

AWS Whitepaper

Hybrid Connectivity



Hybrid Connectivity: AWS Whitepaper

Copyright © 2025 Amazon Web Services, Inc. and/or its affiliates. All rights reserved.

Amazon's trademarks and trade dress may not be used in connection with any product or service that is not Amazon's, in any manner that is likely to cause confusion among customers, or in any manner that disparages or discredits Amazon. All other trademarks not owned by Amazon are the property of their respective owners, who may or may not be affiliated with, connected to, or sponsored by Amazon.

Table of Contents

.....	iv
Abstract and introduction	i
Introduction	1
Are you Well-Architected?	2
AWS hybrid connectivity building blocks	3
Hybrid network connections	3
AWS Direct Connect	3
Site-to-Site VPN	5
Transit Gateway Connect	5
AWS hybrid connectivity services	5
Hybrid connectivity type and design considerations	7
Connectivity type selection	8
Time to deploy	8
Security	10
Service level agreement	12
Performance	14
Cost	16
Connectivity design selection	19
Scalability	20
Connectivity models	21
Reliability	34
Customer-managed VPN and SD-WAN	42
Example Corp. Automotive use case	44
Architecture selected	51
Conclusion	53
Contributors	54
Further reading	55
Document revisions	56
Notices	57
AWS Glossary	58

This whitepaper is for historical reference only. Some content might be outdated and some links might not be available.

Hybrid Connectivity

Publication date: **July 6, 2023** ([Document revisions](#))

Many organizations need to connect their on-premises data centers, remote sites, and the cloud. A hybrid network connects these different environments. This whitepaper describes AWS building blocks and the key requirements to consider when deciding which hybrid connectivity model is right for you. To help you determine the best solution for your business and technical requirements, we provide decision trees to guide you through the logical selection process.

Introduction

A modern organization uses an extensive array of IT resources. In the past, it was common to host these resources in an on-premises data center or a colocation facility. With the increased adoption of cloud computing, organizations deliver and consume IT resources from cloud service providers over a network connection. Organizations can opt to migrate some, or all, of their existing IT resources to the cloud. In either case, a common network is required to connect on-premises and cloud resources. Coexistence of on-premises and cloud resources is called *hybrid cloud*, and the common network connecting them is referred to as a *hybrid network*. Even if your organization keeps all its IT resources in the cloud, it may still require hybrid connectivity to remote sites.

There are several connectivity models to choose from. While having options adds flexibility, selecting the optimal option requires analysis of the business and technical requirements, and elimination of options that are not suitable. You can group requirements together across considerations such as security, time to deploy, performance, reliability, communication model, scalability, and more. Once they have carefully collected, analyzed, and considered the requirements, network and cloud architects can identify the applicable AWS hybrid network building blocks and solutions. To identify and select the optimal model or models, architects must understand advantages and disadvantages of each model. There are also technical limitations that might cause an otherwise suitable model to be excluded.

To simplify the selection process, this whitepaper guides you through each key consideration in a logical order. Under each consideration, there are questions used to collect requirements. Each design decision's impact is identified, along with potential solutions. The whitepaper presents decision trees for some of the considerations as a method to aid decision-making process, eliminate options, and understand consequences of each decision. It concludes with a scenario covering a hybrid use case, applying the end-to-end connectivity model selection and design. You

can use this example to see how to execute the processes laid out in this whitepaper in a practical example.

This whitepaper is intended to help you select and design an optimal hybrid connectivity model. This whitepaper is structured as follows:

- **Hybrid connectivity building blocks** – An overview of AWS services used for hybrid connectivity.
- **Connectivity selection and design considerations** – A definition of each connectivity model, how each affects the design decision, requirement identification questions, solutions, and decision trees.
- **A customer use case** - An example of how to apply the considerations and decision trees in practice.

Are you Well-Architected?

The [AWS Well-Architected Framework](#) helps you understand the pros and cons of the decisions you make when building systems in the cloud. The six pillars of the Framework allow you to learn architectural best practices for designing and operating reliable, secure, efficient, cost-effective, and sustainable systems. Using the [AWS Well-Architected Tool](#), available at no charge in the [AWS Management Console](#), you can review your workloads against these best practices by answering a set of questions for each pillar.

For more expert guidance and best practices for your cloud architecture—reference architecture deployments, diagrams, and whitepapers—refer to the [AWS Architecture Center](#).

