の番号 (A.B.C.D:y) を入力します。

	コマンド	目的
ステップ 5	route-target {export import both} <i>route-target-ext-community</i>	指定した VRF のインポート コミュニティ、エクスポート コミュニティ、またはインポートとエクスポートのルート ターゲット コミュニティのリストを作成します。AS システム番号と任意の番号 (xxx:y) または IP アドレスと任意の番号 (A.B.C.D:y) を入力します。 <i>route-target-ext-community</i> は、ステップ 4 で入力した <i>route-distinguisher</i> と同一にする必要があります。
ステップ 6	import map <i>route-map</i>	(任意) ルート マップを VRF に関連付けます。
ステップ 7	ip multicast-routing vrf vrf-name distributed	(任意) VRF テーブルのグローバル マルチキャスト ルーティングをイネーブルにします。
ステップ 8	interface <i>interface-id</i>	VRF に関連付けるレイヤ 3 インターフェイスを指定し、インターフェイス コンフィギュレーション モードを開始します。インターフェイスはルーテッド ポートまたは SVI に設定できます。
ステップ 9	ip vrf forwarding vrf-name	VRF をレイヤ 3 インターフェイスに関連付けます。
ステップ 10	ip address ip-address mask	レイヤ 3 インターフェイスの IP アドレスを設定します。
ステップ 11	ip pim sparse-dense mode	VRF 関連レイヤ 3 インターフェイス上で PIM をイネーブルにします。
ステップ 12	end	特権 EXEC モードに戻ります。
ステップ 13	show ip vrf [brief detail interfaces] <i>[vrf-name]</i>	設定を確認します。設定した VRF に関する情報を表示します。
ステップ 14	copy running-config startup-config	(任意) コンフィギュレーション ファイルに設定を保存します。

マルチキャスト VRF CE 内でのマルチキャストの設定に関する詳細については、『*Cisco IOS IP Configuration Guide, Release 12.4*』を参照してください。

VRF 認識サービスの設定

IP サービスはグローバル インターフェイス上に設定することが可能で、これらのサービスをグローバル ルーティング インスタンス内で実行することができます。IP サービスは、複数のルーティング インスタンスで実行されるように拡張されていて、これが VRF 認識です。システム内に設定された VRF は、VRF 認識サービス用に指定できます。

VRF 認識サービスは、プラットフォームから独立したモジュールに実装されています。VRF とは、Cisco IOS で複数のルーティング インスタンスのことです。各プラットフォームには独自のサポート VRF 数の制限があります。

VRF 認識サービスには、次の特性があります。

- ユーザは、ユーザ指定の VRF 内のホストに ping を実行することができます。
- ARP エントリは個別の VRF で学習されます。ユーザは、特定の VRF の ARP アドレス解決プロトコル) エントリを表示することができます。

これらのサービスは VRF 認識です。

- ARP
- ping
- Simple Network Management Protocol (SNMP; 簡易ネットワーク管理プロトコル)
- Hot Standby Router Protocol (HSRP; ホットスタンバイ ルータ プロトコル)
- RADIUS

- Syslog
- traceroute
- FTP と TFTP



(注) VRF 認識サービスは Unicast Reverse Path Forwarding (uRPF; ユニキャスト RPF) をサポートしません。

ARP のユーザ インターフェイス

ARP の VRF 認識サービスを設定するには、特権 EXEC モードで次の手順を実行します。コマンドの完全な構文と使用方法については、このリリースのスイッチのコマンド リファレンス、および『Cisco IOS Switching Services Command Reference, Release 12.2』を参照してください。

コマンド	目的
<code>show ip arp vrf vrf-name</code>	指定された VRF 内の ARP テーブルを表示します。

ping のユーザ インターフェイス

ping の VRF 認識サービスを設定するには、特権 EXEC モードで次の手順を実行します。コマンドの完全な構文と使用方法については、このリリースのスイッチのコマンド リファレンス、および『Cisco IOS Switching Services Command Reference, Release 12.2』を参照してください。

コマンド	目的
<code>ping vrf vrf-name ip-host</code>	指定された VRF 内の ARP テーブルを表示します。

SNMP のユーザ インターフェイス

SNMP の VRF 認識サービスを設定するには、特権 EXEC モードで次の手順を実行します。コマンドの完全な構文と使用方法については、このリリースのスイッチのコマンド リファレンス、および『Cisco IOS Switching Services Command Reference, Release 12.2』を参照してください。

	コマンド	目的
ステップ 1	<code>configure terminal</code>	グローバル コンフィギュレーション モードを開始します。
ステップ 2	<code>snmp-server trap authentication vrf</code>	VRF 上のパケットの SNMP トラップをイネーブルにします。
ステップ 3	<code>snmp-server engineID remote <host> vrf <vpn instance> <engine-id string></code>	スイッチ上のリモート SNMP エンジンの名前を指定します。
ステップ 4	<code>snmp-server host <host> vrf <vpn instance> traps <community></code>	SNMP トラップ動作の受信側を指定して、SNMP トラップの送信に使用される VRF テーブルを指定します。
ステップ 5	<code>snmp-server host <host> vrf <vpn instance> informs <community></code>	SNMP 情報動作の受信側を指定して、SNMP 情報の送信に使用される VRF テーブルを指定します。
ステップ 6	<code>snmp-server user <user> <group> remote <host> vrf <vpn instance> <security model></code>	SNMP アクセス用に、VRF 上にあるリモート ホストの SNMP グループにユーザを追加します。
ステップ 7	<code>end</code>	特権 EXEC モードに戻ります。

HSRP のユーザ インターフェイス

VRF の HSRP サポートにより、HSRP 仮想 IP アドレスが確実に正しい IP ルーティング テーブルに追加されます。

HSRP の VRF 認識サービスを設定するには、特権 EXEC モードで次の手順を実行します。コマンドの完全な構文と使用方法については、このリリースのスイッチのコマンド リファレンス、および『Cisco IOS Switching Services Command Reference, Release 12.2』を参照してください。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	interface interface-id	インターフェイス コンフィギュレーション モードを開始し、設定するレイヤ 3 インターフェイスを指定します。
ステップ 3	no switchport	物理インターフェイスの場合、レイヤ 2 コンフィギュレーション モードからインターフェイスを削除します。
ステップ 4	ip vrf forwarding <vrf-name>	インターフェイス上で VRF をイネーブルにします。
ステップ 5	ip address ip address	インターフェイスの IP アドレスを入力します。
ステップ 6	standby 1 ip ip address	HSRP をイネーブルにして、仮想 IP アドレスを設定します。
ステップ 7	end	特権 EXEC モードに戻ります。

VRF 認識 RADIUS のユーザ インターフェイス

VRF 認識 RADIUS を設定するには、まず RADIUS サーバで AAA をイネーブルにする必要があります。スイッチは、**ip vrf forwarding vrf-name** サーバ グループ コンフィギュレーション コマンドと **ip radius source-interface** グローバル コンフィギュレーション コマンドをサポートします。次の URL にある『Per VRF AAA Feature Guide』を参照してください。

http://www.cisco.com/en/US/docs/ios/12_2t/12_2t13/feature/guide/ftvrfaaa.html

Syslog のユーザ インターフェイス

Syslog の VRF 認識サービスを設定するには、特権 EXEC モードで次の手順を実行します。コマンドの完全な構文と使用方法については、このリリースのスイッチのコマンド リファレンス、および『Cisco IOS Switching Services Command Reference, Release 12.2』を参照してください。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	logging on	ストレージルータ イベント メッセージのログギングをイネーブルにしたり、一時的にディセーブルにしたりします。
ステップ 3	logging host ip address vrf vrf name	ログギング メッセージが送信される Syslog サーバのホスト アドレスを指定します。
ステップ 4	logging buffered logging buffered size debugging	内部バッファへのメッセージを記録します。
ステップ 5	logging trap debugging	Syslog サーバに送信されるログギング メッセージを制限します。
ステップ 6	logging facility facility	システム ログギング メッセージをログギング ファシリティに送信します。
ステップ 7	end	特権 EXEC モードに戻ります。

tracertoute のユーザ インターフェイス

tracertoute の VRF 認識サービスを設定するには、特権 EXEC モードで次の手順を実行します。コマンドの完全な構文と使用方法については、このリリースのスイッチのコマンドリファレンス、および『Cisco IOS Switching Services Command Reference, Release 12.2』を参照してください。

コマンド	目的
tracertoute vrf <i>vrf-name</i> <i>ipaddress</i>	VPN VRF 内の宛先アドレスを検索するために VPN VRF の名前を指定します。

FTP および TFTP のユーザ インターフェイス

FTP と TFTP が VRF 認識とするためには、いくつかの FTP/TFTP CLI を設定する必要があります。たとえば、インターフェイスに添付されている VRF テーブルを使用する場合、E1/0 であれば、CLI **ip [t]ftp source-interface E1/0** を設定して、特定のルーティング テーブルを使用するように [t]ftp に通知します。この例では、VRF テーブルが宛先 IP アドレスを検索するために使用されます。これらの変更には下位互換性があり、既存の動作には影響しません。つまり、VRF がそのインターフェイスに設定されていなくても、送信元インターフェイス CLI を使用してパケットを特定のインターフェイスに送信することができます。

