



IP ユニキャスト ルーティングの設定

この章では、Catalyst 3560 スイッチに IP 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。スタティック ルーティングおよび Routing Information Protocol (RIP; ルーティング情報プロトコル) などの基本的なルーティング機能は、IP ベース イメージ (以前の Standard Multilayer Image [SMI; 標準マルチレイヤ イメージ]) と IP サービス イメージ (以前の Enhanced Multilayer Image [EMI; 拡張マルチレイヤ 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、スイッチに 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 36 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の詳細については、『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。この章で使用されるコマンドの構文および使用方法の詳細については、次のコマンド リファレンスを参照してください。これらのマニュアルには、Cisco.com の [Documentation] > [Cisco IOS Software] > [12.2 Mainline] > [Command References] からアクセスできます。

- 『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』

この章の内容は、次のとおりです。

- 「IP ルーティングの概要」(P.35-2)
- 「ルーティングを設定する手順」(P.35-3)
- 「IP アドレス指定の設定」(P.35-4)
- 「IP ユニキャスト ルーティングのイネーブル化」(P.35-19)
- 「RIP の設定」(P.35-20)
- 「OSPF の設定」(P.35-26)
- 「EIGRP の設定」(P.35-35)
- 「BGP の設定」(P.35-43)
- 「Multi-VRF CE の設定」(P.35-65)
- 「プロトコル独立機能の設定」(P.35-81)

- ・「IP ネットワークのモニタリングおよびメンテナンス」(P.35-97)



(注)

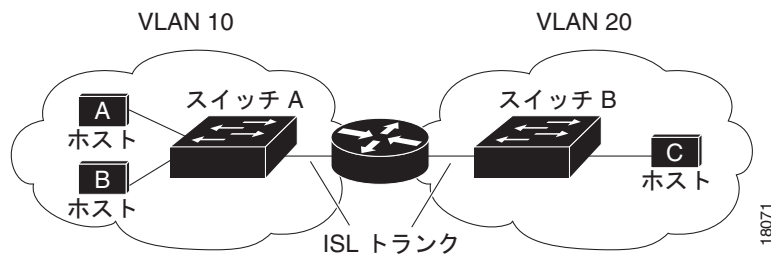
スイッチにルーティング パラメータを設定する場合、使用できるユニキャスト ルート数が最大となるようにシステム リソースを割り当てるには、**sdm prefer routing** グローバル コンフィギュレーション コマンドを使用し、ルーティング テンプレートに Switch Database Management (SDM) 機能を設定します。SDM テンプレートの詳細については、第 7 章「SDM テンプレートの設定」、またはこのリリースのコマンドリファレンスの **sdm prefer** コマンドを参照してください。

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 35-1 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 VLAN 20 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 35-1 ルーティング トポロジの例



VLAN 10 内のホスト A が VLAN 10 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 VLAN 20 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 VLAN 10 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、VLAN20 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。

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。

- ディスタンスベクトル 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。ディスタンスベクトル 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はすばやく対応しますが、ディスタンスベクトル 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トル プロトコルは、RIP および ボーダー ゲートウェイ プロトコル (BGP) 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 ベクトル 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)

サポートされるプロトコルは、スイッチ上で稼働しているソフトウェアによって決まります。スイッチ上で IP ベース イメージが稼働している場合は、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。

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るため、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の詳細については、『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。

以下の手順では、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。詳細については、「[レイヤ 3 EtherChannel の設定](#)」(P.34-13) 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.35-6) を参照してください。



(注)

スイッチは、各ルーテッド ポートおよび SVI に割り当てられた IP アドレスを持つことができます。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 メンバーシップを割り当てます。詳細については、[第 12 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設定します。
- 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 アドレス指定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.35-5)
- 「[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.35-6)
- 「[アドレス解決方法の設定](#)」(P.35-8)
- 「[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.35-11)
- 「[ブロードキャスト パケットの処理方法の設定](#)」(P.35-14)
- 「[IP アドレスのモニタリングおよびメンテナンス](#)」(P.35-18)

アドレス指定のデフォルト設定

表 35-1 に、アドレス指定のデフォルト設定を示します。

表 35-1 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	ARP 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 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 ダイレクトブロードキャスト	ディセーブル（すべての IP ダイレクト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
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、またはユーザ データグラム プロトコル（UDP）フラッド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 スパニングツリー プロトコル（STP）：ディセーブル ターボフラッドディング：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 - アドバタイズ間の最小インターバル：最大インターバルの 0.75 倍 - プリファレンス：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。マスクで、IP アドレス中のネットワーク番号を示すビットが識別できます。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。
ステップ 7	show interfaces [interface-id] show ip interface [interface-id] show running-config interface [interface-id]	入力内容を確認します。
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。

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	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 subnet-zero	インターフェイス アドレスおよびルーティングのアップデート時にサブネット ゼロの使用をイネーブルにします。
ステップ 3	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。

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てに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空間をシミュレートするために使用されるクラス C アドレス空間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 アドレス空間の急速な枯渇を回避するために設計されました。

図 35-2 では、クラスレス ルーティングがイネーブルとなっています。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てにパケットを受信したルータは、パケットを廃棄します。

図 35-2 IP クラスレス ルーティングがイネーブルの場合

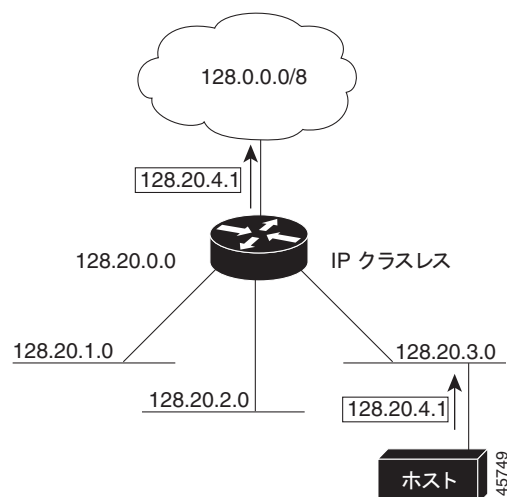
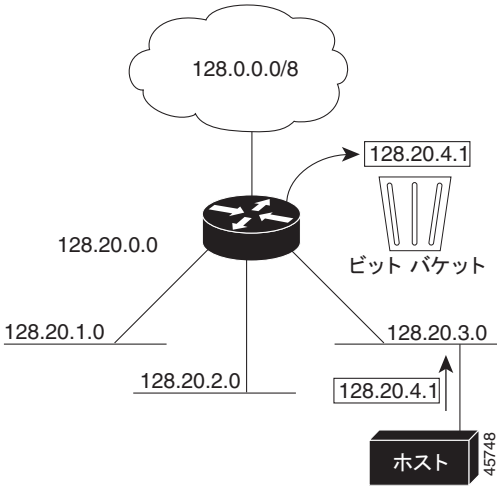


図 35-3 では、ネットワーク 128.20.0.0 のルータはサブネット 128.20.1.0、128.20.2.0、128.20.3.0 に接続されています。ホストがパケットを 128.20.4.1 に送信した場合、ネットワークのデフォルト ルートが存在しないため、ルータはパケットを廃棄します。

図 35-3 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て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。
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て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 (MAC 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。

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 アソシエーションを ARP キャッシュにストアし、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層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、サブネットワーク アクセス プロトコル (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 (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) 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 の詳細については、『*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。

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.35-9)
- 「**ARP カプセル化の設定**」 (P.35-10)
- 「**プロキシ ARP のイネーブル化**」 (P.35-11)

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 アドレス間をダイナミックにマッピングできます。ほとんどのホストではダイナミック 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。デフォルト値は 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 <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	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 [<i>interface-id</i>]	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 <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 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 [<i>interface-id</i>]	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学習できます。

- 「[プロキシ ARP](#)」(P.35-12)

- ・「デフォルト ゲートウェイ」(P.35-12)
- ・「IRDP」(P.35-12)

プロキシ ARP

プロキシ ARP は、他のルートを学習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。最適なルートがある場合、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.35-11) 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かぎり有効です。

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; インターネット制御メッセージプロトコル) 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	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-gateway ip-address	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	end	特権 EXEC モードに戻ります。
ステップ 4	show ip redirects	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	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 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。