AWS hybrid connectivity building blocks

There are three building blocks of a hybrid network connectivity architecture:

- **Hybrid network connections:** The connection types between AWS connectivity services and on-premises customer gateway devices.
- **AWS hybrid connectivity services:** The AWS services that provide connectivity and routing between the customer infrastructure and AWS.
- **On-premises customer gateway device:** The device inside the customer's existing network that is the on-premises endpoint for hybrid network connection. Different connection types have different technical requirements for these devices, which are discussed in the following sections.

Hybrid network connections

There are several ways to connect between your on-premises equipment and AWS. This whitepaper is focused on how these different ways can be combined into overall architectures, however, a brief overview of the different options (AWS Direct Connect, Site-to-Site Virtual Private Network, and Transit Gateway Connect) are provided.

AWS Direct Connect

AWS Direct Connect is a service that establishes a dedicated network connection from your premises to AWS. See [AWS Direct Connect](#) for details.

There are two types of AWS Direct Connect connections: dedicated and hosted. A dedicated connection is a direct link between an AWS device and your on-premises device, whereas a hosted connection is supported by an AWS Partner who can handle connection details for you. See [AWS Direct Connect connections](#) for more information.

A Direct Connect connection uses Virtual Interfaces (VIFs) to isolate different traffic flows. Multiple VIFs can use the same Direct Connect link, separated by VLAN (802.1q) tags. There are three types of VIFs that provide connectivity to the AWS network. See [AWS Direct Connect virtual interfaces](#) for more details. The three types are:

- **Private VIF:** A private VIF is a private connection between your device and your resources inside AWS. These terminate inside AWS on either a Virtual Private Gateway (VGW) directly (that supports a single VPC) or via a Direct Connect Gateway that then connects to multiple VGWs.

- **Public VIF:** A public VIF enables connectivity to any public AWS resources, such as S3, DynamoDB, and public EC2 IP ranges. While a public VIF does not have direct access to the internet, any Amazon public resource can reach it (including other customers' public EC2 instances), which customers should consider during security planning.
- **Transit VIF:** A transit VIF is a private connection between your device and an AWS Transit Gateway, via a Direct Connect Gateway. Transit VIFs are now supported on links with speeds of less than 1 Gbps - see [the launch announcement](#) for details.

Note

Hosted Virtual Interface (Hosted VIF) is a type of Private VIF where the VIF is assigned to a different AWS account than the AWS account which owns the AWS Direct Connect connection (which can include an AWS Direct Connect partner). AWS no longer allows new partners to offer this model. For more information, see [Creating a hosted virtual interface](#).