FTP 接続の IP アドレスを指定するには、**ip ftp source-interface show mode** コマンドを使用します。接続が行われているインターフェイスのアドレスを使用するには、**no** 形式のコマンドを使用します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	ip ftp source-interface <i>interface-type</i> <i>interface-number</i>	FTP 接続の送信元 IP アドレスを指定します。
ステップ 3	end	特権 EXEC モードに戻ります。

TFTP 接続の送信元アドレスとしてインターフェイスの IP アドレスを指定するには、**ip tftp source-interface show mode** コマンドを使用します。デフォルトに戻るには、**no** 形式のコマンドを使用します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	ip tftp source-interface <i>interface-type</i> <i>interface-number</i>	TFTP 接続の送信元 IP アドレスを指定します。
ステップ 3	end	特権 EXEC モードに戻ります。

VPN ルーティング セッションの設定

VPN 内のルーティングは、サポートされている任意のルーティング プロトコル (RIP、OSPF、EIGRP、BGP)、またはスタティック ルーティングで設定できます。ここで説明する設定は OSPF のものですが、その他のプロトコルでも手順は同じです。



(注) VRF インスタンス内で稼動するように EIGRP ルーティング プロセスを設定するには、**autonomous-system autonomous-system-number** アドレスファミリ コンフィギュレーション モード コマンドを入力して AS 番号を設定する必要があります。

VPN 内で OSPF を設定するには、特権 EXEC モードで次の手順を実行します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	router ospf process-id vrf vrf-name	OSPF ルーティングをイネーブルにして VPN 転送テーブルを指定し、ルータ コンフィギュレーション モードを開始します。
ステップ 3	log-adjacency-changes	(任意) 隣接関係ステートの変更をログします。これがデフォルトのステートです。
ステップ 4	redistribute bgp autonomous-system-number subnets	BGP ネットワークから OSPF ネットワークに情報を再配信するようにスイッチを設定します。
ステップ 5	network network-number area area-id	OSPF が動作するネットワーク アドレスとマスク、およびそのネットワーク アドレスのエリア ID を定義します。
ステップ 6	end	特権 EXEC モードに戻ります。
ステップ 7	show ip ospf process-id	OSPF ネットワークの設定を確認します。
ステップ 8	copy running-config startup-config	(任意) コンフィギュレーション ファイルに設定を保存します。

VPN 転送テーブルと OSPF ルーティング プロセスの関連付けを解除するには、**no router ospf process-id vrf vrf-name** グローバル コンフィギュレーション コマンドを使用します。

BGP PE/CE ルーティング セッションの設定

BGP PE/CE ルーティング セッションを設定するには、特権 EXEC モードで次の手順を実行します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	router bgp autonomous-system-number	その他の BGP ルータに AS 番号を渡す BGP ルーティング プロセスを設定し、ルータ コンフィギュレーション モードを開始します。
ステップ 3	network network-number mask network-mask	ネットワークとマスクを指定し、BGP の使用を宣言します。
ステップ 4	redistribute ospf process-id match internal	OSPF 内部ルートを再配信するようにスイッチを設定します。
ステップ 5	network network-number area area-id	OSPF が動作するネットワーク アドレスとマスク、およびそのネットワーク アドレスのエリア ID を定義します。

マルチ VRF CE の設定

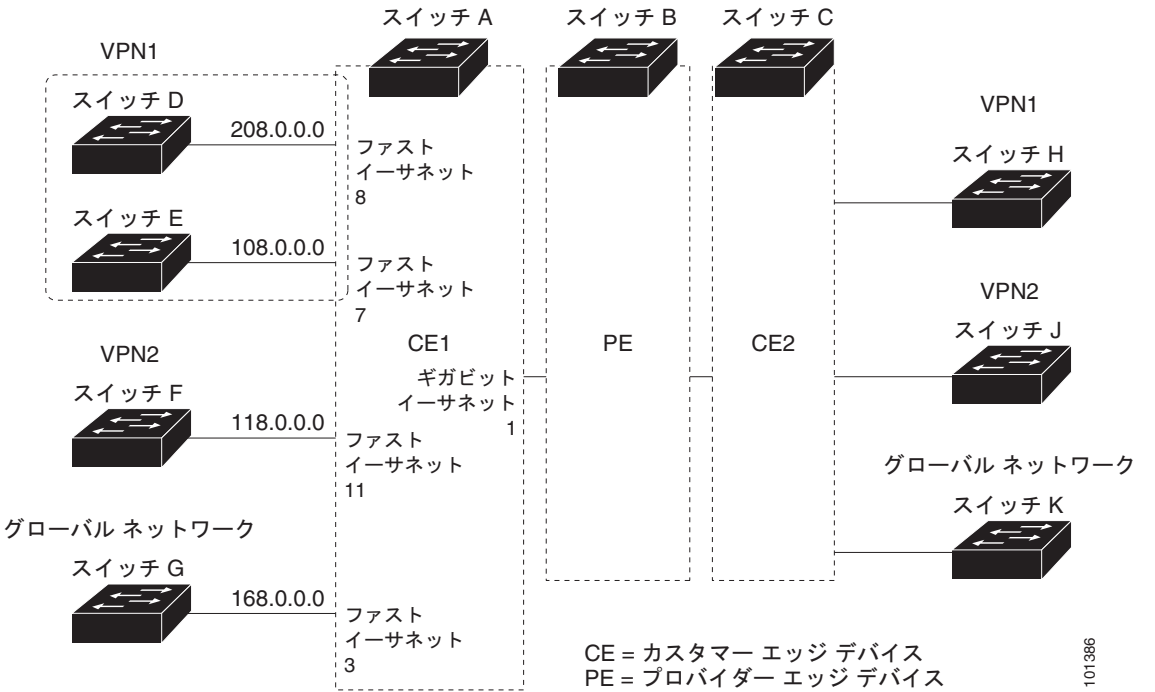
	コマンド	目的
ステップ 6	<code>address-family ipv4 vrf vrf-name</code>	PE/CE ルーティング セッションの BGP パラメータを定義し、VRF アドレスファミリ モードを開始します。
ステップ 7	<code>neighbor address remote-as as-number</code>	PE と CE ルータ間の BGP セッションを定義します。
ステップ 8	<code>neighbor address activate</code>	IPv4 アドレス ファミリのアドバタイズメントをアクティブにします。
ステップ 9	<code>end</code>	特権 EXEC モードに戻ります。
ステップ 10	<code>show ip bgp [ipv4] [neighbors]</code>	BGP 設定を確認します。
ステップ 11	<code>copy running-config startup-config</code>	(任意) コンフィギュレーション ファイルに設定を保存します。

BGP ルーティング プロセスを削除するには、`no router bgp autonomous-system-number` グローバル コンフィギュレーション コマンドを使用します。ルーティング特性を削除するには、コマンドにキーワードを指定してこのコマンドを使用します。

マルチ VRF CE の設定例

図 38-7 は、図 38-6 と同じネットワークの物理接続を簡素化した例です。VPN1、VPN2、およびグローバル ネットワークで使用されるプロトコルは OSPF です。CE/PE 接続には BGP が使用されます。図の後に続く出力は、Catalyst 3750 スイッチを CE スイッチ A として設定する例、およびカスタマー スイッチ D と F の VRF 設定を示しています。CE スイッチ C とその他のカスタマー スイッチを設定するコマンドは含まれていませんが、内容は同じです。この例には、PE ルータとして動作する Catalyst 6000 スイッチまたは Catalyst 6500 スイッチのスイッチ A へのトラフィックを設定するコマンドも含まれています。

図 38-7 マルチ VRF CE の設定例



スイッチ A の設定

スイッチ A では、ルーティングをイネーブルにして VRF を設定します。

```
Switch# configure terminal
Enter configuration commands, one per line. End with CNTL/Z.
Switch(config)# ip routing
Switch(config)# ip vrf v11
Switch(config-vrf)# rd 800:1
Switch(config-vrf)# route-target export 800:1
Switch(config-vrf)# route-target import 800:1
Switch(config-vrf)# exit
Switch(config)# ip vrf v12
Switch(config-vrf)# rd 800:2
Switch(config-vrf)# route-target export 800:2
Switch(config-vrf)# route-target import 800:2
Switch(config-vrf)# exit
```

スイッチ A のループバックおよび物理インターフェイスを設定します。ギガビット イーサネット ポート 1 は PE へのトランク接続です。ファスト イーサネット ポート 8 と 11 は VPN に接続されます。

```
Switch(config)# interface loopback1
Switch(config-if)# ip vrf forwarding v11
Switch(config-if)# ip address 8.8.1.8 255.255.255.0
Switch(config-if)# exit

Switch(config)# interface loopback2
Switch(config-if)# ip vrf forwarding v12
Switch(config-if)# ip address 8.8.2.8 255.255.255.0
Switch(config-if)# exit

Switch(config)# interface gigabitethernet1/0/5
Switch(config-if)# switchport trunk encapsulation dot1q
Switch(config-if)# switchport mode trunk
Switch(config-if)# no ip address
Switch(config-if)# exit

Switch(config)# interface fastethernet1/0/8
Switch(config-if)# switchport access vlan 208
Switch(config-if)# no ip address
Switch(config-if)# exit

Switch(config)# interface fastethernet1/0/11
Switch(config-if)# switchport trunk encapsulation dot1q
Switch(config-if)# switchport mode trunk
Switch(config-if)# no ip address
Switch(config-if)# exit
```