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学習します。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; ルーティング情報プロトコル) ルーティングのアップデートを受信し、この情報を使用して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、スイッチにストアされません。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の両方をルータごとに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) アドバタイズが有効である IRDP 期間を設定します。デフォルト値は 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) アドバタイズ間の IRDP の最大インターバルを設定します。デフォルトは 600 秒です。
ステップ 7	ip irdp minadvertinterval seconds	(任意) アドバタイズ間の IRDP の最小インターバルを設定します。デフォルト値は maxadvertinterval 値の 0.75 倍です。 maxadvertinterval 値を変更すると、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 プリファレンス 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、ルータのプリファレンス レベルも高くなります。
ステップ 9	ip irdp address address [number]	(任意) プロキシアドバタイズを行うために必要な IRDP アドレスとプリファレンスを指定します。
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あとで、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てのデータ 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。

- **ダイレクト 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。ダイレクト 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。詳細については、[第 24 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 ケーブルまでの範囲を制限して、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 ストームを伝播します。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。

- 「[ダイレクト 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.35-14)
- 「[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.35-15)
- 「[IP ブロードキャスト アドレスの確立](#)」(P.35-16)
- 「[IP ブロードキャストのフラッドイング](#)」(P.35-17)

ダイレクト 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 ダイレクト 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 ダイレクト ブロードキャストがドロップされると、ルータが DoS 攻撃（サービス拒絶攻撃）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 ダイレクト 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、ダイレクト 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 32 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 ダイレクト 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、ダイレクト 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; VPN ルーティング/転送) 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。ダイレクト ブロードキャストトラフィックが VRF 内でだけ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。

ダイレクト 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

ユーザ データグラム プロトコル (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 フォワーディング 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 (255.255.255.255) です。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 <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 broadcast-address <i>ip-address</i>	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 [<i>interface-id</i>]	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)、ドメイン ネーム システム (DNS)、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。
- パケットの存続可能時間 (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に伝播されるにつれ、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高速化するには、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 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 アドレスのモニタリング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ている場合、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、すべての内容を削除できます。表 35-2 に、内容をクリアするために使用するコマンドを示します。

表 35-2 キャッシュ、テーブル、データベースをクリアするコマンド

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 キャッシュをクリアします。
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を削除します。

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 35-3 に、IP 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 35-3 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
show arp	ARP テーブルのエントリを表示します。
show hosts	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。
show ip aliases	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)。
show ip arp	IP ARP キャッシュを表示します。
show ip interface <i>[interface-id]</i>	インターフェイスの IP ステータスを表示します。
show ip irdp	IRDp 値を表示します。
show ip masks address	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。
show ip redirects	デフォルト ゲートウェイのアドレスを表示します。
show ip route <i>[address [mask]]</i> <i>[protocol]</i>	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 ip route summary	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	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	router ip_routing_protocol	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。特定の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	end	特権 EXEC モードに戻ります。
ステップ 5	show running-config	入力内容を確認します。
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.35-20)
- 「OSPF の設定」(P.35-26)
- 「EIGRP の設定」(P.35-35)
- 「BGP の設定」(P.35-43)
- 「プロトコル独立機能の設定」(P.35-81) (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) です。RIP は、ブロードキャスト ユーザ データグラム プロトコル (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トル 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、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 秒後もまだ更新がない場合、ルータは更新のないルータのルーティング テーブル エントリをすべて削除します。

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、次の設定について説明します。

- 「RIP のデフォルト設定」(P.35-21)
- 「基本的な RIP パラメータの設定」(P.35-21)
- 「RIP 認証の設定」(P.35-23)

- 「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.35-24)

RIP のデフォルト設定

表 35-4 に、RIP のデフォルト設定を示します。

表 35-4 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	- 更新：30 秒 - 無効：180 秒 - ホールドダウン：180 秒 - フラッシュ：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 3560 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	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だけ、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だけ可能です。 (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隣接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。デフォルトは 30 秒です。 • invalid : ルートが無効と宣言されたあとの時間。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ッ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 だけ)、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	入力内容を確認します。
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トを表示するには、**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キーは、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。「[認証キーの管理](#)」(P.35-96) 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>]	入力内容を確認します。
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、ディスタンスベクトル 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の発信元であるインターフェイス上の、ルータによって、その情報が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、アプリケーションがスプリット ホライズンをディセーブルにする必要がある場合を除き、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 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	入力内容を確認します。
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 gi0/2
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、ディスタンスベクトル 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の発信元であるインターフェイス上の、ルータによって、その情報が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、アプリケーションがスプリット ホライズンをディセーブルにする必要がある場合を除き、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	入力内容を確認します。
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.com ページの [Documentation] > [Cisco IOS Software] > [12.2 Mainline] > [Command References] にある『Cisco IOS IP Command Reference, Volume 2 of 3: Routing Protocols, Release 12.2』の「OSPF Commands」を参照してください。



(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。シスコの実装では、RFC1253 の OSPF 管理情報ベース (MIB) がサポートされています。

シスコの実装は、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。
- エリア内の隣接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、ルータのデッド インターバルと hello インターバル、認証キーなどがあります。
- 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、次の設定について説明します。

- 「OSPF のデフォルト設定」(P.35-27)
- 「基本的な OSPF パラメータの設定」(P.35-28)
- 「OSPF インターフェイスの設定」(P.35-29)
- 「OSPF エリア パラメータの設定」(P.35-30)
- 「その他の OSPF パラメータの設定」(P.35-32)
- 「LSA グループ ペーシングの変更」(P.35-33)
- 「ループバック インターフェイスの設定」(P.35-34)
- 「OSPF のモニタリング」(P.35-34)



(注)

OSPF をイネーブルにするには、スイッチ上で IP サービス イメージが稼働している必要があります。

OSPF のデフォルト設定

表 35-5 に、OSPF のデフォルト設定を示します。

表 35-5 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 デッド 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 Mb/s
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 および 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
NSF ¹ 認識	イネーブル ² 。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、隣接する 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。
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

表 35-5 OSPF のデフォルト設定 (続き)

機能	デフォルト設定
タイマー LSA グループのペーシング	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 デッド インターバル : 40 秒 認証キー : キーは未定義 メッセージ ダイジェスト キー (MD5) : キーは未定義

1. NSF = Nonstop Forwarding

2. OSPF NSF 認識は、IP サービス イメージ Cisco IOS Release 12.2(25)SEC 以降を稼働している Catalyst 3550、3560、および 3750 スイッチの IPv4 に対してイネーブルにされています。

OSPF NSF 認識

Cisco IOS Release 12.2(25) SE 以降、IP サービス イメージは IPv4 の OSPF NSF 認識をサポートします。隣接ルータが NSF 対応で、レイヤ 3 スイッチでは、プライマリ 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できません。この機能の詳細については、次の URL の『*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

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	<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>router ospf process-id</code>	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、内部で使用す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。

	コマンド	目的
ステップ3	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4	end	特権 EXEC モードに戻ります。
ステップ5	show ip protocols	入力内容を確認します。
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。

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、デッド インターバル、認証キーなど) 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変更してください。



(注)

ip ospf インターフェイス コンフィギュレーション コマンドはすべてオプションです。

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示的に指定します。
ステップ4	ip ospf retransmit-interval seconds	(任意) LSA 送信間隔を秒数で指定します。指定できる範囲は 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。指定できる範囲は 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。有効な範囲は 0 ～ 255 です。デフォルトは 1 です。

	コマンド	目的
ステップ 7	ip ospf hello-interval <i>seconds</i>	(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。ネットワークのすべてのノードで同じ値を指定する必要があります。指定できる範囲は 1 ～ 65535 秒です。デフォルトは 10 秒です。
ステップ 8	ip ospf dead-interval <i>seconds</i>	(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。ネットワークのすべてのノードで同じ値を指定する必要があります。指定できる範囲は 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 <i>key</i>	(任意) 隣接 OSPF ルータで使用され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隣接ルータには、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 <i>keyid md5 key</i>	(任意) MDS 認証をイネーブルにします。 • <i>keyid</i> : 1 ～ 255 の ID • <i>key</i> : 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 パケットのフラッディングを阻止します。デフォルトでは、OSPF は、LSA が到着した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で新しい LSA をフラッドします。
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 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。出力には、次のいずれかが表示されます。 • <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 認識です。 • <i>Options is 0x42</i> : ネイバー スイッチが NSF 認識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 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、自律システム (AS) 外の宛先に対するデフォルトの外部ルートが、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 の一部がフラッディングされませんが、再配信することによって、エリア内の AS 外部ルートをインポートできます。