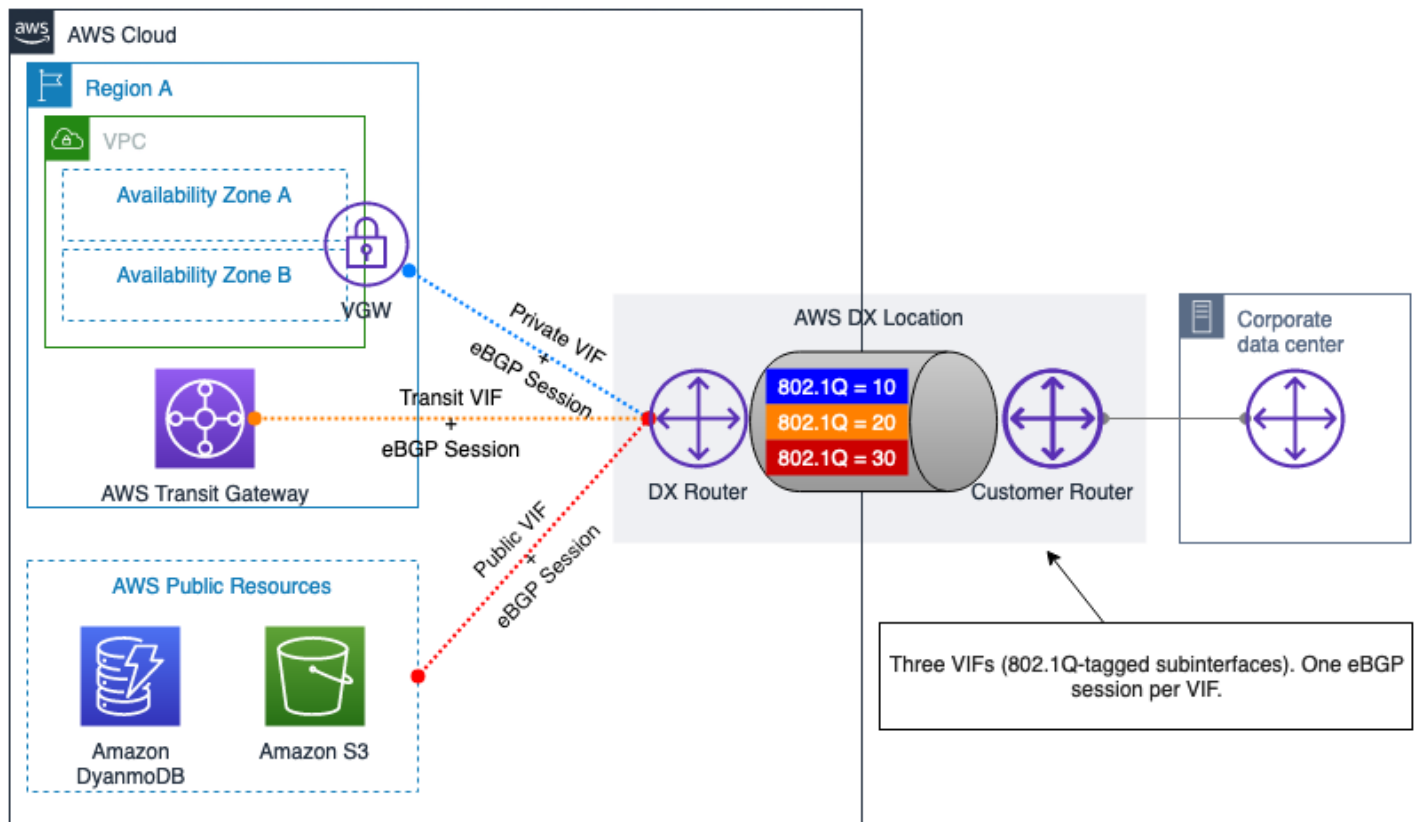


Figure 1 – AWS Direct Connect Private and Public VIFs

Site-to-Site Virtual Private Network (VPN)

A site-to-site VPN enables two networks to communicate securely and can be used over an untrusted transport, such as the internet. Customers can establish VPN connections between on-premises sites and Amazon Virtual Private Clouds (Amazon VPC) via two options:

- **AWS Managed Site-to-Site VPN (AWS S2S VPN):** This is a fully managed and highly available VPN service, using IPsec. See [What is AWS Site-to-Site VPN](#) for more information. You can optionally enable acceleration for your Site-to-Site VPN connection. See [Accelerated Site-to-Site VPN connections](#) for more information. S2S VPN can also use Direct Connect transit VIFs to avoid having the traffic traverse the internet, lowering costs and allowing the use of private IP addresses. For details, see [Private IP VPN with AWS Direct Connect](#).
- **Software Site-to-Site VPN (Customer-managed VPN):** With this VPN connectivity option, you are responsible for provisioning and managing the entire VPN solution, typically by running VPN software on an EC2 instance. For more information, see [Software Site-to-Site VPN](#).

Both options require support on the customer gateway device to terminate the on-premises end of the VPN tunnels. This device can be a physical device or a software appliance. For more information about network devices tested by AWS, refer to the list of [tested customer gateway devices](#).

Transit Gateway Connect (TGW Connect)

Transit Gateway Connect uses GRE tunnels between an AWS Transit Gateway and an on-premises gateway device. BGP is used on top of TGW Connect to enable dynamic routing. Note that TGW Connect is not encrypted. For more information, see [Transit Gateway Connect](#).

AWS hybrid connectivity services

AWS hybrid connectivity services provide highly scalable, highly available networking components. They play an essential role in building hybrid networking solutions. At the time of this whitepaper writing, there are three primary service endpoints:

- **AWS Virtual Private Gateway (VGW)** is a regional, highly redundant service that provides IP routing and forwarding at the VPC level, acting as the gateway for the VPC to communicate with your customer gateway devices. VGW can terminate AWS S2S VPN connections and AWS Direct Connect Private VIFs.

