



ACI Transit Routing, Route Peering, and EIGRP Support

This chapter contains the following sections:

- [ACI Transit Routing, on page 1](#)
- [Transit Routing Use Cases, on page 1](#)
- [ACI Fabric Route Peering, on page 6](#)
- [Transit Route Control, on page 11](#)
- [Default Policy Behavior, on page 13](#)
- [EIGRP Protocol Support, on page 13](#)

ACI Transit Routing

The ACI fabric supports transit routing, which enables border routers to perform bidirectional redistribution with other routing domains. Unlike the stub routing domains of earlier ACI releases, that block transit redistribution, bidirectional redistribution passes routing information from one routing domain to another. Such redistribution lets the ACI fabric provide full IP connectivity between different routing domains. Doing so can also provide redundant connectivity by enabling backup paths between routing domains.

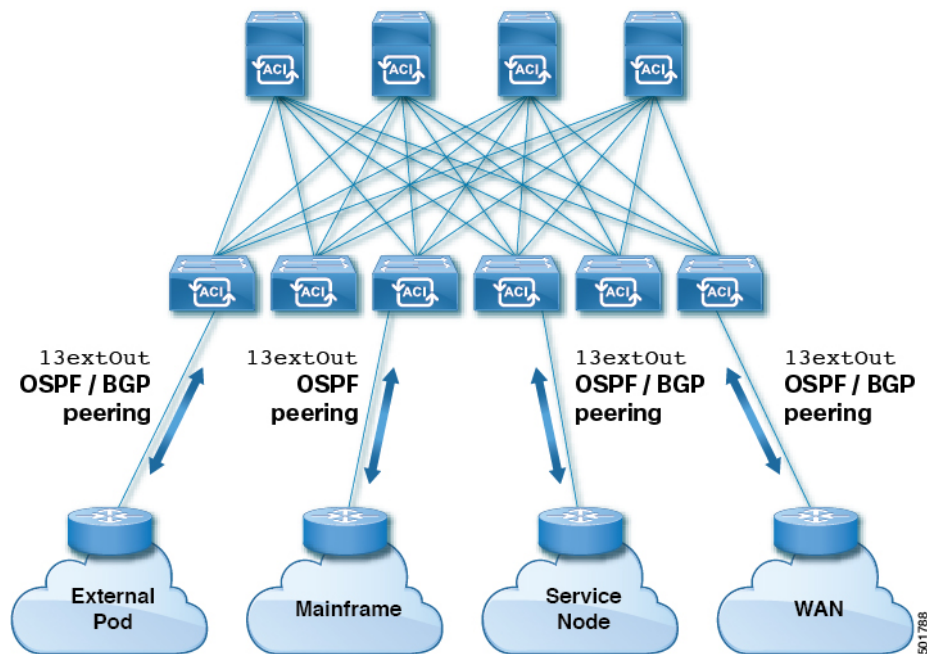
Design transit redistribution policies that avoid suboptimal routing or the more serious problem of routing loops. Typically, transit redistribution does not preserve the original topology and link-state information and redistributes external routes in distance-vector fashion (routes are advertised as vector prefixes and associated distances even with link-state protocols). Under these circumstances, the routers can inadvertently form routing loops that fail to deliver packets to their destination.

Transit Routing Use Cases

Transit Routing Between Layer 3 Domains

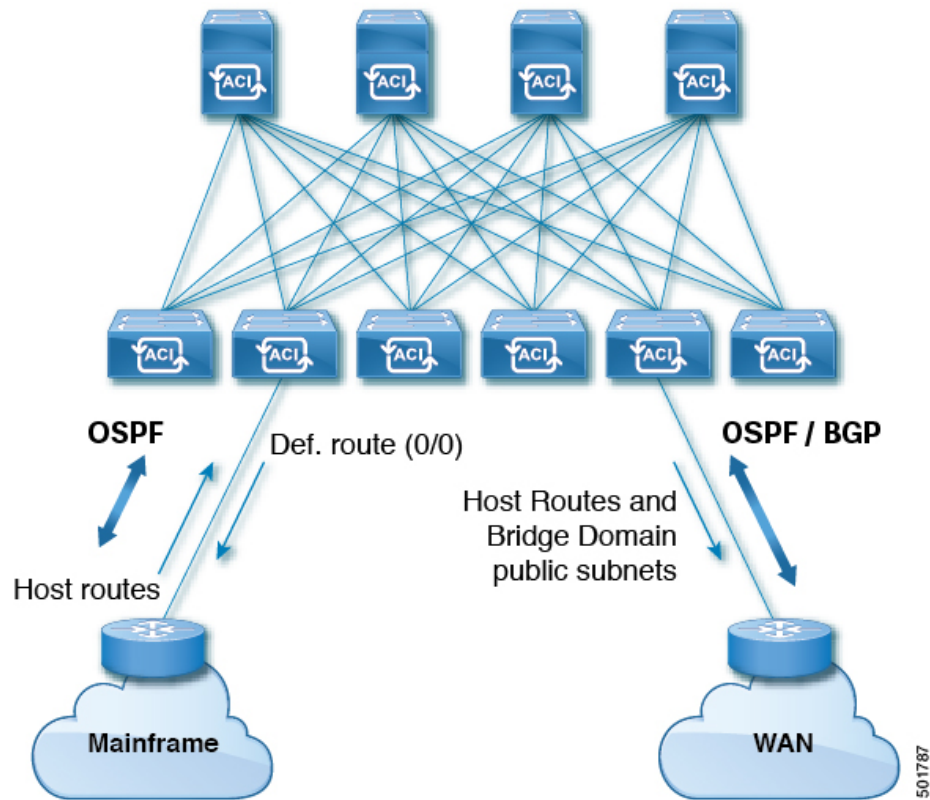
Multiple Layer 3 domains such as external pods, mainframes, service nodes, or WAN routers can peer with the ACI fabric to provide transit functionality between them.