経路集約は、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 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	area area-id 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 area-id 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。
ステップ5	area area-id stub [no-summary]	(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	area area-id nssa [no-redistribution] [default-information-originate] [no-summary]	(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インポートする場合に選択します。 default-information-originate : タイプ 7 LSA を NSSA にインポートするようにする場合に、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	area area-id range address mask	(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	end	特権 EXEC モードに戻ります。
ステップ9	show ip ospf [process-id] show ip ospf [process-id [area-id]] database	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	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。

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。

- 経路集約：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.35-85）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ドメイン ネーム サーバ（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。
- アドミニストレーティブ ディスタンスは、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。アドミニストレーティブ ディスタンスが 255 の場合はルーティング情報の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学習した別のルーティング ドメインからのルート（外部）の 3 つの異なるアドミニストレーティブ ディスタンスが使用されます。どのアドミニストレーティブ ディスタンスの値でも変更できます。
- 受動インターフェイス：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、送信側デバイスを受動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ての **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、ルータの変更を詳細に表示できます。

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。

	コマンド	目的
ステップ 3	summary-address <i>address mask</i>	(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 <i>area-id</i> virtual-link <i>router-id</i> [hello-interval <i>seconds</i>] [retransmit-interval <i>seconds</i>] [trans] [[authentication-key <i>key</i>] message-digest-key <i>keyid md5 key</i>]	(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.35-29)、仮想リンクのデフォルト設定については表 35-5 (P.35-27) を参照してください。
ステップ 5	default-information originate [always] [metric <i>metric-value</i>] [metric-type <i>type-value</i>] [route-map <i>map-name</i>]	(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。デフォルトではディセーブルになっています。
ステップ 7	ip auto-cost reference-bandwidth <i>ref-bw</i>	(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 <i>dist1</i>] [inter-area <i>dist2</i>] [external <i>dist3</i>]}	(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 <i>type number</i>	(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。
ステップ 10	timers throttle spf <i>spf-delay spf-holdtime</i> <i>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 です。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 のモニタリング 」(P.35-34) 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 ペーシングの変更

OSPF LSA グループ ペーシング機能を使用すると、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できるようになります。デフォルトでこの機能はイネーブルとなっています。デフォルトのペーシング 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 ペーシング 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、ペーシング 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、ペーシング インターバルを長くし、10 ~ 20 分に設定してください。

OSPF LSA ペーシ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	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	timers lsa-group-pacing seconds	LSA のグループ ペーシングを変更します。
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	入力内容を確認します。
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、他のインターフェイスにより大きな IP アドレスがある場合でも、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	入力内容を確認します。
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 のモニタリング

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 35-6 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 35-6 IP OSPF 統計情報の表示コマンド

コマンド	目的
<code>show ip ospf [process-id]</code>	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。
<code>show ip ospf [process-id] database [router] [link-state-id]</code> <code>show ip ospf [process-id] database [router] [self-originate]</code> <code>show ip ospf [process-id] database [router] [adv-router [ip-address]]</code> <code>show ip ospf [process-id] database [network] [link-state-id]</code> <code>show ip ospf [process-id] database [summary] [link-state-id]</code> <code>show ip ospf [process-id] database [asbr-summary] [link-state-id]</code> <code>show ip ospf [process-id] database [external] [link-state-id]</code> <code>show ip ospf [process-id area-id] database [database-summary]</code>	OSPF データベースに関連する情報を表示します。
<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 ベクトル 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 テクノロジーには、拡散更新アルゴリズム (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、再計算に含まれません。

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：完全更新パケットを受信ごとに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探索メカニズム：このメカニズムを使用し隣接ルータに関する情報を取得します。
- 可変長サブネット マスク (VLSM)
- 任意のルート集約
- 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。

- **ネイバー探索および回復**：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。また、ネイバーが到達不能または動作不能になっていることを検出するためにも使用されます。ネイバー探索および回復は、サイズの小さな hello パケットを定期的に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かぎり、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、隣接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 パケットとユニキャスト パケットが混在した伝送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。効率化のため、信頼性は必要時にのみ提供されます。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。そのため、EIGRP は、1 つのマルチキャスト hello を送信し、パケットに確認応答が必要ないという通知をそのパケットに含めます。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が追跡されます。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隣接ルータです。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。再計算はプロセッサ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 はフィジブルサクセサの有無を調べます。適切なフィジブルサクセサが存在する場合は、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**：ネットワーク層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。また、EIGRP パケットを解析したり、DUAL に受信した新しい情報を通知したりします。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、次の設定について説明します。

- 「EIGRP のデフォルト設定」 (P.35-37)
- 「基本的な EIGRP パラメータの設定」 (P.35-38)
- 「EIGRP インターフェイスの設定」 (P.35-39)
- 「EIGRP ルート認証の設定」 (P.35-40)
- 「EIGRP スタブ ルーティング」 (P.35-41)
- 「EIGRP のモニタリングおよびメンテナンス」 (P.35-42)



(注) EIGRP をイネーブルにするには、スイッチ上で IP サービス イメージが稼働している必要があります。

EIGRP のデフォルト設定

表 35-7 に、EIGRP のデフォルト設定を示します。

表 35-7 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ッ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 以上の kb/s - 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) - MTU : 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)
ディスタンス	内部距離 : 90 外部距離 : 170
EIGRP の隣接関係変更ログ	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 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 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
メトリック重み	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	イネーブル ² 。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、隣接する 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。

表 35-7 EIGRP のデフォルト設定 (続き)

機能	デフォルト設定
オフセットリスト	ディセーブル
ルータ EIGRP	ディセーブル
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 (等コスト ロード バランシング)。

- 1. NSF = Nonstop Forwarding
- 2. EIGRP NSF 認識は、IP サービス イメージ Cisco IOS Release 12.2(25)SEC 以降を稼働している Catalyst 3550、3560、および 3750 スイッチの IPv4 に対してイネーブルにされています。

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、この次の項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.35-25) 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 NSF 認識

EIGRP NSF 認識機能は、Cisco IOS Release 12.2(25)SEC 以降の IP サービス イメージで IPv4 に対してサポートされています。隣接ルータが 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、隣接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できません。この機能の詳細については、次の URL の『*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

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オプションです。

	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>router eigrp autonomous-system number</code>	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。

	コマンド	目的
ステップ 3	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 4	eigrp log-neighbor-changes	(任意) EIGRP 隣接関係変更のログGINGをイネーブルにし、ルーティング システムの安定性を監視します。
ステップ 5	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。
ステップ 6	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 7	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 8	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 9	end	特権 EXEC モードに戻ります。
ステップ 10	show ip protocols	入力内容を確認します。
ステップ 11	show ip protocols	入力内容を確認します。 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 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。

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。

EIGRP の設定

	コマンド	目的
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。指定できる範囲は 1 ～ 65535 秒です。低速 NBMA ネットワークの場合のデフォルト値は 60 秒、その他のすべてのネットワークでは 5 秒です。
ステップ 6	ip hold-time eigrp <i>autonomous-system-number seconds</i>	(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。指定できる範囲は 1 ～ 65535 秒です。低速 NBMA ネットワークの場合のデフォルト値は 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 <i>autonomous-system-number</i>	(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 <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 authentication mode eigrp <i>autonomous-system md5</i>	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 <i>autonomous-system key-chain</i>	IP EIGRP パケットの認証をイネーブルにします。

	コマンド	目的
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 <i>name-of-chain</i>	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 <i>number</i>	キーチェーン コンフィギュレーション モードで、キー番号を識別します。
ステップ 8	key-string <i>text</i>	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。
ステップ 9	accept-lifetime <i>start-time</i> { infinite <i>end-time</i> duration <i>seconds</i> }	(任意) キー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) キー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	認証キー情報を表示します。
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