- **AWS Transit Gateway (TGW)** is a regional, highly available and scalable service that enables you to connect multiple VPCs with each other, as well as your on-premises networks over Site-to-Site VPN and/or Direct Connect using a single centralized gateway. Conceptually, an AWS Transit Gateway acts as a highly available and redundant virtual cloud router. AWS Transit Gateway supports equal cost multi-path (ECMP) routing over multiple Direct Connect connections, VPN tunnels, or TGW Connect peers. Transit Gateways can peer to each other, both in the same region and cross-region, allowing their connected resources to communicate over the peering links. For more details, see [AWS Transit Gateway scenarios](#).
- **AWS Cloud WAN** provides a central dashboard for making connections between your branch offices, data centers, and Amazon VPCs—building a global network with only a few clicks. You use network policies to automate network management and security tasks in one location. For more details, see the [AWS Cloud WAN documentation](#).
- **Direct Connect Gateway (DXGW)** is a globally available service that distributes routing information across its connections, behaving similarly to BGP route reflectors in a traditional network. Data does not pass through a DXGW – it only handles the routing information. You can create a DXGW in any AWS Region and access it from all other AWS Regions. You can connect Direct Connect VIFs to a DXGW, then associate the DXGW with either VGWs (using private VIFs) or an AWS Transit Gateway (using transit VIFs). See [Direct Connect gateways](#) for more information. You do not need to create multiple DXGWs for redundancy as it is a globally availability service. However, you might choose to use multiple DXGWs to separate routing domains, for example, a production and a testing network you want to keep completely isolated.

Hybrid connectivity type and design considerations

This section of the whitepaper covers the considerations that affect your choices when selecting a hybrid network to connect your on-premises environments to AWS. It follows a logical thought process to support you selecting an optimal hybrid connectivity solution. The considerations affecting your design are categorized into considerations that impact your *connectivity type*, and considerations that affect your *connectivity design*. Connectivity type considerations will support you deciding between using an internet-based VPN or Direct Connect. Connectivity design considerations will support you deciding how to set up the connections.

The following considerations that impact your *connectivity type* are covered: time to deploy, security, SLA, performance, and cost. After reviewing those considerations, and how they affect your design choices, you will be able to decide if using an internet-based connection or Direct Connect is recommended to meet your requirements.

The following considerations that impact your *connectivity design* are covered: scalability, communication model, reliability, and third-party SD-WAN integration. After reviewing those considerations, and how they affect your design choices, you will be able to decide the optimal logical design recommended to meet your requirements.

The following structure is used to discuss and analyze each of the selection and design considerations:

- **Definition** - Brief definition of what is the consideration.
- **Key questions** - Provides a set of questions to enable you to collect the requirements associated with the consideration.
- **Capabilities to consider** - Solutions to address the requirements associated with the consideration.
- **Decision tree** - For some considerations or a group of considerations, a decision tree is provided to help you select the optimal hybrid network solution.

The considerations affecting your hybrid network design are covered in an order where the output of one consideration is part of the input for the subsequent consideration. As illustrated in Figure 2, the first step is to decide on the connectivity type, followed by refining it with the design selection considerations.

Figure 2 demonstrates the two consideration categories, the individual considerations, and the logical order in which the considerations are covered in the subsequent sub-sections. Those are the essential considerations when making a hybrid network design decision. If the targeted design does not require all these considerations, you can focus on the considerations that apply to your requirements.