スイッチ A で使用する VLAN を設定します。VLAN 10 は、CE と PE 間の VRF 11 によって使用されます。VLAN 20 は、CE と PE 間の VRF 12 によって使用されます。VLAN 118 と 208 は、それぞれスイッチ F とスイッチ D を含む VPN に使用されます。

```
Switch(config)# interface vlan10
Switch(config-if)# ip vrf forwarding v11
Switch(config-if)# ip address 38.0.0.8 255.255.255.0
Switch(config-if)# exit

Switch(config)# interface vlan20
Switch(config-if)# ip vrf forwarding v12
Switch(config-if)# ip address 83.0.0.8 255.255.255.0
Switch(config-if)# exit
```

■ マルチ VRF CE の設定

```
Switch(config)# interface vlan118
Switch(config-if)# ip vrf forwarding v12
Switch(config-if)# ip address 118.0.0.8 255.255.255.0
Switch(config-if)# exit

Switch(config)# interface vlan208
Switch(config-if)# ip vrf forwarding v11
Switch(config-if)# ip address 208.0.0.8 255.255.255.0
Switch(config-if)# exit
```

VPN1 と VPN2 で OSPF ルーティングを設定します。

```
Switch(config)# router ospf 1 vrf v11
Switch(config-router)# redistribute bgp 800 subnets
Switch(config-router)# network 208.0.0.0 0.0.0.255 area 0
Switch(config-router)# exit
Switch(config)# router ospf 2 vrf v12
Switch(config-router)# redistribute bgp 800 subnets
Switch(config-router)# network 118.0.0.0 0.0.0.255 area 0
Switch(config-router)# exit
```

CE/PE ルーティングに BGP を設定します。

```
Switch(config)# router bgp 800
Switch(config-router)# address-family ipv4 vrf v12
Switch(config-router-af)# redistribute ospf 2 match internal
Switch(config-router-af)# neighbor 83.0.0.3 remote-as 100
Switch(config-router-af)# neighbor 83.0.0.3 activate
Switch(config-router-af)# network 8.8.2.0 mask 255.255.255.0
Switch(config-router-af)# exit

Switch(config-router)# address-family ipv4 vrf v11
Switch(config-router-af)# redistribute ospf 1 match internal
Switch(config-router-af)# neighbor 38.0.0.3 remote-as 100
Switch(config-router-af)# neighbor 38.0.0.3 activate
Switch(config-router-af)# network 8.8.1.0 mask 255.255.255.0
Switch(config-router-af)# end
```

スイッチ D の設定

スイッチ D は VPN 1 に属します。次のコマンドを使用して、スイッチ A への接続を設定します。

```
Switch# configure terminal
Enter configuration commands, one per line. End with CNTL/Z.
Switch(config)# ip routing
Switch(config)# interface fastethernet1/0/2
Switch(config-if)# no switchport
Switch(config-if)# ip address 208.0.0.20 255.255.255.0
Switch(config-if)# exit

Switch(config)# router ospf 101
Switch(config-router)# network 208.0.0.0 0.0.0.255 area 0
Switch(config-router)# end
```

スイッチ F の設定

スイッチ F は VPN 2 に属します。次のコマンドを使用して、スイッチ A への接続を設定します。

```
Switch# configure terminal
Enter configuration commands, one per line. End with CNTL/Z.
Switch(config)# ip routing
Switch(config)# interface fastethernet1/0/1
Switch(config-if)# switchport trunk encapsulation dot1q
Switch(config-if)# switchport mode trunk
```

```
Switch(config-if)# no ip address
Switch(config-if)# exit

Switch(config)# interface vlan118
Switch(config-if)# ip address 118.0.0.11 255.255.255.0
Switch(config-if)# exit

Switch(config)# router ospf 101
Switch(config-router)# network 118.0.0.0 0.0.0.255 area 0
Switch(config-router)# end
```

PE スイッチ B の設定

このコマンドをスイッチ B (PE ルータ) で使用すると、CE デバイス、スイッチ A に対する接続だけが設定されます。

```
Router# configure terminal
Enter configuration commands, one per line. End with CNTL/Z.
Router(config)# ip vrf v1
Router(config-vrf)# rd 100:1
Router(config-vrf)# route-target export 100:1
Router(config-vrf)# route-target import 100:1
Router(config-vrf)# exit

Router(config)# ip vrf v2
Router(config-vrf)# rd 100:2
Router(config-vrf)# route-target export 100:2
Router(config-vrf)# route-target import 100:2
Router(config-vrf)# exit

Router(config)# ip cef
Router(config)# interface Loopback1
Router(config-if)# ip vrf forwarding v1
Router(config-if)# ip address 3.3.1.3 255.255.255.0
Router(config-if)# exit

Router(config)# interface Loopback2
Router(config-if)# ip vrf forwarding v2
Router(config-if)# ip address 3.3.2.3 255.255.255.0
Router(config-if)# exit

Router(config)# interface gigabitethernet1/1/0.10
Router(config-if)# encapsulation dot1q 10
Router(config-if)# ip vrf forwarding v1
Router(config-if)# ip address 38.0.0.3 255.255.255.0
Router(config-if)# exit

Router(config)# interface gigabitethernet1/1/0.20
Router(config-if)# encapsulation dot1q 20
Router(config-if)# ip vrf forwarding v2
Router(config-if)# ip address 83.0.0.3 255.255.255.0
Router(config-if)# exit

Router(config)# router bgp 100
Router(config-router)# address-family ipv4 vrf v2
Router(config-router-af)# neighbor 83.0.0.8 remote-as 800
Router(config-router-af)# neighbor 83.0.0.8 activate
Router(config-router-af)# network 3.3.2.0 mask 255.255.255.0
Router(config-router-af)# exit
Router(config-router)# address-family ipv4 vrf v1
Router(config-router-af)# neighbor 38.0.0.8 remote-as 800
Router(config-router-af)# neighbor 38.0.0.8 activate
```



```
Router(config-router-af)# network 3.3.1.0 mask 255.255.255.0
Router(config-router-af)# end
```

マルチ VRF CE ステータスの表示

マルチ VRF CE の設定とステータスに関する情報を表示するには、表 38-15 の特権 EXEC コマンドを使用します。

表 38-15 マルチ VRF CE 情報を表示するコマンド

コマンド	目的
show ip protocols vrf vrf-name	VRF に関するルーティング プロトコル情報を表示します。
show ip route vrf vrf-name [connected] [protocol [as-number]] [list] [mobile] [odr] [profile] [static] [summary] [supernets-only]	VRF に関する IP ルーティング テーブル情報を表示します。
show ip vrf [brief detail interfaces] [vrf-name]	定義した VRF インスタンスに関する情報を表示します。

表示される情報の詳細については、『Cisco IOS Switching Services Command Reference, Release 12.2』を参照してください。

プロトコル独立機能の設定

ここでは、IP ルーティング プロトコルに依存しない機能の設定方法について説明します。これらの機能は、IP ベース イメージまたは IP サービス イメージが稼動するスイッチ上で使用できますが、IP ベース イメージ付属のプロトコル関連機能は RIP だけで使用できます。この章の IP ルーティング プロトコル独立コマンドの詳細については、『Cisco IOS IP Command Reference, Volume 2 of 3: Routing Protocols, Release 12.2』の「IP Routing Protocol-Independent Commands」の章を参照してください。このマニュアルには、Cisco.com ホームページ ([Documentation] > [Cisco IOS Software] > [12.2 Mainline] > [Command References]) からアクセスできます。