EIGRP スタブ ルーティング機能は、すべてのイメージで使用でき、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 イメージには EIGRP スタブ ルーティングだけが含まれています。IP サービス イメージには、完全な EIGRP ルーティングが含まれています。
IP ベース イメージが稼働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。

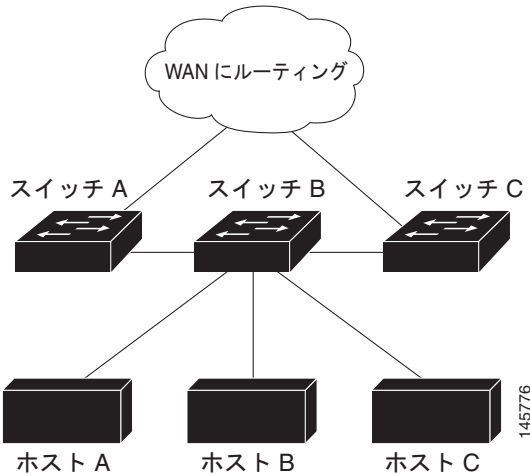
EIGRP スタブ ルーティングを使用するネットワークでは、ユーザに対する IP トラフィックの唯一の許容ルートは、EIGRP スタブ ルーティングを設定しているスイッチ経由です。スイッチは、ユーザ インターフェイスとして設定されているインターフェイスまたは他のデバイスに接続されているインターフェイスにルーテッド トラフィックを送信します。

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 ルータの状態を通知するパケットを受信した隣接ルータは、ルートについてはスタブ ルータに照会しません。また、スタブ ピアを持つルータは、そのピアについては照会しません。スタブ ルータは、ディストリビューション ルータを使用して適切なアップデートをすべてのピアに送信します。

図 35-4 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 と 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（逆の場合も同様です）。

図 35-4 EIGRP スタブ ルータ設定



EIGRP スタブ ルーティングの詳細については、Cisco.com ページの [Documentation] > [Cisco IOS Software] > [12.2 Mainline] > [Configuration Guides] にある『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。

EIGRP のモニタリング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 35-8 に、ネイバー削除および統計情報表示用の特権 EXEC コマンドを示します。表示されるフィールドの詳細については、Cisco.com ページの [Documentation] > [Cisco IOS Software] > [12.2 Mainline] > [Command References] にある『Cisco IOS IP Command Reference, Volume 2 of 3: Routing Protocols, Release 12.2』を参照してください。

表 35-8 IP EIGRP の clear および show コマンド

コマンド	目的
clear ip eigrp neighbors [if-address interface]	ネイバー テーブルからネイバーを削除します。
show ip eigrp interface [interface] [as number]	EIGRP に設定されているインターフェイスに関する情報を表示します。
show ip eigrp neighbors [type-number]	EIGRP によって検出されたネイバーを表示します。

表 35-8 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。

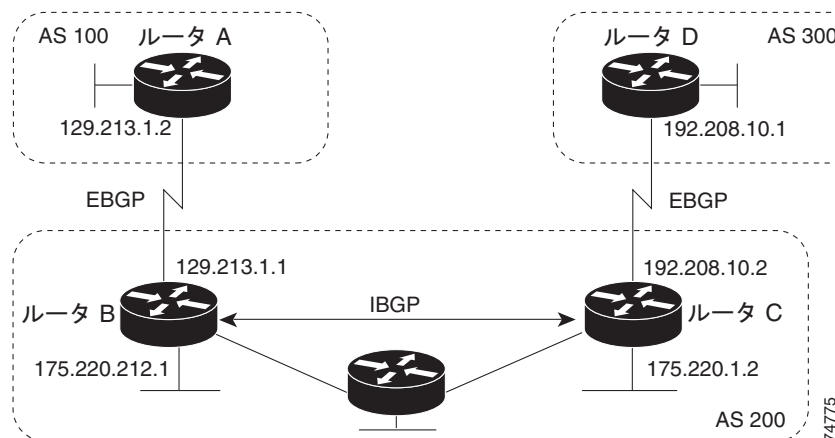
BGP の設定

ボーダー ゲートウェイ プロトコル (BGP) は、Exterior Gateway Protocol です。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。自律システムは、同じ管理下で動作して RIP や OSPF などの Interior Gateway Protocol (IGP) を境界内で実行し、Exterior Gateway Protocol (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。BGP の詳細については、『*Internet Routing Architectures*』(Cisco Press 刊)、および Cisco.com ページの [Documentation] > [Cisco IOS Software] > [12.2 Mainline] > [Configuration Guides] にある『*Cisco IP and IP Routing Configuration Guide*』の「Configuring BGP」の章を参照してください。

BGP コマンドおよびキーワードの詳細については、[Documentation] > [Cisco IOS Software] > [12.2 Mainline] > [Command References] にある『*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。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 C「Cisco IOS リリース 12.2(40)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、同じ自律システムに属するルータは *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 アップデートが自律システム間で交換されるか (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 35-5 に、EBGP と IBGP の両方が稼働するネットワークを示します。

図 35-5 EBGP、IBGP、および複数の自律システム



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 内で稼働する 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 プロトコルとして伝送制御プロトコル (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 35-5 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、ループのない自律システム 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 が稼働しています。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 が稼働し、2 つのネイバーが相互に到達するかぎり、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応答) 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、情報が通過した自律システムのリスト (自律システム パス)、および他のパス属性 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 ループをプルーニングしたり、AS レベル ポリシー判断を行うために使用できます。

Cisco IOS が稼働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 は属性値に基づいてパスを選択します。BGP 属性の詳細については、「BGP 判断属性の設定」(P.35-52) を参照してください。

BGP バージョン 4 では Classless Interdomain Routing (CIDR; クラスレス ドメイン間ルーティング) 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をサポートします。

ここでは、次の設定について説明します。

- 「BGP のデフォルト設定」(P.35-45)
- 「BGP ルーティングのイネーブル化」(P.35-47)
- 「ルーティング ポリシー変更の管理」(P.35-50)
- 「BGP 判断属性の設定」(P.35-52)
- 「ルート マップによる BGP フィルタリングの設定」(P.35-54)
- 「ネイバーによる BGP フィルタリングの設定」(P.35-55)
- 「BGP フィルタリング用のプレフィックス リストの設定」(P.35-56)
- 「BGP コミュニティ フィルタリングの設定」(P.35-57)
- 「BGP ネイバーおよびピア グループの設定」(P.35-58)
- 「集約アドレスの設定」(P.35-60)
- 「ルーティング ドメイン連合の設定」(P.35-61)
- 「BGP ルート リフレクタの設定」(P.35-62)

- 「ルート ダンプニングの設定」 (P.35-63)
- 「BGP のモニタリングおよびメンテナンス」 (P.35-64)

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 Protocols, Release 12.2』を参照してください。これらのマニュアルは、Cisco.com ページの [Documentation] > [Cisco IOS Software] > [12.2 Mainline] > [Configuration Guides] または [Command References] から入手できます。

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 C 「Cisco IOS リリース 12.2(40)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 35-9 に、BGP の基本的なデフォルト設定を示します。すべての特性のデフォルトについては、『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 35-9 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、外部 BGP ピアからの類似ルートは比較されない - 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なし
BGP ルート ダンプニング	デフォルトでは、ディセーブルです。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)
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)

表 35-9 BGP のデフォルト設定 (続き)

機能	デフォルト設定
ディスタンス	- 外部ルート アドミニストレーティブ ディスタンス : 20 (有効値は 1 ~ 255) - 内部ルート アドミニストレーティブ ディスタンス : 200 (有効値は 1 ~ 255) - ローカル ルート アドミニストレーティブ ディスタンス : 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 プレフィックス リスト	未定義
Multi Exit Discriminator (MED)	- 常に比較 : ディセーブル。異なる自律システム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
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 : 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 - コミュニティ属性送信 : 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 : 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 - 重み : 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

表 35-9 BGP のデフォルト設定 (続き)

機能	デフォルト設定
NSF ¹ 認識	ディセーブル ² 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、隣接する 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。
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