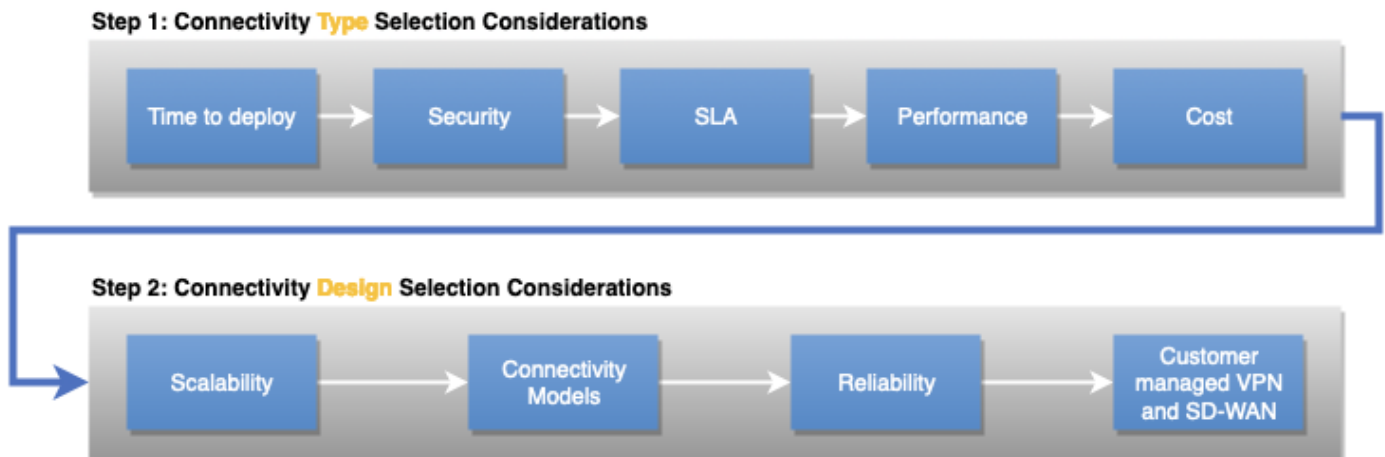


Figure 2 – Consideration categories, individual considerations, and the logical order between them

Connectivity type selection

This section covers considerations that affect the connectivity type you select for your workload. This includes time to deploy, security, SLA, performance, and cost.

Considerations

- [Time to deploy](#)
- [Security](#)
- [Service level agreement \(SLA\)](#)
- [Performance](#)
- [Cost](#)

Time to deploy

Definition

Time to deploy can be an important factor in selecting a suitable connectivity type for a workload. Depending on the type of connectivity and on-premises locations, connectivity can be established

within hours, however, it may take weeks or months if additional circuits must be installed. This will influence your decision to use an internet-based connection, a private dedicated connection, or a private hosted connection provided as a managed service by an AWS Direct Connect Partner.

Key questions

- What is the required timeline for the deployment – hours, days, weeks, or months?
- How long will the connection be needed – will it be a short-lived project or permanent infrastructure?

Capabilities to consider

When you require AWS connectivity within hours or days, you will most likely need to use an existing network connection. This often means establishing a VPN connection to AWS over the public internet. If an existing AWS DX partner is providing you with private AWS connectivity, a new hosted connection could be provisioned within hours.

When you have days to weeks, you can work with an AWS Direct Connect Partner to establish private connectivity to AWS. AWS Direct Connect Partners help you establish network connectivity between AWS Direct Connect locations and your data center, office, or co-location environment. Certain [AWS Direct Connect Partners](#) are approved to offer [Direct Connect Hosted Connections](#). Hosted Connections can often be provisioned faster than Dedicated Connections. AWS Direct Connect Partner will provision each Hosted Connection using their existing infrastructure that is connected to the AWS backbone.

When you have several weeks to months, you can investigate establishing a dedicated private connection with AWS. Service providers and AWS Direct Connect Partners facilitate AWS Direct Connect Dedicated Connections. It's common for service providers to install networking equipment at the customer's premises to facilitate a Direct Connect Dedicated Connection. Depending on the service provider, location of your site, and other physical factors, the installation of a Direct Connect Dedicated Connection can take from several weeks to a few months.

If you already have your network equipment installed in the same colocation facility where the AWS Direct Connect location exists, then you can quickly establish an AWS Direct Connect Dedicated Connection via a cross-connect at the co-location site. After you request the connection, AWS makes a Letter of Authorization and Connecting Facility Assignment (LOA-CFA) available to you to download, or emails you with a request for more information. The LOA-CFA is the

authorization to connect to AWS, and is required by your network provider to order a cross connect for you.

Table 1 – Cost effectiveness comparison

	Internet-based connectivity	DX Dedicated Connection (existing equipment within DX location)	DX Dedicated Connection (net-new)	DX Hosted Connection (existing port with DX Partner)	DX Hosted Connection (net-new)
Provisioning time	Hours to days	Days	Several weeks to months	Hours to days	Several days to weeks to months

Note

The provided provision time guidelines are based on real-world observation and only serve as an illustration. When taking into considerations your site location, proximity to direct connect locations, and pre-existing infrastructure, and will all impact provisioning time. Your AWS Direct Connect Partner will advise you on the precise provisioning time.

Security

Definition

Security requirements will influence your hybrid connectivity type. These considerations include:

- Transport type – internet or private network connection
- Encryption requirements

Key questions

- Do your security requirements and policies allow the use of encrypted connections over the internet to connect to AWS, or do they mandate the use of private network connections?
- When leveraging private network connections, does the network layer have to provide encryption in transit?