Figure 1: Transit Routing Between Layer 3 Domains



Mainframe Traffic Transiting the ACI Fabric

Mainframes can function as IP servers running standard IP routing protocols that accommodate requirements from Logical Partitions (LPARs) and Virtual IP Addressing (VIPA).

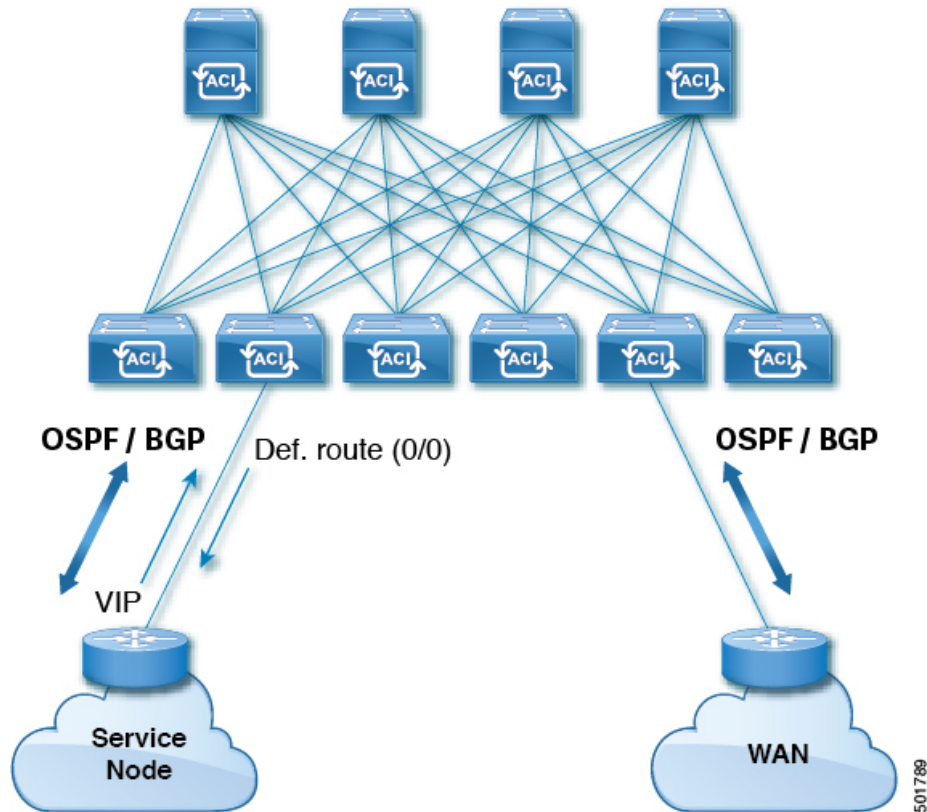
Figure 2: Mainframe Transit Connectivity

In this topology, mainframes require the ACI fabric to be a transit domain for external connectivity through a WAN router and for east-west traffic within the fabric. They push host routes to the fabric to be redistributed within the fabric and out to external interfaces.

Service Node Transit Connectivity

Service nodes can peer with the ACI fabric to advertise a Virtual IP (VIP) route that is redistributed to an external WAN interface.