2. NSF 認識は、グレースフル リスタートをイネーブルにすることにより、Cisco IOS Release 12.2(25) SEC IP サービス イメージを稼働している Catalyst 3550、3560、および 3750 スイッチの IPv4 に対してイネーブルにできます。

NSF 認識

BGP NSF 認識機能は、Cisco IOS Release 12.2(25)SEC 以降の IP サービス イメージで IPv4 に対してサポートされています。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。隣接ルータが 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、隣接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 の『*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_guide09186a008015fede.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自律システム内にあります。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ドバ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、アドバタイズ対象のルートに矛盾が存在しないことが重要です。BGP がルートをアドバタイズしてから、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、または自律システム内のすべてのルータで BGP が稼働している場合は、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、スイッチ上で IP サービス イメージが稼働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だけ必須)。
ステップ 3	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。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 <i>network-number</i> [<i>mask network-mask</i>] [<i>route-map route-map-name</i>]	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。
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。

	コマンド	目的
ステップ 12	show ip bgp network network-number 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 35-5 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 ピアが稼働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、ピアは稼働していません。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 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 BGP コマンドについては、付録 C「Cisco IOS リリース 12.2(40)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、インバウンドまたはアウトバウンド 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。このあとで BGP フィルタ、重量、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。Cisco IOS Release 12.1 以降では、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、それぞれのアウトバウンドルーティングテーブルをあと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インバウンド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インバウンド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、このリセットはアウトバウンドソフトリセットといいます。

ソフトインバウンドリセットが発生すると、新規インバウンドポリシーが有効になります。ソフトアウトバウンド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、新規ローカルアウトバウンドポリシーが有効になります。アウトバウンド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、新規インバウンドポリシーも有効になる場合があります。

表 35-10 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 35-10 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 FIB 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。	インバウンド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。
ダイナミックインバウンドソフトリセット	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 (Cisco IOS Release 12.1 以降)。

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インバウンド ルーティング テーブルをリセットするには、アウトバウンド 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 判断属性の設定

BGP スピーカーが複数の自律システム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属性値、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 からプレフィックスに対する 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、同じネイバー自律システムから複数の 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属性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 ネクスト ホップ属性 (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. 最大の重みのパスを推奨します (シスコ独自のパラメータ)。ウェイト属性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属性は 32768 で、それ以外のパスのウェイト属性は 0 です。最大の重みのルートを推奨します。重みを設定するには、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 プリファレンス値が最大のルートを推奨します。ローカル プリファレンスはルーティング 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。ローカル初期設定属性のデフォルト値は 100 です。ローカル プリファレンスを設定するには、**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 ルータ上で稼働する 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 メトリック属性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、同じネイバー自律システム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。

同じ判断属性を設定するには、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 <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	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 {<i>ip-address</i> <i>peer-group-name</i>} 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 {<i>ip-address</i> <i>peer-group-name</i>} weight <i>weight</i>	(任意) ネイバー接続に重みを割り当てます。指定できる値は 0 ～ 65535 です。最大の重み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の重みは 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の重みは 32768 です。
ステップ6	default-metric <i>number</i>	(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) 異なる自律システム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 内のパス間でだけ比較されます。
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バ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 プリファレンス値を変更します。指定できる範囲は 0 ～ 4294967295 で、デフォルト値は 100 です。最大のローカル プリファレンス値を推奨します。
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.35-85) 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。 - インバウンド 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 ピアのアウトバウンド 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。
ステップ 5	show route-map [map-name]	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。

	コマンド	目的
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および処理の制御](#)」(P.35-94) 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、さまざまな属性を変更したりできます。ルート マップは、インバウンドアップデートまたはアウトバウンド 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。インバウンドおよびアウトバウンド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。自律システム 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 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	neighbor {ip-address peer-group name} distribute-list {access-list-number name} {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。 (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 {ip-address peer-group name} route-map map-tag {in out}	(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 自律システム 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、正規表現を使用するアクセス リストです。(正規表現の作成方法の詳細については、『*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、自律システム 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ックス 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ックス リストを使用できます。プレフィックス 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、コマンドライン インターフェイス (CLI) 設定が簡素化され、柔軟性が増すなどの利点が生じます。

プレフィックス 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、プレフィックス リストに記載されたプレフィックスとルートとのプレフィックスが照合されます。一致すると、一致したルートが使用されます。プレフィックスが許可されるか、または拒否されるかは、次に示すルールに基づいて決定されます。

- 空のプレフィックス リストはすべてのプレフィックスを許可します。
- 特定のプレフィックスがプレフィックス リストのどのエン트리とも一致しなかった場合、実質的に拒否されたものと見なされます。
- 指定されたプレフィックスと一致するエントリがプレフィックス 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ックス 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ックス リストを使用する場合は、あらかじめプレフィックス リストを設定しておく必要があります。プレフィックス リストを作成したり、プレフィックス 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) するプレフィックス 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、照合するプレフィッ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) プレフィックス 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]	プレフィックス リストまたはプレフィックス 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。

プレフィックス 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。プレフィックス 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。プレフィックス 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 属性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。この属性によって、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属性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属性である、COMMUNITIES 属性 (1 ~ 4294967200 の数値) によって識別されます。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** : どのピア (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** : ローカルな自律システム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 属性に、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** 句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 属性および **match** 句を設定するには、「[ルート マップによるルーティング情報の再配信](#)」(P.35-85) 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 属性はネイバーに送信されません。COMMUNITIES 属性が特定の 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 community-list-number {permit deny} community-number	コミュニティ リストを作成し、番号を割り当てます。 community-list-number 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。 community-number 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 属性を指定します。
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属性から、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 の形式で、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 バイト長形式で表示されます。シスコのデフォルトのコミュニティ形式は、NNAA です。BGP に関する最新の RFC では、コミュニティは AA:NN の形式をとります。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（同じアウトバウンドルート 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、アウトバウンド 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</i> <i>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 属性を指定します。
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。
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ッ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 (パーセンテージ) です。デフォルトは 75% です。

	コマンド	目的
ステップ 14	neighbor { <i>ip-address</i> <i>peer-group-name</i> } next-hop-self	(任意) ネイバー宛での 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 属性を指定します。
ステップ 18	neighbor { <i>ip-address</i> <i>peer-group-name</i> } timers <i>keepalive 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。デフォルト値は 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。デフォルト値は 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重みを指定します。
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に使用する 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ストア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、アトミック集約属性が設定されます。
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ルール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、アップデートされます。
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属性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、自律システムを複数のサブ自律システムに分割して、単一の自律システム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。各自律システムは内部で完全にメッシュ化されていて、同じコンフェデレーション内の他の自律システムとの間には数本の接続があります。異なる自律システム内にあるピアでは 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。具体的には、ネクスト ホップ、MED、およびローカル プリファレンス情報は維持されます。すべての自律システムで単一の IGP を使用できます。

BGP 連合を設定するには、自律システム グループの自律システム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 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 3	bgp confederation identifier autonomous-system	BGP 連合 ID を設定します。
ステップ 4	bgp confederation peers autonomous-system [autonomous-system ...]	連合に属する自律システム、および特殊な EBGp ピアとして処理する自律システムを指定します。
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、内部ネイバーから学習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、学習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 によって学習されたルートを一連の IBGP ネイバーに送信ようになります。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（自律システム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。

- 外部 BGP 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属性を表示します。
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 ダンプニングの設定

ルート フラップ ダンプニングは、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。ルートの状態が使用可能、使用不可能、使用可能、使用不可能という具合に、繰り返し変化する場合、ルートはフラッピングと見なされます。ルート ダンプニング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、ルートが稼働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、ダンプニングが適用されません。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 ダンプニングを設定するには、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 ダンプニング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> [route-map <i>map</i>]	(任意) ルート ダンプニング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 [{ regexp <i>regexp</i> } { filter-list <i>list</i> } { address mask [longer-prefix]}]	(任意) フラッピングしているすべてのパスのフラップを監視します。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、減衰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 <i>regex</i>} {filter-list <i>list</i>} {address mask [<i>longer-prefix</i>]}]</code>	(任意) BGP フラップ統計情報を消去して、ルートが減衰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 ダンプニ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 ダンプニング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。ダンプニ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 のモニタリング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 35-8 に、BGP を消去および表示するために使用する特権 EXEC コマンドを示します。表示フィールドの詳細については、Cisco.com のホームページにアクセスして、『*Cisco IOS IP Command Reference, Volume 2 of 3: Routing Protocols, Release 12.2*』を参照してください ([Documentation] > [Cisco IOS Software] > [12.2 Mainline] > [Command References])。