Technical solutions

Your security requirements and policies might permit use of internet or require use of a private network connection between AWS and your company network. They also affect the decision if the network must provide encryption in transit, or if performing encryption at application layer is acceptable.

If you can leverage the internet, then AWS Site-to-Site VPN can be used to create encrypted tunnels between your network and your Amazon VPCs or AWS Transit Gateways over the internet. Extending your [SD-WAN](#) solution into AWS over the internet is also an option if you are leveraging an internet-based connection. The section Customer-managed VPN and SD-WAN later in this whitepaper covers the specific considerations for SD-WAN.

If you require a private network connection between AWS and your company network, then AWS recommends using AWS Direct Connect Dedicated Connections or Hosted Connections. If encryption in transit is required over a private network connection, then you should establish a VPN over Direct Connect (either over public VIF or transit VIF), or consider using MACsec on a 10Gbps or 100Gbps Dedicated connection.

Table 2 – Example Automotive Corp connectivity type requirements

	Site-to-Site VPN	Direct Connect
Transport	Internet	Private network connection
Encryption in transit	Yes	Requires S2S VPN over DX, S2S VPN over a transit VIF, or MACsec on a 10Gbps or 100Gbps Dedicated Connection

Service level agreement (SLA)

Definition

Enterprise organizations often require a service provider to fulfil an SLA for each service the organization consumes. The organization in turn builds its own services on top and may offer their own consumers an SLA. The SLA is important as it describes how the service is provided and operated, and it often includes specific measurable characteristics, such as availability. Should the service break the defined SLA, a service provider usually offers financial compensation specified by the agreement. An SLA defines the type of measure, the requirement, and the measurement period. As an example, refer to uptime target definition under the [AWS Direct Connect SLA](#).

Key questions

- Is a hybrid connectivity connection SLA with service credits required?
- Does the entire hybrid network need to adhere to an uptime target?

Capabilities to consider

Connectivity type: Internet connectivity can be unpredictable. While AWS takes great care with multiple links in place with a diverse set of ISPs, the administration of the internet is simply outside of AWS or a single provider's administrative domain. There is a limited amount of route engineering and traffic influence a cloud provider can do once traffic has left the border of their network. That said, there is an [AWS Site-to-Site VPN SLA](#) that provides availability targets for AWS Site-to-Site VPN endpoints.

AWS [Direct Connect offers a formal SLA](#) with service credits calculated as a percentage of the total AWS Direct Connect Port Hour charges paid by you for the applicable connections experiencing unavailability for the monthly billing cycle in which the SLA was not met. This is the recommended transport if an SLA is required. AWS Direct Connect lists [specific minimal configuration requirements](#) for each uptime target such as number of AWS Direct Connect locations, connections, and other configuration details. The failure to satisfy the requirements means that service credits cannot be offered should the service break defined SLAs.

Importantly, even if the service selected to provide hybrid connectivity is configured to meet the SLA requirements, the rest of the network may not provide the same level of SLA. The AWS responsibility ends at the AWS Direct Connect location at the AWS Direct Connect port. Once AWS

hands traffic off to your organization's network, it is no longer the responsibility of AWS. If you use a service provider between AWS and your on-premises network, connectivity is subject to SLA between yourself and the service provider, if applicable. Keep in mind that the entire hybrid network is just as good as the weakest part of it when designing hybrid connectivity.

AWS Direct Connect partners offer AWS Direct Connect connectivity. The partner may offer an SLA with service credits based on their product offering up to the demarcation point with AWS. The option should be evaluated and further researched directly with APN Partners. AWS publishes [a list of validated delivery partners](#).

Logical design: In addition to the connectivity type, you also must consider other building blocks as part of your overall design. As an example, [AWS Transit Gateway](#) has its own SLA, as does [AWS S2S VPN](#). You might be using AWS Transit Gateway for scale and AWS S2S VPN for security reasons, but you must design both in manner consistent with each SLAs to be eligible for service credits with each respective service.

Review [AWS Direct Connect Resiliency Recommendations](#) and [Resiliency Toolkit](#).