Figure 3: Service Node Transit Connectivity

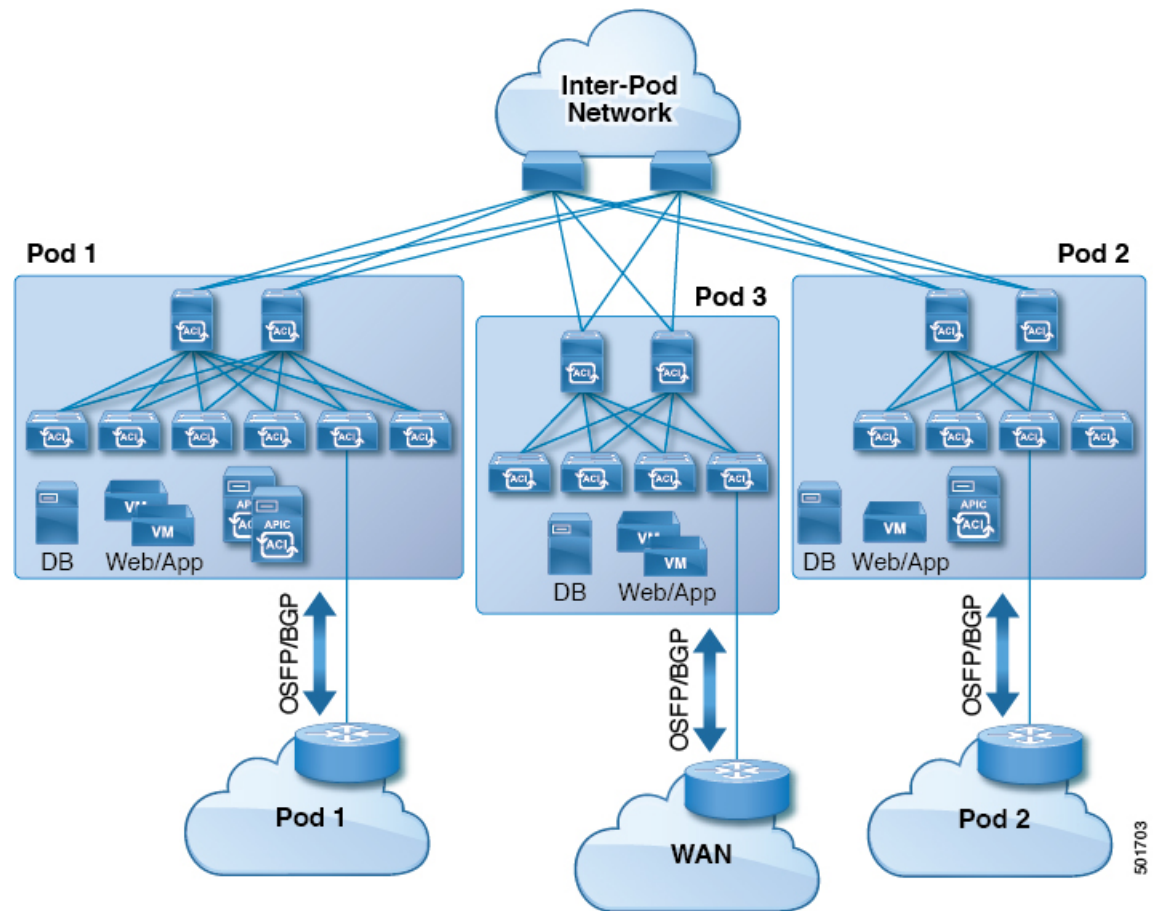


The VIP is the external facing IP address for a particular site or service. A VIP is tied to one or more servers or nodes behind a service node.

Multipod in a Transit-Routed Configuration

In a multipod topology, the fabric acts as a transit for external connectivity and interconnection between multiple pods. Cloud providers can deploy managed resource pods inside a customer datacenter. The demarcation point can be an L3Out with OSPF or BGP peering with the fabric.

Figure 4: Multiple Pods with L3Outs in a Transit-Routed Configuration



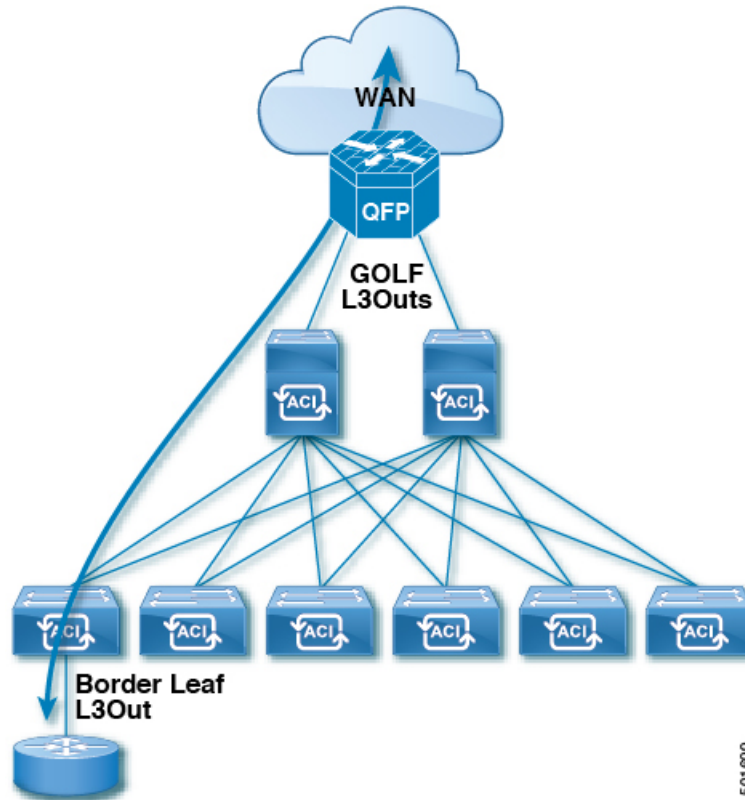
In such scenarios, the policies are administered at the demarcation points and ACI policies need not be imposed.

Layer 4 to Layer 7 route peering is a special use case of the fabric as a transit where the fabric serves as a transit OSPF or BGP domain for multiple pods. You configure route peering to enable OSPF or BGP peering on the Layer 4 to Layer 7 service device so that it can exchange routes with the leaf node to which it is connected. A common use case for route peering is Route Health Injection where the SLB VIP is advertised over OSPF or iBGP to clients outside the fabric. See *L4-L7 Route Peering with Transit Fabric - Configuration Walkthrough* for a configuration walk-through of this scenario.

GOLF in a Transit-Routed Configuration

In APIC, release 2.0 and later, the Cisco ACI supports transit routing with GOLF L3Outs (with BGP and OSPF). For example, the following diagram shows traffic transiting the fabric with GOLF L3Outs and a border leaf L3Out.