ここでは、次の設定情報について説明します。

- 「dCEF の設定」(P.38-95)
- 「等価コスト ルーティング パスの個数の設定」(P.38-96)
- 「スタティック ユニキャスト ルートの設定」(P.38-97)
- 「デフォルトのルートおよびネットワークの指定」(P.38-98)
- 「ルート マップによるルーティング情報の再配信」(P.38-99)
- 「PBR の設定」(P.38-102)
- 「ルーティング情報のフィルタリング」(P.38-106)
- 「認証鍵の管理」(P.38-109)

dCEF の設定

Cisco Express Forwarding (CEF) は、ネットワーク パフォーマンスを最適化するために使用されるレイヤ 3 IP スイッチング技術です。CEF には高度な IP 検索および転送アルゴリズムが実装されているため、レイヤ 3 スイッチングのパフォーマンスを最大化できます。高速スイッチング ルート キャッシュよりも CPU にかかる負担が少ないため、CEF はより多くの CPU 処理能力をパケット転送に振り分けることができます。Catalyst 3750 スイッチ スタックでは、ハードウェアはスタックの dCEF を使用します。動的なネットワークでは、ルーティングの変更によって、高速スイッチング キャッシュ エントリが頻繁に無効となります。高速スイッチング キャッシュ エントリが無効になると、トラフィックがルート キャッシュによって高速スイッチングされずに、ルーティング テーブルによってプロセス スイッチングされることがあります。CEF および dCEF は Forwarding Information Base (FIB; 転送情報ベース) 検索テーブルを使用して、宛先ベースの IP パケット スイッチングを実行します。

dCEF の 2 つの主要な構成要素は、分散 FIB と分散隣接テーブルです。

- FIB はルーティング テーブルや情報ベースと同様、IP ルーティング テーブルに転送情報のミラー イメージが保持されます。ネットワーク内でルーティングまたはトポロジが変更されると、IP ルーティング テーブルがアップデートされ、これらの変更が FIB に反映されます。FIB には、IP ルーティング テーブル内の情報に基づいて、ネクスト ホップのアドレス情報が保持されます。FIB にはルーティング テーブル内の既知のルートがすべて格納されているため、CEF はルート キャッシュをメンテナンスする必要がなく、トラフィックのスイッチングがより効率化され、トラフィック パターンの影響も受けません。
- リンク レイヤ上でネットワーク内のノードが 1 ホップで相互に到達可能な場合、これらのノードは隣接関係にあると見なされます。CEF は隣接テーブルを使用し、レイヤ 2 アドレッシング情報を付加します。隣接テーブルには、すべての FIB エントリに対する、レイヤ 2 のネクスト ホップのアドレスが保持されます。

スイッチ スタックは、ギガビット速度の回線レート IP トラフィックを達成するため Application Specific Integrated Circuit (ASIC; 特定用途向け IC) を使用しているので、dCEF 転送はソフトウェア転送パス (CPU により転送されるトラフィック) にだけ適用されます。

デフォルトで、dCEF はグローバルでイネーブルに設定されています。何らかの理由でこれがディセーブルになった場合は、**ip cef distributed** グローバル コンフィギュレーション コマンドを使用し、再度イネーブルに設定できます。

デフォルト設定では、すべてのレイヤ 3 インターフェイスで dCEF がイネーブルです。**no ip route-cache cef** インターフェイス コンフィギュレーション コマンドを入力すると、ソフトウェアが転送するトラフィックに対して CEF がディセーブルになります。このコマンドは、ハードウェア転送パスには影響しません。CEF をディセーブルにして **debug ip packet detail** 特権 EXEC コマンドを使用すると、ソフトウェア転送トラフィックをデバッグするのに便利です。ソフトウェア転送パス用のインターフェイスで CEF をイネーブルにするには、**ip route-cache cef** インターフェイス コンフィギュレーション コマンドを使用します。



注意

CLI には、インターフェイス上で CEF をディセーブルにする **no ip route-cache cef** インターフェイス コンフィギュレーション コマンドが表示されますが、デバッグ以外の目的でインターフェイス上で dCEF をディセーブルにしないでください。

ディセーブルである dCEF をグローバルにイネーブルにしたり、ソフトウェア転送トラフィックのインターフェイス上でイネーブルにするには、特権 EXEC モードで、次の手順を実行します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	ip cef distributed	CEF の動作をイネーブルにします。
ステップ 3	interface interface-id	インターフェイス コンフィギュレーション モードを開始し、設定するレイヤ 3 インターフェイスを指定します。
ステップ 4	ip route-cache cef	ソフトウェア転送トラフィック用のインターフェイスで CEF をイネーブルにします。
ステップ 5	end	特権 EXEC モードに戻ります。
ステップ 6	show ip cef	すべてのインターフェイスの CEF ステータスを表示します。
ステップ 7	show cef linecard [slot-number] [detail]	スタック内のすべてのスイッチ、または指定されたスイッチに対して、スタック メンバー別に CEF 関連インターフェイス情報を表示します。 (任意) <i>slot-number</i> には、スタック メンバーのスイッチ番号を入力します。
ステップ 8	show cef interface [interface-id]	すべてのインターフェイスまたは指定されたインターフェイスの詳細な CEF 情報を表示します。
ステップ 9	show adjacency	CEF の隣接テーブル情報を表示します。
ステップ 10	copy running-config startup-config	(任意) コンフィギュレーション ファイルに設定を保存します。

等価コスト ルーティング パスの個数の設定

同じネットワークへ向かう同じメトリックのルートが複数ルータに格納されている場合、これらのルートは等価コストを保有していると見なされます。ルーティング テーブルに複数の等価コスト ルートが含まれる場合は、これらをパラレルパスと呼ぶこともあります。ネットワークへの等価コストパスがルータに複数格納されている場合、ルータはこれらを同時に使用できます。パラレルパスを使用すると、パスに障害が発生した場合に冗長性を確保できます。また、使用可能なパスにパケットの負荷を分散し、使用可能な帯域幅を有効利用することもできます。等価コスト ルートは、スタック内の各スイッチでサポートされます。

等価コスト ルートはルータによって自動的に取得、設定されますが、ルーティング テーブルの IP ルーティング プロトコルでサポートされるパラレルパスの最大数は制御可能です。スイッチ ソフトウェアでは最大 32 の等価コスト ルーティングが許可されていますが、スイッチ ハードウェアはルートあたり 17 パス以上は使用しません。

ルーティング テーブルに格納されるパラレルパスのデフォルトの最大数を変更するには、特権 EXEC モードで次の手順を実行します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	router {bgp rip ospf eigrp}	ルータ コンフィギュレーション モードを開始します。
ステップ 3	maximum-paths maximum	プロトコル ルーティング テーブルのパラレルパスの最大数を設定します。指定できる範囲は 1 ～ 16 です。ほとんどの IP ルーティング プロトコルでデフォルトは 4 ですが、BGP の場合は 1 です。
ステップ 4	end	特権 EXEC モードに戻ります。

	コマンド	目的
ステップ 5	show ip protocols	<i>Maximum path</i> フィールドの設定を確認します。
ステップ 6	copy running-config startup-config	(任意) コンフィギュレーション ファイルに設定を保存します。

デフォルト値に戻すには、**no maximum-paths** ルータ コンフィギュレーション コマンドを使用します。

スタティック ユニキャスト ルートの設定

スタティック ユニキャスト ルートは、特定のパスを通過して送信元と宛先間でパケットを送受信するユーザ定義のルートです。ルータが特定の宛先へのルートを構築できない場合、スタティック ルートは重要で、到達不能なすべてのパケットが送信される最終ゲートウェイを指定する場合に有効です。

スタティック ルートを設定するには、特権 EXEC モードで次の手順を実行します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	ip route prefix mask {address interface} [distance]	スタティック ルートを確立します。
ステップ 3	end	特権 EXEC モードに戻ります。
ステップ 4	show ip route	設定を確認するため、ルーティング テーブルの現在のステータスを表示します。
ステップ 5	copy running-config startup-config	(任意) コンフィギュレーション ファイルに設定を保存します。

スタティック ルートを削除するには、**no ip route prefix mask {address | interface}** グローバル コンフィギュレーション コマンドを使用します。

ユーザによって削除されるまで、スタティック ルートはスイッチに保持されます。ただし、管理距離の値を割り当て、スタティック ルートをダイナミック ルーティング情報で上書きできます。各ダイナミック ルーティング プロトコルには、デフォルトの管理距離が設定されています (表 38-16 を参照)。ダイナミック ルーティング プロトコルの情報でスタティック ルートを上書きする場合は、スタティック ルートの管理距離がダイナミック プロトコルの管理距離よりも大きな値になるように設定します。

表 38-16 ダイナミック ルーティング プロトコルのデフォルトの管理距離

ルート送信元	デフォルト距離
接続されたインターフェイス	0
スタティック ルート	1
EIGRP サマリー ルート	5
EBGP	20
内部 EIGRP	90
IGRP	100
OSPF	110
IBGP	200
不明	225

インターフェイスを指し示すスタティック ルートは、RIP、IGRP、およびその他のダイナミック ルーティング プロトコルを通してアドバタイズされます。**redistribute** スタティック ルータ コンフィギュレーション コマンドが、これらのルーティング プロトコルに対して指定されているかどうかは関係ありません。これらのスタティック ルートがアドバタイズされるのは、インターフェイスを指し示すスタティック ルートが接続された結果、静的な性質を失ったとルーティング テーブルで見なされるためです。ただし、**network** コマンドで定義されたネットワーク以外のインターフェイスに対してスタティック ルートを定義する場合は、ダイナミック ルーティング プロトコルに **redistribute** スタティック コマンドを指定しない限り、ルートはアドバタイズされません。

インターフェイスがダウンすると、ダウンしたインターフェイスを経由するすべてのスタティック ルートが IP ルーティング テーブルから削除されます。転送ルータのアドレスとして指定されたアドレスへ向かう有効なネクスト ホップがスタティック ルート内に見つからない場合は、IP ルーティング テーブルからそのスタティック ルートも削除されます。