表 35-11 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>	プレフィッ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 プレフィックスなどのプレフィックス属性も表示されます。
<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>	送信元の自律システム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。
<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。

表 35-11 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] [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ギングをイネーブルにすることもできます。

Multi-VRF CE の設定

バーチャル プライベート ネットワーク (VPN) 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、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 3560 スイッチは、スイッチで IP サービスまたは拡張 IP サービス イメージが稼働している場合に、カスタマー エッジ (CE) デバイスで複数の VPN ルーティング/転送 (Multi-VRF CE) インスタンスをサポートします。IP ベース イメージが稼働しているスイッチでこれを設定しようとすると、エラー メッセージが表示されます。サービス プロバイダーは、Multi-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 イメージが稼働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マルチプロトコル ラベル スイッチング (MPLS) が使用されません。MPLS VRF の詳細については、Cisco.com のホームページにアクセスして、『Cisco IOS Switching Services Configuration Guide, Release 12.2』を参照してください ([Documentation] > [Cisco IOS Software] > [12.2 Mainline] > [Command References])。

ここでは、次の情報について説明します。

- 「Multi-VRF CE の概要」 (P.35-66)
- 「Multi-VRF CE のデフォルト設定」 (P.35-68)
- 「Multi-VRF CE の設定時の注意事項」 (P.35-68)
- 「VRF の設定」 (P.35-69)
- 「VRF 認識サービスの設定」 (P.35-71)
- 「VPN ルーティング セッションの設定」 (P.35-75)
- 「BGP PE/CE ルーティング セッションの設定」 (P.35-76)
- 「Multi-VRF CE の設定例」 (P.35-77)
- 「Multi-VRF CE ステータスの表示」 (P.35-80)

Multi-VRF CE の概要

Multi-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。Multi-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)

Multi-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。

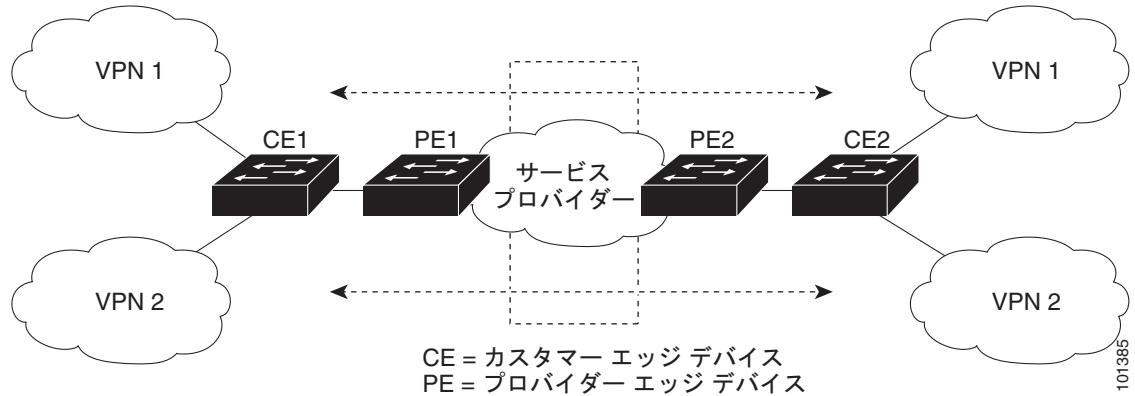
Multi-VRF CE には、次のデバイスが含まれます。

- お客様は、CE デバイスにより、1 つまたは複数のプロバイダー エッジ (PE) 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 3560 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 から学習したあとで、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。

Multi-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。Multi-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ブランチ オフィスに拡張します。

図 35-6 は、Catalyst 3560 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 3560 スイッチでは Multi-VRF CE のサポートが必要です。Multi-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 35-6 複数の仮想 CE として機能する Catalyst 3560 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、Multi-VRF CE 関連のデータ構造で VLAN ID と Policy Label (PL) 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。

Multi-VRF CE を設定すると、レイヤ 3 フォワーディング 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。

- Multi-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 PL 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 PL を取得し、Multi-VRF CE ルーティング セクションに PL 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。

Multi-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 PL 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。ルートが見つかったら、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 PL 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。ルートが見つかったら、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 が優先ルーティング プロトコルです。

Multi-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。

Multi-VRF CE のデフォルト設定

表 35-12 に、VRF のデフォルト設定を示します。

表 35-12 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。

Multi-VRF CE の設定時の注意事項



(注)

Multi-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、次の内容に注意してください。

- Multi-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。
- Multi-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。
- Multi-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、Multi-VRF CE の使用と複数の CE の使用に違いはありません。図 35-6 では、複数の仮想レイヤ 3 インターフェイスが Multi-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かぎり、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 3560 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働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、ルートの属性を CE に簡単に渡すことができます。
- Multi-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。
- Multi-VRF CE 内のライン レート マルチキャスト転送は、Cisco IOS Release 12.2(40) SE 以降でサポートされています。
- 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、104 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 をイネーブルにできません。
- VRF とポリシーベース ルーティング (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 をイネーブルにはできません。その反対の場合も同じで、VRF がインターフェイスでイネーブルになっているときは、PBR をイネーブルにできません。
- 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。

	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。

	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	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 8	ip vrf forwarding <i>vrf-name</i>	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] [<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 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 テーブル内で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 <i>vrf-name</i>	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 <i>route-distinguisher</i>	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} <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 <i>vrf-name</i> distributed	(任意) VRF テーブルのグローバルなマルチキャスト ルーティングをイネーブルにします。

	コマンド	目的
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 <i>vrf-name</i>	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 <i>ip-address mask</i>	レイヤ 3 インターフェイスに IP アドレスを設定します。
ステップ 11	ip pim sparse-dense mode	VRF に関連付けられたレイヤ 3 インターフェイスで 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。

Multi-VRF CE 内でのマルチキャスト設定の詳細については、『*Cisco IOS IP Multicast 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 インターフェイスに設定可能で、グローバル ルーティング インスタンスで稼働します。IP サービスは複数のルーティング インスタンス上で稼働するように拡張されます。これが、VRF 認識です。システム内の任意の設定済み VRF であればいずれも、VRF 認識サービス用に指定できます。

VRF 認識サービスは、プラットフォームに依存しないモジュールに実装されます。VRF とは、Cisco IOS 内の複数のルーティング インスタンスを意味します。各プラットフォームには、サポートする VRF 数に関して独自の制限があります。

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 エントリを表示できます。

次のサービスは、VRF 認識です。

- ARP
- ping
- 簡易ネットワーク管理プロトコル (SNMP)
- ホットスタンバイ ルータ プロトコル (HSRP)
- Unicast Reverse Path Forwarding (uRPF; ユニキャスト RPF)
- Syslog
- traceroute
- FTP および TFTP

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
show ip arp vrf <i>vrf-name</i>	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
ping vrf <i>vrf-name</i> ip-host	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	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	snmp-server trap authentication vrf	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	snmp-server engineID remote <host> vrf <vpn instance> <engine-id string>	スイッチ上のリモート SNMP エンジンの名前を設定します。
ステップ 4	snmp-server host <host> vrf <vpn instance> traps <community>	SNMP トラップ動作の受信側、および 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	snmp-server host <host> vrf <vpn instance> informs <community>	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	snmp-server user <user> <group> remote <host> vrf <vpn instance> <security model>	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	end	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、確実に適切な 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 <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	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 vrf forwarding <<i>vrf-name</i>>	インターフェイス上に VRF を設定します。
ステップ5	ip address <i>ip address</i>	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 <i>ip address</i>	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。

uRPF のユーザ インターフェイス

uRPF は、VRF に割り当てられたインターフェイス上で設定でき、送信元検索が VRF テーブルで実行されます。

uRPF 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	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 vrf forwarding <vrf-name>	インターフェイス上に VRF を設定します。
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	ip verify unicast reverse-path	インターフェイス上で uRPF をイネーブルにします。
ステップ 7	end	特権 EXEC モードに戻ります。