Figure 5: GOLF L3Outs and a Border Leaf L3Out in a Transit-Routed Configuration



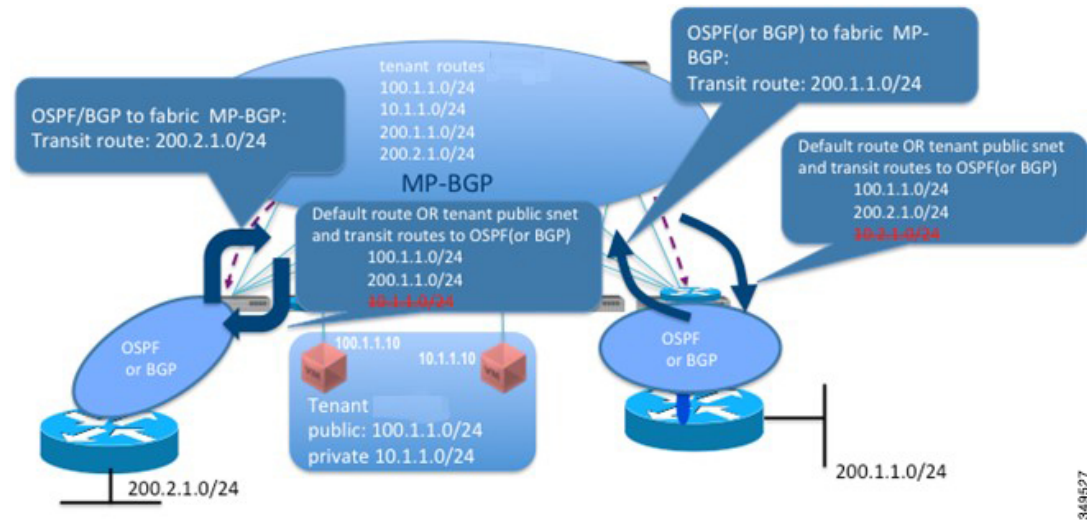
ACI Fabric Route Peering

Layer 3 connectivity and peering with the fabric is configured using a Layer 3 external outside network (`l3extOut`) interface. The peering protocol configuration along with route redistribution and inbound/outbound-filtering rules is associated with an `l3extOut`. The ACI fabric does not appear as a giant router to the external peers, but rather as a transit between separate Layer 3 domains. The peering considerations on one `l3extOut` need not affect the peering considerations on other `l3extOut` policies. The ACI fabric uses MP-BGP for distributing external routes inside the fabric.

Route Redistribution

Inbound routes from external peers are redistributed into the ACI fabric using MP-BGP, subject to inbound filtering rules. These can be transit routes or external routes in the case of WAN connectivity. MP-BGP distributes the routes to all the leaves (including other border leaves) where the tenant is deployed.

Figure 6: Route Redistribution



Inbound route filtering rules select a subset of routes advertised by the external peers to the fabric on the `l3extOut` interfaces. The import filter route-map is generated by using the prefixes in the prefix based EPG. The import filter list is associated only with MP-BGP to restrict the prefixes distributed into the fabric. Set actions can also be associated with import route-maps.

In the outbound direction, an administrator has the option to advertise default routes or transit routes and bridge domain public subnets. If default route advertisement is not enabled, outbound route filtering selectively advertises routes as configured by the administrator.

Currently, route-maps are created with a prefix-list on a per-tenant basis to indicate the bridge domain public subnets to be advertised to external routers. In addition, a prefix-list has to be created to allow all transit routes to be advertised to an external router. The prefix-list for transit routes are configured by an administrator. The default behavior is to deny all transit route advertisement to an external router.

The following options are available for the route-maps associated with transit routes:

- *Permit-all*: Allow all transit routes to be redistributed and advertised outside.
- *Match prefix-list*: Only a subset of transit routes are redistributed and advertised outside.
- *Match prefix-list and set action*: A set action can be associated with a subset of transit routes to tag routes with a particular attribute.

The bridge domain public subnets and transit route prefixes can be different prefix-lists but combined into a single route-map with different sequence numbers. Transit routes and bridge domain public subnets are not expected to have the same prefixes, so prefix-list matches are mutually exclusive.

Route Peering by Protocol

Route peering can be configured per protocol when combining BGP and OSPF with static routes.