デフォルトのルートおよびネットワークの指定

ルータが他のすべてのネットワークへのルートを学習することはできません。完全なルーティング機能を実現するには、一部のルータをスマート ルータとして使用し、それ以外のルータのデフォルト ルートをスマート ルータ宛に指定します（スマート ルータには、インターネットワーク全体のルーティング テーブル情報が格納されます）。これらのデフォルト ルートは動的に学習されるか、ルータごとに設定されます。ほとんどのダイナミックな内部ルーティング プロトコルには、スマート ルータを使用してデフォルト情報をダイナミックに生成し、他のルータに転送するメカニズムがあります。

指定されたデフォルト ネットワークに直接接続されたインターフェイスがルータに存在する場合は、そのデバイス上で動作するダイナミック ルーティング プロトコルによってデフォルト ルートが生成されます。RIP の場合は、疑似ネットワーク 0.0.0.0 がアドバタイズされます。

ネットワークのデフォルトを生成しているルータには、そのルータ自身のデフォルト ルートも指定する必要があります。ルータが自身のデフォルト ルートを生成する方法の 1 つは、適切なデバイスを経由してネットワーク 0.0.0.0 に至るスタティック ルートを指定することです。

ネットワークへのスタティック ルートをスタティック デフォルト ルートとして定義するには、特権 EXEC モードで次の手順を実行します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	ip default-network <i>network number</i>	デフォルト ネットワークを指定します。
ステップ 3	end	特権 EXEC モードに戻ります。
ステップ 4	show ip route	最終ゲートウェイで選択されたデフォルト ルートを表示します。
ステップ 5	copy running-config startup-config	(任意) コンフィギュレーション ファイルに設定を保存します。

ルートを削除するには、**no ip default-network *network number*** グローバル コンフィギュレーション コマンドを使用します。

ダイナミック ルーティング プロトコルによってデフォルト情報を送信するときは、特に設定する必要はありません。ルーティング テーブルは定期的にスキャンされ、デフォルト ルートとして最適なデフォルト ネットワークが選択されます。IGRP ネットワークでは、システムのデフォルト ネットワークの候補が複数存在する場合があります。Cisco ルータでは、デフォルト ルートまたは最終ゲートウェイを設定するため、管理距離およびメトリック情報を使用します。

ダイナミックなデフォルト情報がシステムに送信されない場合は、**ip default-network** グローバル コンフィギュレーション コマンドを使用し、デフォルト ルートの候補を指定します。このネットワークが任意の送信元のルーティング テーブルに格納されている場合は、デフォルト ルートの候補としてフ

ラグ付けされます。ルータにデフォルト ネットワークのインターフェイスが存在しなくても、そこへのパスが格納されている場合、そのネットワークは 1 つの候補と見なされ、最適なデフォルト パスへのゲートウェイが最終ゲートウェイになります。

ルート マップによるルーティング情報の再配信

スイッチでは複数のルーティング プロトコルを同時に実行し、ルーティング プロトコル間で情報を再配信できます。ルーティング プロトコル間での情報の再配信は、サポートされているすべての IP ベース ルーティング プロトコルに適用されます。

2 つのドメイン間で拡張パケット フィルタまたはルート マップを定義することにより、ルーティング ドメイン間でルートの再配信を条件付きで制御することもできます。**match** および **set** ルート マップ コンフィギュレーション コマンドは、ルート マップの条件部分を定義します。**match** コマンドは、一致しなければならない条件を指定します。**set** コマンドは、ルーティング アップデートが **match** コマンドによって定義される条件と一致した場合に実行されるアクションを指定します。再配信はプロトコルに依存しない機能ですが、**match** および **set** ルート マップ コンフィギュレーション コマンドの一部は特定のプロトコル固有のものです。

route-map コマンドの後に、**match** コマンドおよび **set** コマンドをそれぞれ 1 つまたは複数指定します。**match** コマンドを指定しない場合は、すべて一致すると見なされます。**set** コマンドを指定しない場合、一致以外の処理はすべて実行されません。このため、少なくとも 1 つの **match** または **set** コマンドを指定する必要があります。



(注)

set ルート マップ コンフィギュレーション コマンドが指定されていないルート マップは CPU に送信され、CPU の利用率が高くなります。

ルートマップ ステートメントは、**permit** または **deny** として識別することもできます。ステートメントが拒否としてマークされている場合、一致基準を満たすパケットは通常の転送チャネルを通じて送り返されます (宛先ベース ルーティング)。ステートメントが許可としてマークされている場合は、一致基準を満たすパケットに **set** コマンドが適用されます。一致基準を満たさないパケットは、通常のルーティング チャネルを通じて転送されます。

BGP ルート マップ **continue** コマンドを使用すると、**match** および **set** コマンドが正常に実行された後、ルート マップの他のエントリを実行することができます。**continue** コマンドを使用することで、よりモジュール化したポリシー定義の構成と編成ができるので、同じルート マップ内に特定のポリシー設定を繰り返す必要がなくなります。スイッチでは、発信ポリシーに **continue** コマンドを使用できます。ルート マップ **continue** コマンドの使用の詳細については、次の URL にある『BGP Route-Map Continue Support for an Outbound Policy feature guide for Cisco IOS Release 12.4(4)T』を参照してください。

http://www.cisco.com/en/US/products/ps6441/products_feature_guides_list.html



(注)

次に示すステップ 3 ~ 14 はそれぞれ任意ですが、少なくとも 1 つの **match** ルート マップ コンフィギュレーション コマンド、および 1 つの **set** ルート マップ コンフィギュレーション コマンドを入力する必要があります。

再配信用のルート マップを設定するには、特権 EXEC モードで次の手順を実行します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	route-map <i>map-tag</i> [permit deny] [<i>sequence number</i>]	再配信を制御するために使用するルート マップを定義し、ルートマップ コンフィギュレーション モードを開始します。 <i>map-tag</i> : ルート マップ用のわかりやすい名前を指定します。 redistribute ルータ コンフィギュレーション コマンドはこの名前を使用して、このルート マップを参照します。複数のルート マップで同じマップ タグ名を共有できます。 (任意) permit が指定され、このルート マップの一致条件が満たされている場合は、 set アクションの制御に従ってルートが再配信されます。 deny が指定されている場合、ルートは再配信されません。 <i>sequence number</i> (任意) : 同じ名前によってすでに設定されているルート マップのリスト内で、新しいルート マップの位置を指定する番号です。
ステップ 3	match as-path <i>path-list-number</i>	BGP AS パス アクセス リストと一致させます。
ステップ 4	match community-list <i>community-list-number</i> [exact]	BGP コミュニティ リストと一致させます。
ステップ 5	match ip address { <i>access-list-number</i> <i>access-list-name</i> } [... <i>access-list-number</i> ... <i>access-list-name</i>]	名前または番号を指定し、標準アクセス リストと一致させます。1 ～ 199 の整数を指定できます。
ステップ 6	match metric <i>metric-value</i>	指定されたルート メトリックと一致させます。 <i>metric-value</i> には、0 ～ 4294967295 の値が指定された、EIGRP のメトリックを指定できます。
ステップ 7	match ip next-hop { <i>access-list-number</i> <i>access-list-name</i> } [... <i>access-list-number</i> ... <i>access-list-name</i>]	指定されたアクセス リスト (番号 1 ～ 199) のいずれかで送信される、ネクスト ホップのルータ アドレスと一致させます。
ステップ 8	match tag <i>tag value</i> [... <i>tag-value</i>]	1 つまたは複数のルート タグ値からなるリスト内の指定されたタグ値と一致させます。0 ～ 4294967295 の整数を指定できます。
ステップ 9	match interface <i>type number</i> [... <i>type number</i>]	指定されたインターフェイスの 1 つから、指定されたネクスト ホップへのルートと一致させます。
ステップ 10	match ip route-source { <i>access-list-number</i> <i>access-list-name</i> } [... <i>access-list-number</i> ... <i>access-list-name</i>]	指定されたアドバタイズ済みアクセス リストによって指定されるアドレスと一致させます。
ステップ 11	match route-type { local internal external [type-1 type-2]}	指定された route-type と一致させます。 • local : ローカルに生成された BGP ルート • internal : OSPF エリア内およびエリア間ルート、または EIGRP 内部ルート • external : OSPF 外部ルート (タイプ 1 またはタイプ 2) または EIGRP 外部ルート
ステップ 12	set dampening <i>halflife reuse suppress max-suppress-time</i>	BGP ルート ダンピング係数を設定します。