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 メッセージのロギングをイネーブルにする、または一時的にディセーブルにします。
ステップ 3	logging host ip address vrf vrf name	ロギング 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 ロギング メッセージをロギング 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。

traceroute のユーザ インターフェイス

traceroute 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
traceroute 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、そ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 および TFTP を VRF 認識とするには、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付加される 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 [t]ftp に通知する必要があります。この例では、VRF テーブルが宛先 IP アドレスを検索するのに使用されます。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を及ぼしません。つまり、VRF がそのインターフェイスに設定されていない場合でも、送信元インターフェイス CLI を使用して、特定のインターフェイスにパケットを送信できます。

FTP 接続の送信元 IP アドレスを指定するには、**ip ftp source-interface show** モード コマンドを使用します。接続が確立さ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特定の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** モード 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。

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) 隣接状態の変更をログします。これは、デフォルトの状態です。
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。

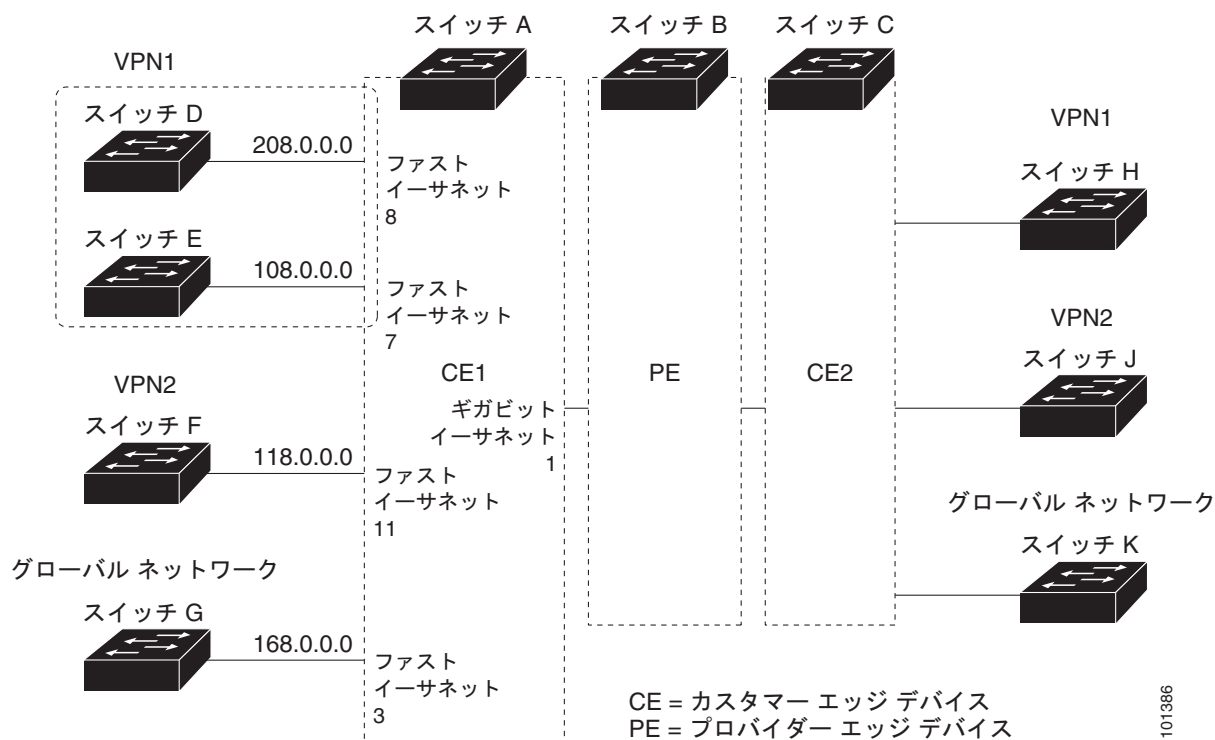
	コマンド	目的
ステップ 1	configure terminal	グローバル コンフィギュレーション モードを開始します。
ステップ 2	router bgp autonomous-system-number	その他の BGP ルータに AS 番号を渡す BGP ルーティング プロセスを設定し、ルータ コンフィギュレーション モードを開始します。
ステップ 3	network network-number mask network-mask	ネットワークとマスクを指定し、BGP の使用を宣言します。
ステップ 4	redistribute ospf process-id match internal	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。
ステップ 6	address-family ipv4 vrf vrf-name	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	neighbor address remote-as as-number	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	neighbor address activate	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	end	特権 EXEC モードに戻ります。
ステップ 10	show ip bgp [ipv4] [neighbors]	BGP 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。

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。

Multi-VRF CE の設定例

図 35-7 は、図 35-6 と同じネットワークの物理接続を単純化した例です。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 3560 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 35-7 Multi-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 gigabitethernet0/5
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 fastethernet0/8
Switch(config-if)# switchport access vlan 208
Switch(config-if)# no ip address
Switch(config-if)# exit

Switch(config)# interface fastethernet0/11
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