OSPF	BGP
Various host types require OSPF to enable connectivity and provide redundancy. These include mainframes, external pods, and service nodes that use ACI as a layer-3 transit within the fabric and to the WAN. Such external devices peer with the fabric through a nonborder leaf running OSPF. Ideally, the OSPF areas are configured as a Not-So-Stubby Area (NSSA) or totally stub area to enable them to receive a default route, so they do not participate in full area routing. For existing deployments where the administrator prefers not to change routing configurations, a stub area configuration is not mandated. Two fabric leaf switches do not establish OSPF adjacency with each other unless they share the same external SVI interface.	External pods and service nodes can use BGP peering with the fabric. BGP peers are associated with an <code>l3extOut</code> and multiple BGP peers can be configured per <code>l3extOut</code>. BGP peer reachability can be through OSPF, EIGRP, connected interfaces, static routes, or loopback. iBGP or eBGP is used for peering with external routers. BGP route attributes from the external router are preserved since MP-BGP is used for distributing the external routes in the fabric. A configuration that contains a match for both transitive and nontransitive BGP extended communities with the same value is not supported; the APIC rejects this configuration.
OSPF Route Redistribution The default-information originate policy in OSPF generates a default route to an external router. Enabling the policy is recommended when peering with mainframes, external pods, and service nodes. When the default-information originate policy is not enabled, configure redistribute-static and redistribute-BGP in the OSPF domain to advertise static bridge domain (BD) public subnets and transit routes respectively. Associate a route-map with the redistribution policy for outbound filtering. It is recommended to not enable the default-information originate option, when peering with external WAN routers. In the inbound direction, OSPF routes are redistributed into the ACI fabric using MP-BGP.	BGP Route Redistribution In the outbound direction, a default route is generated by BGP on a per-peer basis by the default-originate policy. The default route is injected to the peer by BGP even in the absence of a default route in the local routing table. If a default-originate policy is not configured, then static redistribution is enabled for bridge domain public subnets. Transit routes from MP-BGP are available to BGP for advertising. These routes are conditionally advertised outside, subject to outbound filtering policies. In the inbound direction, the advertised routes are available to MP-BGP for redistribution in the fabric, subject to inbound filtering rules. If BGP is used for external peering, then all the BGP attributes of the route are preserved across the fabric.

OSPF	BGP
OSPF Route filtering You can configure OSPF to limit the number of Link-State Advertisements (LSAs) accepted from an external peer to avoid over consumption of the route table, caused by a rogue external router. Inbound route filtering is supported for Layer 3 external outside tenant networks using OSPF. It is applied using a route-map in the in-direction, to filter transit routes allowed in the fabric. In the outbound direction, configure redistribute-static and redistribute-BGP at the OSPF domain level. Configure a route-map to filter the bridge domain public subnets and transit routes. Optionally, some prefixes in the route-map can also be configured with a set action to add route tags. Inter-area prefixes are also filtered using the outbound filter list and associating it with an OSPF area.	BGP Route filtering Inbound route filtering in BGP is applied using a route-map on a per-peer basis. A route-map is configured at the <i>peer-af</i> level in the in-direction, to filter the transit routes to be allowed in the fabric. In the outbound direction, static routes are redistributed into BGP at the <i>dom-af</i> level. Transit routes from MP-BGP are available to the external BGP peering sessions. A route-map is configured at the <i>peer-af</i> level in the out-direction to allow only public subnets and selected transit routes outside. Optionally, a set action to advertise a community value for selected prefixes is configured on the route-map. The bridge domain public subnets and transit route prefixes can be different prefix-lists but combined into a single route-map at the <i>peer-af</i> level with different sequence numbers.
OSPF Name Lookup, Prefix Suppression, and Type 7 Translation OSPF can be configured to enable name lookup for router IDs and suppress prefixes. The APIC system performs OSPF Forwarding Address Suppression in the Translated Type-5 LSAs feature, which causes an NSSA ABR to translate Type-7 LSAs to Type-5 LSAs. To avoid this, use the 0.0.0.0 subnet as the forwarding address instead of the one that is specified in the Type-7 LSA. This feature causes routers that are configured not to advertise forwarding addresses into the backbone to direct forwarded traffic to the translating NSSA ASBRs.	BGP Dynamic Neighbor Support and Private AS Control Instead of providing a specific neighbor address, a dynamic neighbor range of addresses can be provided. Private Autonomous System (AS) numbers can be in the range from 64512 to 65535. They cannot be advertised to a global BGP table. Private AS numbers can be removed from the AS path on a per peer basis and can only be used for eBGP peers according to the following options: • <code>Remove Private AS</code> – Remove if the AS path only has private AS numbers. • <code>Remove All</code> – Remove if the AS path has both private and public AS numbers. • <code>Replace AS</code> – Replace the private AS with the local AS number. Note <code>Remove all</code> and <code>replace AS</code> can only be set if <code>remove private AS</code> is set.

BGP dampening minimizes propagation into the fabric of flapping e-BGP routes that were received from external routers that are connected to border leaf switches (BLs). Frequently flapping routes from external routers are suppressed on BLs based on criteria you configure. They are then prohibited from redistribution to iBGP peers (ACI spine switches). Suppressed routes are reused after a configured time. Each flap penalizes the e-BGP route with a penalty of 1000. When the flap penalty reaches a defined suppress-limit threshold