	コマンド	目的
ステップ 13	<code>set local-preference value</code>	ローカル BGP パスに値を割り当てます。
ステップ 14	<code>set origin {igp egp as incomplete}</code>	BGP の送信元コードを設定します。
ステップ 15	<code>set as-path {tag prepend as-path-string}</code>	BGP AS パスを変更します。
ステップ 16	<code>set level {level-1 level-2 level-1-2 stub-area backbone}</code>	ルーティング ドメインの指定エリアにアドバタイズされるルートのレベルを設定します。 stub-area および backbone は、OSPF NSSA およびバックボーン エリアです。
ステップ 17	<code>set metric metric value</code>	再配信されるルートに指定するメトリック値を設定します (EIGRP 専用)。 <i>metric value</i> は -294967295 ~ 294967295 の整数です。
ステップ 18	<code>set metric bandwidth delay reliability loading mtu</code>	再配信されるルートに指定するメトリック値を設定します (EIGRP 専用)。 <i>bandwidth</i> : 0 ~ 4294967295 の範囲のルートのメトリック値または IGRP 帯域幅 (キロビット/秒単位) <i>delay</i> : 0 ~ 4294967295 の範囲のルート遅延 (10 マイクロ秒単位) <i>reliability</i> : 0 ~ 255 の数値で表されるパケット伝送の成功可能性。255 は信頼性が 100% であること、0 は信頼性がないことを意味します。 <i>loading</i> : 0 ~ 255 の数値で表されるルートの有効帯域幅 (255 は 100% の負荷) <i>mtu</i> : ルートの MTU の最小サイズ (バイト単位)。範囲は 0 ~ 4294967295 です。
ステップ 19	<code>set metric-type {type-1 type-2}</code>	再配信されるルートに OSPF 外部メトリック タイプを設定します。
ステップ 20	<code>set metric-type internal</code>	ネクスト ホップの IGP メトリックと一致するように、EBGP ネイバーにアドバタイズされるプレフィックスの MED 値を設定します。
ステップ 21	<code>set weight</code>	ルーティング テーブルの BGP ウェイトを設定します。指定できる値は 1 ~ 65535 です。
ステップ 22	<code>end</code>	特権 EXEC モードに戻ります。
ステップ 23	<code>show route-map</code>	設定を確認するため、設定されたすべてのルート マップ、または指定されたルート マップだけを表示します。
ステップ 24	<code>copy running-config startup-config</code>	(任意) コンフィギュレーション ファイルに設定を保存します。

エントリを削除するには、**no route-map map tag** グローバル コンフィギュレーション コマンド、または **no match** や **no set** ルートマップ コンフィギュレーション コマンドを使用します。

ルーティング ドメイン間でルートを配信したり、ルート再配信を制御できます。

ルート再配信を制御するには、特権 EXEC モードで次の手順を実行します。キーワードは前述の手順で定義されたキーワードと同じです。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	router {bgp rip ospf eigrp}	ルータ コンフィギュレーション モードを開始します。
ステップ 3	redistribute protocol [<i>process-id</i>] { level-1 level-1-2 level-2 } [metric <i>metric-value</i>] [metric-type <i>type-value</i>] [match internal external <i>type-value</i>] [tag <i>tag-value</i>] [route-map <i>map-tag</i>] [weight <i>weight</i>] [subnets]	ルーティング プロトコル間でルートを手配します。 route-map を指定しないと、すべてのルートが再配信されます。キーワード route-map に <i>map-tag</i> を指定しないと、ルートは配信されません。
ステップ 4	default-metric number	現在のルーティング プロトコルが、再配信されたすべてのルートに対して同じメトリック値を使用するように設定します (BGP、RIP、OSPF)。
ステップ 5	default-metric bandwidth delay reliability loading mtu	EIGRP ルーティング プロトコルが、EIGRP 以外で再配信されたすべてのルートに対して同じメトリック値を使用するように設定します。
ステップ 6	end	特権 EXEC モードに戻ります。
ステップ 7	show route-map	設定を確認するため、設定されたすべてのルートマップ、または指定されたルート マップだけを表示します。
ステップ 8	copy running-config startup-config	(任意) コンフィギュレーション ファイルに設定を保存します。

再配信をディセーブルにするには、そのコマンドの **no** 形式を使用します。

ルーティング プロトコルのメトリックを、必ずしも別のルーティング プロトコルのメトリックに変換する必要はありません。たとえば、RIP メトリックはホップ カウントで、IGRP メトリックは 5 つの特性の組み合わせです。このような場合は、メトリックを独自に設定し、再配信されたルートに割り当てます。ルーティング情報を制御せずにさまざまなルーティング プロトコル間で交換するとルーティング ループが発生し、ネットワーク動作が著しく低下することもあります。

メトリック変換の代わりに使用されるデフォルトの再配信メトリックが定義されていない場合は、ルーティング プロトコル間で自動的にメトリック変換が発生することもあります。

- RIP はスタティック ルートを自動的に再配信できます。スタティック ルートにはメトリック 1 (直接接続) が割り当てられます。
- デフォルト モードになっている場合、どのプロトコルも他のルーティング プロトコルを再配信できます。

PBR の設定

PBR を使用すると、トラフィック フローに定義済みポリシーを設定できます。PBR を使用してルーティングをより細かく制御するには、ルーティング プロトコルから取得したルートの信頼度を小さくします。PBR は、次の基準に基づいて、パスを許可または拒否するルーティング ポリシーを設定したり、実装したりできます。

- 特定のエンド システムの ID
- 説明
- プロトコル

PBR を使用すると、等価アクセスや送信元依存ルーティング、双方向対パッチ トラフィックに基づくルーティング、専用リンクに基づくルーティングを実現できます。たとえば、在庫記録を本社に送信する場合は広帯域で高コストのリンクを短時間使用し、電子メールなど日常的に使用するアプリケーション データは狭帯域で低コストのリンクで送信できます。

PBR がイネーブルの場合は、Access Control List (ACL; アクセス コントロール リスト) を使用してトラフィックを分類し、各トラフィックがそれぞれ異なるパスを経由するようにします。PBR は着信パケットに適用されます。PBR がイネーブルのインターフェイスで受信されたすべてのパケットは、ルート マップを通過します。ルート マップで定義された基準に基づいて、パケットは適切なネクストホップに転送 (ルーティング) されます。

- パケットがルート マップ ステートメントと一致しない場合は、すべての **set** コマンドが適用されます。
- ステートメントが許可とマークされている場合、どのルートマップ ステートメントとも一致しないパケットは通常の転送チャネルを通じて送信され、宛先ベースのルーティングが実行されます。
- PBR に対して、拒否のマークが付いているルートマップ ステートメントはサポートされていません。

ルート マップの設定の詳細については、「[ルート マップによるルーティング情報の再配信](#)」(P.38-99) を参照してください。

標準 IP ACL を使用すると、アプリケーション、プロトコル タイプ、またはエンド ステーションに基づいて一致基準を指定するように、送信元アドレスまたは拡張 IP ACL の一致基準を指定できます。一致が見つかるまで、ルート マップにこのプロセスが行われます。不一致が見つからない場合は、通常の宛先ベース ルーティングが発生します。**match** ステートメント リストの末尾には、暗黙の拒否エントリがあります。

match コマンドが満たされた場合は、**set** コマンドを使用して、パス内のネクスト ホップ ルータを識別する IP アドレスを指定できます。

PBR コマンドおよびキーワードの詳細については、『*Cisco IOS IP Command Reference, Volume 2 of 3: Routing Protocols, Release 12.2*』を参照してください。表示されているにもかかわらずスイッチでサポートされない PBR コマンドについては、[付録 D「Cisco IOS Release 12.2\(55\)SE でサポートされていないコマンド」](#)を参照してください。

PBR 設定はスタック全体に適用され、すべてのスイッチでスタック マスターの設定が使用されます。



(注) このソフトウェア リリースは、IPv4 および IPv6 トラフィック処理時に PBR をサポートしません。

PBR 設定時の注意事項

PBR の設定を開始する前に、次の点に注意してください。

- PBR を使用するには、スタック マスター上で IP サービス イメージが稼動している必要があります。
- マルチキャスト トラフィックには、ポリシーによるルーティングが行われません。PBR が適用されるのはユニキャスト トラフィックだけです。
- ルーテッド ポートまたは SVI 上で、PBR をイネーブルにできます。
- スイッチは、PBR の **route-map deny** ステートメントをサポートしていません。
- レイヤ 3 モードの EtherChannel ポート チャンネルにはポリシー ルート マップを適用できますが、EtherChannel のメンバーである物理インターフェイスには適用できません。適用しようとする、コマンドが拒否されます。ポリシー ルート マップが適用されている物理インターフェイスは、EtherChannel のメンバーになることができません。
- スイッチ スタックには最大 246 個の IP ポリシー ルート マップを定義できます。