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 fastethernet0/2
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 fastethernet0/1
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 gigabitethernet1/0.10
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 gigabitethernet1/0.20
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
```

Multi-VRF CE ステータスの表示

Multi-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 35-13 の特権 EXEC コマンドを使用します。

表 35-13 Multi-VRF CE 情報を表示するコマンド

コマンド	目的
show ip protocols vrf <i>vrf-name</i>	VRF に関するルーティング プロトコル情報を表示します。

表 35-13 Multi-VRF CE 情報を表示するコマンド（続き）

コマンド	目的
<code>show ip route vrf vrf-name [connected] [protocol [as-number]] [list] [mobile] [odr] [profile] [static] [summary] [supernets-only]</code>	VRF に関する IP ルーティング テーブル情報を表示します。
<code>show ip vrf [brief detail interfaces] [vrf-name]</code>	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 イメージが稼働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 でだけ使用できます。この章に記載されている 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.com のホームページにアクセスして、『Cisco IOS IP Command Reference, Volume 2 of 3: Routing Protocols, Release 12.2』の「IP Routing Protocol-Independent Commands」を参照してください（[Documentation] > [Cisco IOS Software] > [12.2 Mainline] > [Command References]）。

ここでは、次の設定について説明します。

- 「CEF の設定」(P.35-81)
- 「等コスト ルーティング パスの個数の設定」(P.35-83)
- 「スタティック ユニキャスト ルートの設定」(P.35-83)
- 「デフォルトのルートおよびネットワークの指定」(P.35-85)
- 「ルート マップによるルーティング情報の再配信」(P.35-85)
- 「ポリシーベース ルーティングの設定」(P.35-89)
- 「ルーティング情報のフィルタリング」(P.35-93)
- 「認証キーの管理」(P.35-96)

CEF の設定

Cisco Express Forwarding (CEF; シスコ エクスプレス フォワーディング) 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割り当てることができます。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 は FIB 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。

CEF 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。
- リンク層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、CEF 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、CEF はグローバルに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** 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 CEF 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は、インターフェイス上で CEF をディセーブルにしないようにしてください。

ディセーブルである CEF 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したりするには、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	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 [detail]	CEF 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 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。

	コマンド	目的
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 の場合だけ 1 です。
ステップ 4	end	特権 EXEC モードに戻ります。
ステップ 5	show ip 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。

	コマンド	目的
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、アドミニストレーティブ ディスタンスの値を割り当て、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、デフォルトのアドミニストレーティブ ディスタンスが設定されています (表 35-14 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 ルートのアドミニストレーティブ ディスタンスがダイナミック プロトコルのアドミニストレーティブ ディスタンスよりも大きな値になるように設定します。

表 35-14 ダイナミック ルーティング プロトコルのデフォルトのアドミニストレーティブ ディスタンス

ルートの送信元	デフォルト距離
接続されているインターフェイス	0
スタティック ルート	1
EIGRP サマリー ルート	5
外部 BGP	20
内部 EIGRP	90
IGRP	100
OSPF	110
内部 BGP	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かぎり、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

ルータは、他のすべてのネットワークへのルートを学習できません。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 ルータ宛てに指定します（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 ルートは動的に取得されるか、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動的に生成し、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 network number	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もあります。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、アドミニストレーティブ ディスタンス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 マップの条件部を定義します。**match** コマンドは、条件

が一致する必要があることを指定しています。**set** コマンドは、ルーティング アップデートが **match** コマンドで定義した条件を満たす場合に行われる処理を指定します。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あります。

route-map コマンドのあとに、**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。

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** 句を使用して、エントリが無事に一致して句が設定された後、ルートマップの追加のエントリを実行できます。**continue** 句を使用すれば、同じルート マップ内で特定のポリシー コンフィギュレーションを繰り返す必要がないように、より多くのモジュラ ポリシー定義を設定および編成できます。Cisco IOS Release 12.2(37)SE 以降のスイッチでは、発信ポリシーに対して **continue** 句がサポートされています。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 で、『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 map-tag [permit deny] [sequence number]	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 path-list-number	BGP AS パス アクセス リストと一致させます。

	コマンド	目的
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 type <i>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 ダンプニング係数を設定します。
ステップ 13	set local-preference <i>value</i>	ローカル BGP パスに値を割り当てます。
ステップ 14	set origin { igp egp <i>as</i> incomplete }	BGP の送信元コードを設定します。
ステップ 15	set as-path { tag prepend <i>as-path-string</i> }	BGP 自律システム パスを変更します。
ステップ 16	set level { level-1 level-2 level-1-2 stub-area backbone }	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	set metric <i>metric value</i>	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。

	コマンド	目的
ステップ 18	set metric <i>bandwidth delay reliability loading mtu</i>	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	set metric-type { <i>type-1</i> <i>type-2</i> }	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	set metric-type <i>internal</i>	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 multi-exit discriminator (MED) 値を設定します。
ステップ 21	set weight	ルーティング テーブルの BGP 重みを設定します。指定できる値は 1 ~ 65535 です。
ステップ 22	end	特権 EXEC モードに戻ります。
ステップ 23	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 24	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** 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 { <i>bgp</i> <i>rip</i> <i>ospf</i> <i>eigrp</i> }	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>] { <i>level-1</i> <i>level-1-2</i> <i>level-2</i> } [<i>metric metric-value</i>] [<i>metric-type type-value</i>] [<i>match internal</i> <i>external type-value</i>] [<i>tag tag-value</i>] [<i>route-map map-tag</i>] [<i>weight weight</i>] [<i>subnets</i>]	ルーティング プロトコル間でルートを再配信します。 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)。

	コマンド	目的
ステップ5	default-metric <i>bandwidth delay reliability loading mtu</i>	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 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があります。

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があります。

- 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。

ポリシーベース ルーティングの設定

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
- アプリケーション
- 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、インタラクティブ対バッチ 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高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 データは低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、アクセス コントロール リスト (ACL) 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** 句が適用されます。

- ・ ステートメントが許可として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 では、拒否としてマークされているルートマップ ステートメントはサポートされません。

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.35-85)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。一致が見つからない場合、通常の宛先ベース ルーティングが行われます。match ステートメント リストの末尾には、暗黙の拒否ステートメントがあります。

match 句が満たされた場合は、set 句を使用して、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 C「Cisco IOS リリース 12.2\(40\)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。



(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、スイッチ上で IP サービス イメージが稼働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。
- ・ PBR には、**route-map deny** ステートメントはサポートされません。
- ・ 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。適用しようとすると、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。
- ・ スイッチには、PBR 用として最大 512 個のアクセス コントロール エントリ (ACE) を定義できます。
- ・ ルート マップに一致基準を設定する場合は、次の注意事項に従ってください。
 - － ローカル アドレス宛てのパケットを許可する ACL と照合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 を含む ACL と照合させないでください。拒否 ACE と一致するパケットが CPU に送られるため、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 コマンドを使用して、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 7 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 をイネーブルにはできません。その反対の場合も同じで、VRF がインターフェイスでイネーブルになっているときは、PBR をイネーブルにできません。
- WCCP と PBR は、スイッチ インターフェイス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 をイネーブルにできません。その反対の場合も同じで、WCCP がインターフェイスでイネーブルになっているときは、PBR をイネーブルにできません。
- PBR で使用される Ternary CAM (TCAM; 3 値連想メモリ) 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; タイプ オブ サービス)、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。
- Cisco IOS Release 12.2(35) SE 以降のスイッチでは、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、次の制限があります。
 - 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できません。
 - DSCP 透過性と PBR DSCP ルート マップを同じスイッチ上に設定できません。
 - QoS DSCP を含む PBR を設定する場合、QoS がイネーブル (**mls qos** グローバル コンフィギュレーション コマンドを使用) またはディセーブル (**no mls qos** コマンドを使用) になるように設定できます。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ないトラフィックは DSCP 値が 0 としてマークされることになります。

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** 句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** 句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、スイッチ上で IP サービス イメージが稼働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 map-tag [permit] [sequence number]	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する PBR ルート マップではサポートされません。 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 {access-list-number access-list-name} [...access-list-number ...access-list-name]	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 を含む ACL またはローカル アドレス宛てのパケットを許可する ACL は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 ip-address [...ip-address]	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 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 7	ip policy route-map map-tag	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ていると、その設定は失敗します。
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。

	コマンド	目的
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、以下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。

受動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。このあとで、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。

受動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 モニタリング用特権 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イズおよび処理の制御

アクセス コントロール リストと **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のアドバイズを抑制し、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にだけ適用されるため、インターフェイス名を指定できません。

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イズ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イズ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。

	コマンド	目的
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。

ルーティング情報発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。「アドミニストレーティブ ディスタンス」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。アドミニストレーティブ ディスタンスの値を指定すると、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 プロトコルのアドミニストレーティブ ディスタンスが最短（値が最小）であるルートが選択されます。[表 35-14 \(P.35-84\)](#) に、さまざまなルーティング情報送信元のデフォルトのアドミニストレーティブ ディスタンスを示します。

各ネットワークには独自の要件があるため、アドミニストレーティブ ディスタンス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]	アドミニストレーティブ ディスタンスを定義します。 <i>weight</i> : アドミニストレーティブ ディスタンスは 10 ～ 255 の整数です。単独で使した場合、 <i>weight</i> はデフォルトのアドミニストレーティブ ディスタンスを指定します。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。アドミニストレーティブ ディスタンスが 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。
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アドミニストレーティブ ディスタンスを表示します。
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。

アドミニストレーティブ ディスタンスを削除するには、**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。

認証キーの管理

キー管理を使用すると、ルーティング プロトコルで使用される認証キーを制御できます。一部のプロトコルでは、キー管理を使用できません。認証キーは EIGRP および RIP バージョン 2 で使用できません。

認証キーを管理する前に、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。認証キーを管理するには、キー チェーンを定義してそのキー チェーンに属するキーを識別し、各キーの有効期間を指定します。各キーには、ローカルにストアされる独自のキー 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。キー 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) 認証キーが一意に識別されます。

有効期間が指定された複数のキーを設定できます。存在する有効なキーの数にかかわらず、送信される認証パケットは 1 つだけです。最小の番号から順にキー番号が調べられ、最初に見つかった有効なキーが使用されます。キー変更中は、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。

認証キーを管理するには、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	キー番号を識別します。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できません。
ステップ 5	accept-lifetime start-time {infinite end-time duration seconds}	(任意) キー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 6	send-lifetime start-time {infinite end-time duration seconds}	(任意) キー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	end	特権 EXEC モードに戻ります。
ステップ 8	show key chain	認証キー情報を表示します。
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。

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リング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。ルートを削除したり、ステータスを表示するには、表 35-15 に示す特権 EXEC コマンドを使用します。

表 35-15 IP ルートの削除またはルート ステータスの表示を行うコマンド

コマンド	目的
clear ip route { <i>network</i> [<i>mask</i> *]}	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削除します。
show ip protocols	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トを表示します。
show ip route [<i>address</i> [<i>mask</i>] [longer-prefixes]] [<i>protocol</i> [<i>process-id</i>]]	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 ip route summary	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。
show ip route supernets-only	スーパーネットを表示します。
show ip cache	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。
show route-map [<i>map-name</i>]	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。