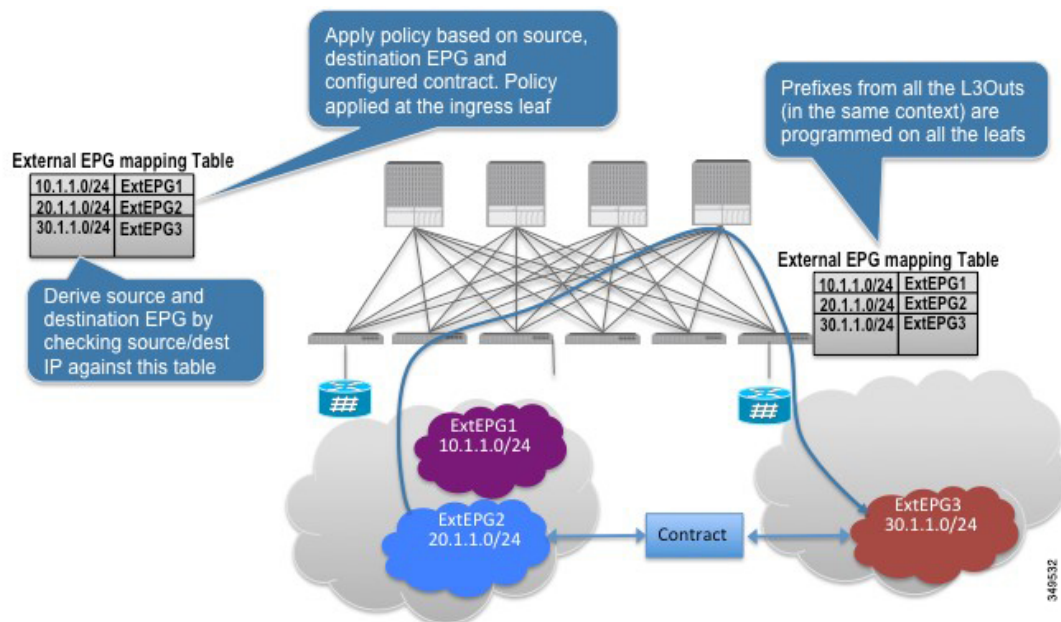
(default 2000), the e-BGP route is marked as dampened. Dampened routes are not advertised to other BGP peers. The penalty is decremented to half after every half-life interval (the default is 15 minutes). A dampened route is reused if the penalty falls below a specified reuse-limit (the default is 750). A dampened route is suppressed at most for a specified maximum suppress time (maximum of 45 minutes).

Use the BGP weight attribute to select a best path. The weight (from 0 to 65,535) is assigned locally to a specific router. The value is not propagated or carried through any of the route updates. By default, paths that the router originates have a weight of 32,768, and other paths have a weight of 0. Routes with a higher weight value have preference when there are multiple routes to the same destination. Set the weight under the BGP neighbor or under the route map.

BGP peering is typically configured to the neighbor's loopback address. In such cases, loopback reachability is statically configured or advertised (more commonly) through OSPF. The loopback interface is configured as a passive interface and added into the OSPF area. There are no redistribution policies that are attached to OSPF. The route redistribution implementation is through BGP. Route filtering can be configured in L3Outs for tenant networks that use either BGP or OSPF.

External routes can also be programmed as static routes on the border leaf in the respective tenants. A peering protocol is not required if the external routes are programmed as static routes on the border leaf. External static routes are redistributed to other leaf switches in the fabric through MP-BGP, subject to import filtering. Starting with release 1.2(1x), static route preference incoming within the ACI fabric is carried in MP-BGP using a cost extended community. On an L3Out connection, an MP-BGP route coming from Layer 4 wins over a local static route. A route is installed in the Unicast Routing Information Base (URIB) with the preference that is specified by an administrator. On an ACI nonborder leaf switch, a route is installed with Layer 4 as nexthop. When nexthop on Layer 4 is not available, the Layer 3 static route becomes the best route in fabric.

Figure 7: Static Route Policy Model for Transit



For `l3extOut` connections, external endpoints are mapped to an external EPG based on IP prefixes. For each `l3extOut` connection, you can create one or more external EPGs, based on whether different endpoints require different policies.

Each external EPG is associated with a class-id. Each prefix in the external EPG is programmed in the hardware to derive the corresponding class-id. The prefixes are only qualified VRF instance and not by the `l3extOut` interface on which they are deployed.

The union of prefixes from all the `l3extOut` policies in the same VRF is programmed on all the leaf switches where the `l3extOut` policies are deployed. The source and destination class-ids corresponding to the source and destination IP address in a packet are derived at the ingress leaf and the policy is applied on the ingress leaf itself, based on the configured contract. If a contract allows traffic between two prefixes on two L3Out interfaces, then packets with any combination of the source and destination IP address (belonging to the configured prefixes) is allowed between the L3Out interfaces. If there is no contract between the EPGs, the traffic is dropped at the ingress leaf.

Since prefixes are programmed on every leaf switch where `l3extOut` policies are deployed, the total number of prefixes APIC supports for prefix-based EPGs is limited to 1000 for the fabric.

Overlapping or equal subnets cannot be configured on different `l3extOut` interfaces in the same VRF. If overlapping or equal subnets are required, then a single `l3extOut` is used for transit with appropriate export prefixes.

Transit Route Control

A route transit is defined to import traffic through a Layer 3 outside network `L3extOut` profile (`l3extInstP`), where it is to be imported. A different route transit is defined to export traffic through another `l3extInstP` where it is to be exported.

Since multiple `l3extOut` policies can be deployed on a single node or multiple nodes in the fabric, a variety of protocol combinations are supported. Every protocol combination can be deployed on a single node using multiple `l3extOut` policies or multiple nodes using multiple `l3extOut` policies. Deployments of more than two protocols in different `l3extOut` policies in the fabric are supported.

Export route-maps are made up of prefix-list matches. Each prefix-list consists of bridge domain (BD) public subnet prefixes in the VRF and the export prefixes that need to be advertised outside.

Route control policies are defined in an `l3extOut` policy and controlled by properties and relations associated with the `l3extOut`. APIC uses the `enforceRtctrl` property of the `l3extOut` to enforce route control directions. The default is to enforce control on export and allow all on import. Imported and exported routes (`l3extSubnets`), are defined in the `l3extInstP`. The default scope for every route is import. These are the routes and prefixes which form a prefix-based EPG.