- スイッチ スタックには、PBR 用として最大 512 個の Access Control Entry (ACE; アクセス コントロール エントリ) を定義できます。
- ルート マップに一致基準を設定するときには、次の注意事項に従ってください。
 - ローカル アドレス宛のパケットを許可する ALC と一致させないでください。PBR はこれらのパケットを転送しますが、ping や Telnet 障害またはルート プロトコル フラッピングが発生する可能性があります。
 - 拒否 ACE のある ACL と一致させないでください。拒否 ACE と一致するパケットが CPU に送信されると、CPU の利用率が高くなる可能性があります。
- PBR を使用するには、**sdm prefer routing** グローバル コンフィギュレーション コマンドを使用し、まずルーティング テンプレートをイネーブルにする必要があります。VLAN またはデフォルト テンプレートでは、PBR がサポートされません。SDM テンプレートの詳細については、[第 8 章「SDM テンプレートの設定」](#)を参照してください。
- VRF と PBR は、スイッチ インターフェイス上で相互に排他的です。PBR がインターフェイスでイネーブルになっているときは、VRF をイネーブルにすることはできません。同じように、インターフェイスで VRF がイネーブルになっているときは、PBR をイネーブルにはできません。
- WCCP と PBR は、スイッチ インターフェイス上で相互に排他的です。PBR がインターフェイスでイネーブルになっているときは、WCCP をイネーブルにすることはできません。同じように、インターフェイスで WCCP がイネーブルになっているときは、PBR をイネーブルにすることはできません。
- PBR で使用される Ternary CAM (TCAM) エントリ数は、ルート マップ自体、使用される ACL、ACL およびルート マップ エントリの順序によって異なります。
- パケット長、Type of Service (ToS; サービス タイプ)、set interface、set default next hop、または set default interface に基づく PBR は、サポートされていません。有効な set アクションがないか、または set アクションが *Don't Fragment* に設定されているポリシー マップは、サポートされていません。
- スイッチは PBR ルート マップでの Quality of Service (QoS) DSCP および IP precedence の一致をサポートしていて、次のような制限事項があります。
 - QoS DSCP 変換マップと PBR ルート マップを同じインターフェイスに適用することができません。
 - 透過的な DSCP と PBR DSCP ルート マップを同一スイッチに設定することはできません。
 - PBR と QoS DSCP を設定する際に、QoS をイネーブルに設定 (**mls qos** グローバル コンフィギュレーション コマンドを入力) するか、ディセーブルに設定 (**no mls qos** グローバル コンフィギュレーション コマンドを入力) することができます。QoS がイネーブルの場合、トラフィックの DSCP 値が変更されないようにするには、**mls qos trust dscp** インターフェイス コンフィギュレーション コマンドを入力して、スイッチの入力トラフィック ポートで DSCP 信頼状態を設定します。信頼状態が DSCP でない場合、デフォルトですべての信頼されていないトラフィックの DSCP 値が 0 に設定されます。

PBR のイネーブル化

デフォルトでは、PBR はスイッチ上でディセーブルです。PBR をイネーブルにするには、一致基準およびすべての **match** コマンドと一致した場合の動作を指定するルート マップを作成する必要があります。次に、特定のインターフェイスでそのルート マップ用の PBR をイネーブルにします。指定したインターフェイスに着信したパケットのうち、**match** コマンドと一致したものはすべて PBR の対象になります。

PBR は、スイッチの速度低下を引き起こさない速度で、高速転送したり実装したりできます。高速スイッチングされた PBR では、ほとんどの **match** および **set** コマンドを使用できます。PBR の高速スイッチングをイネーブルにするには、事前に PBR をイネーブルにする必要があります。PBR の高速スイッチングは、デフォルトでディセーブルです。

スイッチで生成されたパケットまたはローカル パケットは、通常どおりにポリシー ルーティングされません。スイッチ上でローカル PBR をグローバルにイネーブルにすると、そのスイッチから送信されたすべてのパケットがローカル PBR の影響を受けます。ローカル PBR は、デフォルトでディセーブルに設定されています。



(注) PBR をイネーブルにするには、スタック マスター上で IP サービス イメージが稼動している必要があります。

PBR を設定するには、特権 EXEC モードで次の手順を実行します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	route-map <i>map-tag</i> [permit] [<i>sequence number</i>]	パケットの出力場所を制御するために使用するルート マップを定義し、ルート マップ コンフィギュレーション モードを開始します。 map-tag : ルート マップ用のわかりやすい名前を指定します。ip policy route-map インターフェイス コンフィギュレーション コマンドはこの名前を使用して、このルート マップを参照します。複数のルート マップで同じマップ タグ名を共有できます。 (任意) permit が指定され、このルート マップの一致条件が満たされている場合は、set アクションの制御に従ってルートがポリシールーティングされます。 (注) route-map deny ステートメントは、インターフェイスに適用される PBR ルート マップでサポートされていません。 sequence number (任意) : 同じ名前によってすでに設定されているルート マップのリスト内で、新しいルート マップの位置を示す番号です。
ステップ 3	match ip address { <i>access-list-number</i> <i>access-list-name</i> } [... <i>access-list-number</i> ... <i>access-list-name</i>]	1 つまたは複数の標準または拡張アクセス リストで許可されている送信元および宛先 IP アドレスを照合します。 (注) 拒否 ACE のある ACL またはローカル アドレス宛のパケットを許可する ACL を入力しないでください。 match コマンドを指定しない場合、ルート マップはすべてのパケットに適用されます。
ステップ 4	set ip next-hop <i>ip-address</i> [... <i>ip-address</i>]	基準と一致するパケットの動作を指定します。パケットのルーティング先となるネクスト ホップを設定します (ネクスト ホップは隣接していなければなりません)。
ステップ 5	exit	グローバル コンフィギュレーション モードに戻ります。
ステップ 6	interface <i>interface-id</i>	インターフェイス コンフィギュレーション モードを開始し、設定するインターフェイスを指定します。

	コマンド	目的
ステップ 7	ip policy route-map <i>map-tag</i>	レイヤ 3 インターフェイス上で PBR をイネーブルにし、使用するルート マップを識別します。1 つのインターフェイスに設定できるルート マップは、1 つだけです。ただし、異なるシーケンス番号を持つ複数のルート マップ エントリを設定できます。これらのエントリは、最初の一致が見つかるまで、シーケンス番号順に評価されます。一致が見つからない場合、パケットは通常どおりにルーティングされます。 (注) IP ポリシー ルート マップに deny ステートメントが含まれる場合、設定に失敗します。
ステップ 8	ip route-cache policy	(任意) PBR の高速スイッチングをイネーブルにします。PBR の高速スイッチングをイネーブルにするには、まず PBR をイネーブルにする必要があります。
ステップ 9	exit	グローバル コンフィギュレーション モードに戻ります。
ステップ 10	ip local policy route-map <i>map-tag</i>	(任意) ローカル PBR をイネーブルにして、スイッチから送信されるパケットに PBR を実行します。ローカル PBR は、スイッチによって生成されるパケットに適用されます。着信パケットには適用されません。
ステップ 11	end	特権 EXEC モードに戻ります。
ステップ 12	show route-map [<i>map-name</i>]	(任意) 設定を確認するため、設定されたすべてのルート マップ、または指定されたルート マップだけを表示します。
ステップ 13	show ip policy	(任意) インターフェイスに付加されたポリシー ルート マップを表示します。
ステップ 14	show ip local policy	(任意) ローカル PBR がイネーブルであるかどうか、およびイネーブルである場合は使用されているルート マップを表示します。
ステップ 15	copy running-config startup-config	(任意) コンフィギュレーション ファイルに設定を保存します。

エントリを削除するには、**no route-map** *map-tag* グローバル コンフィギュレーション コマンド、または **no match** または **no set** ルート マップ コンフィギュレーション コマンドを使用します。インターフェイス上で PBR をディセーブルにするには、**no ip policy route-map** *map-tag* インターフェイス コンフィギュレーション コマンドを使用します。PBR の高速スイッチングをディセーブルにするには、**no ip route-cache policy** インターフェイス コンフィギュレーション コマンドを使用します。スイッチから送信されるパケットに対して PBR をディセーブルにするには、**ip local policy route-map** *map-tag* グローバル コンフィギュレーション コマンドを使用します。

ルーティング情報のフィルタリング

ルーティング プロトコル情報をフィルタリングする場合は、次の作業を実行します。



(注)

OSPF プロセス間でルートが再配信される場合、OSPF メトリックは保持されません。

パッシブ インターフェイスの設定

ローカル ネットワーク上の他のルータが動的にルートを取得しないようにするには、**passive-interface** ルータ コンフィギュレーション コマンドを使用し、ルーティング アップデート メッセージがルータ インターフェイスから送信されないようにします。OSPF プロトコルでこのコマンドを使用すると、パッシブに指定したインターフェイス アドレスが OSPF ドメインのスタブ ネットワークとして表示されます。OSPF ルーティング情報は、指定されたルータ インターフェイスから送受信されません。

多数のインターフェイスが存在するネットワークで、インターフェイスを手動でパッシブに設定する作業を回避するには、**passive-interface default** ルータ コンフィギュレーション コマンドを使用し、すべてのインターフェイスをデフォルトでパッシブになるように設定します。この後で、隣接関係が必要なインターフェイスを手動で設定します。

パッシブ インターフェイスを設定するには、特権 EXEC モードで次の手順を実行します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	router {bgp rip ospf eigrp}	ルータ コンフィギュレーション モードを開始します。
ステップ 3	passive-interface interface-id	指定されたレイヤ 3 インターフェイス経由のルーティング アップデートの送信を抑制します。
ステップ 4	passive-interface default	(任意) すべてのインターフェイスを、デフォルトでパッシブとなるように設定します。
ステップ 5	no passive-interface interface type	(任意) 隣接関係を送信する必要があるインターフェイスだけをアクティブにします。
ステップ 6	network network-address	(任意) ルーティング プロセス用のネットワーク リストを指定します。 <i>network-address</i> は IP アドレスです。
ステップ 7	end	特権 EXEC モードに戻ります。
ステップ 8	copy running-config startup-config	(任意) コンフィギュレーション ファイルに設定を保存します。