All the import routes form the import route map and are used by BGP and OSPF to control import. All the export routes form the export route map used by OSPF and BGP to control export.

Import and export route control policies are defined at different levels. All IPv4 policy levels are supported for IPv6. Extra relations that are defined in the `l3extInstP` and `l3extSubnet` MOs control import.

Default route leak is enabled by defining the `l3extDefaultRouteLeakP` MO under the `l3extOut`.

`l3extDefaultRouteLeakP` can have Virtual Routing and Forwarding (VRF) scope or `L3extOut` scope per area for OSPF and per peer for BGP.

The following set rules provide route control:

- `rtctrlSetPref`
- `rtctrlSetRtMetric`
- `rtctrlSetRtMetricType`

Additional syntax for the `rtctrlSetComm` MO includes the following:

- `no-advertise`
- `no-export`
- `no-peer`

BGP

The ACI fabric supports BGP peering with external routers. BGP peers are associated with an `l3extOut` policy and multiple BGP peers can be configured per `l3extOut`. BGP can be enabled at the `l3extOut` level by defining the `bgpExtP` MO under an `l3extOut`.



Note Although the `l3extOut` policy contains the routing protocol (for example, BGP with its related VRF), the `L3Out` interface profile contains the necessary BGP interface configuration details. Both are needed to enable BGP.

BGP peer reachability can be through OSPF, EIGRP, a connected interface, static routes, or a loopback. iBGP or eBGP can be used for peering with external routers. The BGP route attributes from the external router are preserved since MP-BGP is used for distributing the external routes in the fabric. BGP enables IPv4 and/or IPv6 address families for the VRF associated with an `l3extOut`. The address family to enable on a switch is determined by the IP address type defined in `bgpPeerP` policies for the `l3extOut`. The policy is optional; if not defined, the default will be used. Policies can be defined for a tenant and used by a VRF that is referenced by name.

You must define at least one peer policy to enable the protocol on each border leaf (BL) switch. A peer policy can be defined in two places:

- Under `l3extRsPathL3OutAtt`—a physical interface is used as the source interface.
- Under `l3extLNodeP`—a loopback interface is used as the source interface.

OSPF

Various host types require OSPF to enable connectivity and provide redundancy. These include mainframe devices, external pods and service nodes that use the ACI fabric as a Layer 3 transit within the fabric and to the WAN. Such external devices peer with the fabric through a nonborder leaf switch running OSPF. Configure the OSPF area as an NSSA (stub) area to enable it to receive a default route and not participate in full-area routing. Typically, existing routing deployments avoid configuration changes, so a stub area configuration is not mandated.

You enable OSPF by configuring an `ospfExtP` managed object under an `l3extOut`. OSPF IP address family versions configured on the BL switch are determined by the address family that is configured in the OSPF interface IP address.



Note Although the `l3extOut` policy contains the routing protocol (for example, OSPF with its related VRF and area ID), the Layer 3 external interface profile contains the necessary OSPF interface details. Both are needed to enable OSPF.

You configure OSPF policies at the VRF level by using the `fvRsCtxToOspfCtxPol` relation, which you can configure per address family. If you do not configure it, default parameters are used.

You configure the OSPF area in the `ospfExtP` managed object, which also exposes IPv6 the required area properties.

Default Policy Behavior

When there are no contracts between two prefix-based EPGs, traffic between unknown-source and unknown-destination prefixes are dropped. These drops are achieved by implicitly programming different class-ids for unknown source and destination prefixes. Since the class-ids are different, they are impacted by the class-unequal rule and packets are denied. The class-unequal drop rule also causes packets to be dropped from known source and destination IP addresses to unknown source and destination IP addresses and vice versa.

Due to this change in the default behavior, the class-id programming for catch-all (0/0) entries has been changed as illustrated in the example below:

- Unknown source IP address is EPG1.
- Unknown destination IP address is EPG2.
- Unknown source IP $\longleftrightarrow$ Unknown destination IP $\Rightarrow$ class-unequal rule $\Rightarrow$ DROP.
- User-configured default prefixes (0/0) = EPG3 and (10/8) = EPG4. Contract between EPG3 and EPG4 is set to ALLOW.
- Programmed rules:
 - EPG1 $\longleftrightarrow$ EPG4 $\Rightarrow$ class-unequal rule $\Rightarrow$ DROP
 - EPG4 $\longleftrightarrow$ EPG2 $\Rightarrow$ class-unequal rule $\Rightarrow$ DROP

EIGRP Protocol Support

EIGRP protocol is modeled similar to other routing protocols in the Cisco Application Centric Infrastructure (ACI) fabric.

Supported Features

The following features are supported:

- IPv4 and IPv6 routing
- Virtual routing and forwarding (VRF) and interface controls for each address family
- Redistribution with OSPF across nodes
- Default route leak policy per VRF
- Passive interface and split horizon support
- Route map control for setting tag for exported routes
- Bandwidth and delay configuration options in an EIGRP interface policy
- Authentication support

Unsupported Features

The following features are not supported:

- Stub routing
- EIGRP used for BGP connectivity
- Multiple EIGRP `L3extOut`s on the same node
- Per-interface summarization (an EIGRP summary policy will apply to all interfaces configured under an `L3Out`)
- Per interface distribute lists for import and export

Categories of EIGRP Functions

EIGRP functions can be broadly categorized as follows:

- Protocol policies
- `L3extOut` configurations
- Interface configurations
- Route map support
- Default route support
- Transit support

Primary Managed Objects That Support EIGRP

The following primary managed objects provide EIGRP support:

- **EIGRP Address Family Context Policy** `eigrpCtxAfPol`: Address Family Context policy configured under `fvTenant` (Tenant/Protocols).
- `fvRsCtxToEigrpCtxAfPol`: Relation from a VRF to a `eigrpCtxAfPol` for a given address family (IPv4 or IPv6). There can be only one relation for each address family.
- `eigrpIfPol`: EIGRP Interface policy configured in `fvTenant`.
- `eigrpExtP`: Enable flag for EIGRP in an `L3extOut`.
- `eigrpIfP`: EIGRP interface profile attached to an `l3extLIfP`.
- `eigrpRsIfPol`: Relation from EIGRP interface profile to an `eigrpIfPol`.
- `Defrtleak`: Default route leak policy under an `l3extOut`.

EIGRP Protocol Policies Supported Under a Tenant

The following EIGRP protocol policies are supported under a tenant:

- **EIGRP Interface policy** (`eigrpIfPol`)—contains the configuration that is applied for a given address family on an interface. The following configurations are allowed in the interface policy:
 - *Hello interval* in seconds

- *Hold interval* in seconds
- One or more of the following interface control flags:
 - *split horizon*
 - *passive*
 - *next hop self*
- **EIGRP Address Family Context Policy (`eigrpCtxAfPol`)**—contains the configuration for a given address family in a given VRF. An `eigrpCtxAfPol` is configured under tenant protocol policies and can be applied to one or more VRFs under the tenant. An `eigrpCtxAfPol` can be enabled on a VRF through a relation in the VRF-per-address family. If there is no relation to a given address family, or the specified `eigrpCtxAfPol` in the relation does not exist, then the default VRF policy created under the `common` tenant is used for that address family.

The following configurations are allowed in the `eigrpCtxAfPol`:

- Administrative distance for internal route
- Administrative distance for external route
- Maximum ECMP paths allowed
- Active timer interval
- Metric version (32-bit / 64-bit metrics)

EIGRP L3extOut Configuration

EIGRP is the main protocol used for advertising the fabric public subnets, connect routes, static routes, and transit routes configured on the leaf switches.

There is an enable/disable flag for EIGRP for a given Layer 3 external outside network (`l3extOut`) routed domain.



Note The autonomous system number that is a tag used for EIGRP and is not the same as the fabric ASN used by BGP.

EIGRP cannot be enabled with BGP and OSPF on the same `L3extOut`.

The following EIGRP transit scenarios are supported:

- EIGRP running in an `L3extOut` on one node and OSPF running in another `L3extOut` on a different node.



Note Multiple EIGRP `L3extOuts` are not supported on the same node in the same Virtual Routing and Forwarding (VRF).

- EIGRP to static route transit.

EIGRP Interface Profile

To enable EIGRP on an interface, an EIGRP profile needs to be configured under the interface profile in the `L3extOut->Node->Interface` hierarchy. An EIGRP profile has a relation to an EIGRP interface policy enabled in the tenant. In the absence of a relation or interface policy in the tenant, the default EIGRP interface policy in the `common` tenant is used. EIGRP is enabled on all interfaces contained in the interface profile. This includes L3-Port, Sub-interfaces, External SVI on ports, port-channels, and VPCs contained in the interface profile.

Route-map infrastructure and settings in the policy model are common across all the protocols. The route-map set actions are a superset of actions to cover BGP, OSPF, and EIGRP. The EIGRP protocol supports the *set tag* option in route maps used for interleaf/redistribution. These route maps are configured on a per-VRF basis. If the `L3extOut` has both IPv4 and IPv6 interfaces, then an interleaf policy is applied on both IPv4 and IPv6 address families for that VRF.



Note At this time, VRF-level route maps are supported, but interface route maps are not supported.

The default route leak policy on the `L3extOut` is protocol agnostic in terms of the configuration. Properties enabled in the default route leak policy are a superset of the individual protocols. Supported configurations in the default route leak are as follows:

- **Scope:** VRF is the only scope supported for EIGRP.
- **Always:** The switch advertises the default route only if present in the routing table or advertises it regardless.
- **Criteria:** only or in-addition. With the only option, only the default route is advertised by EIGRP. With in-addition, the public subnets and transit routes are advertised along with the default route.

The default route leak policy is enabled in the domain per VRF per address family.

By default, the protocol redistribution interleaf policies with appropriate route maps are set up for all valid configurations. The administrator enables transit routing purely by virtue of creating `L3extInstP` subnets with *scope=export-route control* to allow certain routes to be transmitted between two `L3extOuts` in the same VRF. Apart from the scope of `L3extInstP` subnets, there are no special protocol specific configurations for covering transit cases. Apart from the scope, which is protocol-specific, other parameters of the default route leak policy are common across all the protocols.

The OSPF on another `L3extOut` on a different node transit scenario is supported with EIGRP.

Observe the following EIGRP guidelines and limitations:

- At this time, multiple EIGRP L3Outs are not supported on the same leaf switch.
- All routes are imported on an `L3extOut` that uses EIGRP. Import subnet scope is disabled in the GUI if EIGRP is the protocol on the `L3extOut`.