パッシブとしてイネーブルにしたインターフェイスを確認するには、**show ip ospf interface** などのネットワーク モニタ用特権 EXEC コマンドを使用します。アクティブとしてイネーブルにしたインターフェイスを確認するには、**show ip interface** 特権 EXEC コマンドを使用します。

ルーティング アップデートの送信を再度イネーブルにするには、**no passive-interface interface-id** ルータ コンフィギュレーション コマンドを使用します。**default** キーワードを指定すると、すべてのインターフェイスがデフォルトでパッシブに設定されます。次に、**no passive-interface** ルータ コンフィギュレーション コマンドを使用し、隣接関係を必要とする各インターフェイスを個別に設定します。**default** キーワードは、ほとんどの配信ルータに 200 を超えるインターフェイスが備わっているインターネット サービス プロバイダーや大規模な企業ネットワークの場合に役立ちます。

ルーティング アップデートのアドバタイズメントおよび処理の制御

ACL と **distribute-list** ルータ コンフィギュレーション コマンドを組み合わせると、ルーティング アップデート中にルート of アドバタイズメントを抑制し、他のルータが 1 つまたは複数のルートを取得しないようにできます。この機能を OSPF で使用した場合は外部ルートだけに適用されるため、インターフェイス名を指定することはできません。

distribute-list ルータ コンフィギュレーション コマンドを使用し、着信したアップデートのリストのうち特定のルート进行处理しないようにすることもできます (OSPF にこの機能は適用されません)。

ルーティング アップデートのアドバタイズメントまたは処理を制御するには、特権 EXEC モードで次の手順を実行します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	router {bgp rip eigrp}	ルータ コンフィギュレーション モードを開始します。
ステップ 3	distribute-list {access-list-number access-list-name} out [interface-name routing process autonomous-system-number]	アクセス リスト内のアクションに応じて、ルーティング アップデート内のルートのアドバタイズメントを許可または拒否します。
ステップ 4	distribute-list {access-list-number access-list-name} in [type-number]	アップデートにリストされたルートの処理を抑制します。
ステップ 5	end	特権 EXEC モードに戻ります。
ステップ 6	copy running-config startup-config	(任意) コンフィギュレーション ファイルに設定を保存します。

フィルタを変更またはキャンセルするには、**no distribute-list in** ルータ コンフィギュレーション コマンドを使用します。アップデート中のネットワーク アドバタイズメントの抑制をキャンセルするには、**no distribute-list out** ルータ コンフィギュレーション コマンドを使用します。

ルーティング情報の送信元のフィルタリング

一部のルーティング情報が他の情報よりも正確な場合があるため、フィルタリングを使用して、さまざまな送信元から送られる情報にプライオリティを設定できます。「管理距離」は、ルータやルータのグループなど、ルーティング情報の送信元の信頼性を示す数値です。大規模ネットワークでは、他のルーティング プロトコルよりも信頼できるルーティング プロトコルが存在する場合があります。管理距離の値を指定すると、ルータはルーティング情報の送信元をインテリジェントに区別できるようになります。常にルーティング プロトコルの管理距離が最短（値が最小）であるルートが選択されます。

表 38-16 (P.38-97) に、さまざまなルーティング情報送信元のデフォルトの管理距離を示しています。

各ネットワークには独自の要件があるため、管理距離を割り当てる一般的な注意事項はありません。

ルーティング情報の送信元をフィルタリングするには、特権 EXEC モードで、次の手順を実行します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	router {bgp rip ospf eigrp}	ルータ コンフィギュレーション モードを開始します。
ステップ 3	distance weight {ip-address {ip-address mask}} [ip access list]	管理距離を定義します。 <i>weight</i> : 管理距離は 10 ～ 255 の整数です。単独で使 用した場合、 <i>weight</i> はデフォルトの管理距離を指定しま す。ルーティング情報の送信元に他の指定がない場合 に使用されます。管理距離が 255 のルートはルーティ ング テーブルに格納されません。 (任意) <i>ip access list</i> : 着信ルーティング アップデート に適用される IP 標準または IP 拡張アクセス リストで す。
ステップ 4	end	特権 EXEC モードに戻ります。

	コマンド	目的
ステップ 5	show ip protocols	指定されたルーティング プロセス用のデフォルトの管理距離を表示します。
ステップ 6	copy running-config startup-config	(任意) コンフィギュレーション ファイルに設定を保存します。

管理距離の定義を削除するには、**no distance** ルータ コンフィギュレーション コマンドを使用します。

認証鍵の管理

鍵管理を使用すると、ルーティング プロトコルで使用される認証鍵を制御できます。一部のプロトコルでは、鍵管理を使用することができません。認証鍵は EIGRP および RIP バージョン 2 で使用できます。

認証鍵を管理する前に、認証をイネーブルにする必要があります。プロトコルに対して認証をイネーブルにする方法については、該当するプロトコルについての説明を参照してください。認証鍵を管理するには、キー チェーンを定義してそのキー チェーンに属する鍵を識別し、各鍵の有効期間を指定します。各鍵には、ローカルに格納される独自の鍵 ID (**key number** キー チェーン コンフィギュレーション コマンドで指定) があります。鍵 ID、およびメッセージに関連付けられたインターフェイスの組み合わせにより、使用中の認証アルゴリズムおよび Message Digest 5 (MD5) 認証鍵が一意に識別されます。

有効期間が指定された複数の鍵を設定できます。存在する有効な鍵の個数に関係なく、1 つの認証パケットだけが送信されます。鍵番号は小さい方から大きい方へ順に調べられ、最初に見つかった有効な鍵が使用されます。鍵変更中は、有効期間が重なっても問題ありません。これらの有効期間は、ルータに通知する必要があります。

認証鍵を管理するには、特権 EXEC モードで次の手順を実行します。

	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	key chain name-of-chain	キー チェーンを識別し、キー チェーン コンフィギュレーション モードを開始します。
ステップ 3	key number	鍵番号を識別します。指定できる範囲は 0 ～ 2147483647 です。
ステップ 4	key-string text	キー スtring を識別します。String には 1 ～ 80 文字の大文字および小文字の英数字を指定できますが、最初の文字に数字を指定することはできません。
ステップ 5	accept-lifetime start-time {infinite end-time duration seconds}	(任意) 鍵を受信する期間を指定します。 <i>start-time</i> および <i>end-time</i> 構文には、 <i>hh:mm:ss Month date year</i> または <i>hh:mm:ss date Month year</i> のいずれかを使用できます。デフォルトはデフォルトの <i>start-time</i> 以降、無制限です。指定できる最初の日付は 1993 年 1 月 1 日です。デフォルトの <i>end-time</i> および <i>duration</i> は infinite です。

	コマンド	目的
ステップ 6	<code>send-lifetime start-time {infinite end-time duration seconds}</code>	(任意) 鍵を送信する期間を指定します。 <i>start-time</i> および <i>end-time</i> 構文には、 <i>hh:mm:ss Month date year</i> または <i>hh:mm:ss date Month year</i> のいずれかを使用できます。デフォルトはデフォルトの <i>start-time</i> 以降、無制限です。指定できる最初の日付は 1993 年 1 月 1 日です。デフォルトの <i>end-time</i> および duration は infinite です。
ステップ 7	<code>end</code>	特権 EXEC モードに戻ります。
ステップ 8	<code>show key chain</code>	認証鍵情報を表示します。
ステップ 9	<code>copy running-config startup-config</code>	(任意) コンフィギュレーション ファイルに設定を保存します。

キー チェーンを削除するには、**no key chain name-of-chain** グローバル コンフィギュレーション コマンドを使用します。

IP ネットワークのモニタおよびメンテナンス

特定のキャッシュ、テーブル、またはデータベースのすべての内容を削除できます。特定の統計情報を表示することもできます。ルートを消去したり、ステータスを表示するには、表 38-17 に示す特権 EXEC コマンドを使用します。

表 38-17 IP ルートの消去またはルート ステータスの表示を行うコマンド

コマンド	目的
<code>clear ip route {network [mask *]}</code>	IP ルーティング テーブルから 1 つまたは複数のルートを消去します。
<code>show ip protocols</code>	アクティブなルーティング プロトコル プロセスのパラメータおよびステータスを表示します。
<code>show ip route [address [mask] [longer-prefixes]] [protocol [process-id]]</code>	ルーティング テーブルの現在のステータスを表示します。
<code>show ip route summary</code>	ルーティング テーブルの現在のステータスをサマリー形式で表示します。
<code>show ip route supernets-only</code>	スーパーネットを表示します。
<code>show ip cache</code>	IP トラフィックのスイッチングに使用されるルーティング テーブルを表示します。
<code>show route-map [map-name]</code>	設定されたすべてのルート マップ、または指定されたルート マップだけを表示します。