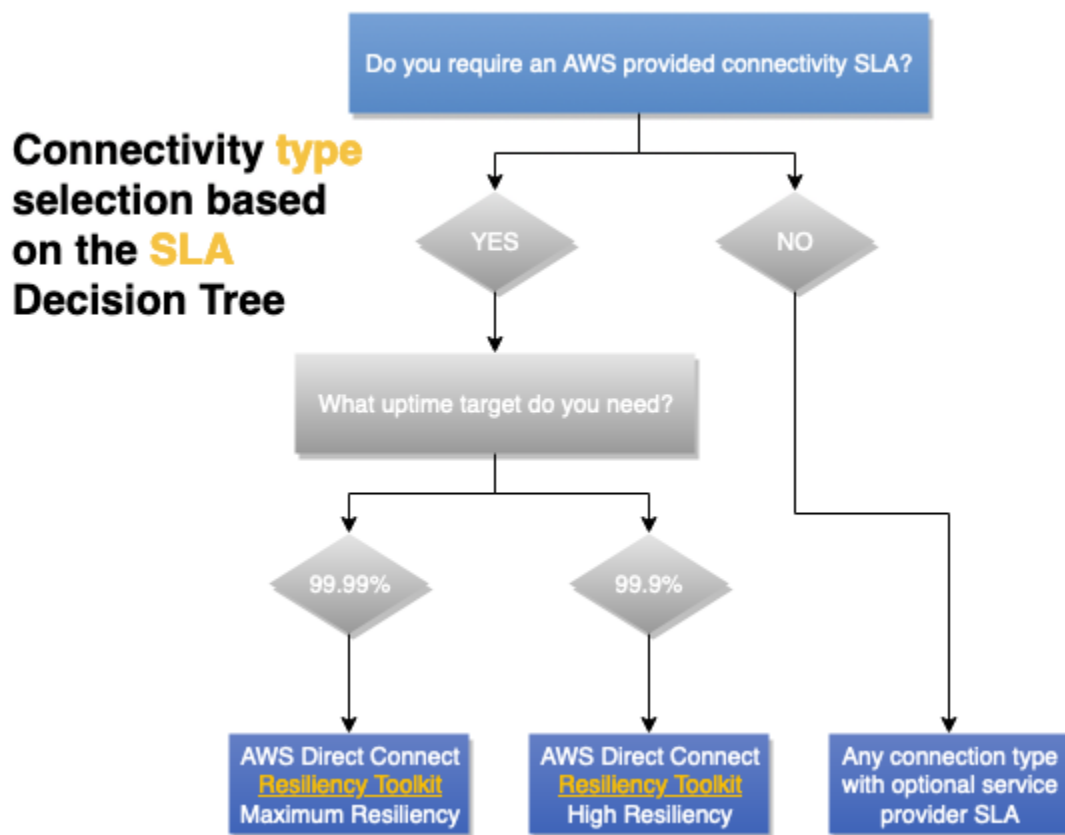


Figure 3 – SLA consideration decision tree

Performance

Definition

There are multiple factors which influence network performance, such as latency, packet loss, jitter, and bandwidth. Depending on application requirements, the importance of each of these factors can vary.

Key questions

Based on your application requirements, you need to identify and prioritize the network performance factors that impact your application behavior and user experience.

Bandwidth

Bandwidth refers to the data transfer rate of a connection, and is usually measured in bits per second (bps). Megabits per second (Mbps) and gigabits per second (Gbps) are common scaling, and are base 10 (1,000,000 bits per second = 1 Mbps) as opposed to base 2 (2^{10}) seen elsewhere.

When evaluating the bandwidth needs of applications, keep in mind that the bandwidth requirements can change over time. Initial deployment into the cloud, normal operations, new workloads, and failover scenarios can all have different bandwidth requirements.

Applications can have their own bandwidth considerations. Some applications might require deterministic performance over a high-bandwidth connection, while others can require both deterministic performance and high bandwidth. An application may need special configuration to use multiple traffic flows (sometimes referred to as streams or sockets) in parallel if it is hitting per traffic flow bandwidth limits, allowing it to use more of the connection's bandwidth. VPNs can limit throughput because of tunneling overheads, lower MTU limits, or hardware bandwidth limitations.

Latency

Latency is the time needed for a packet to go from source to destination over a network connection, and is usually measured in milliseconds (ms), with low latency requirements sometimes expressed in microseconds (μ s). Latency is a function of the speed of light, hence latency increases with distance.

Application latency requirements can take different forms. A highly interactive application, such as a virtual desktop, can have a latency target measured from when a user performs an input until the user sees the virtual desktop react to that input. Voice over IP (VoIP) applications can have similar requirements. A second type of workload to consider are ones that are highly transactional,

needing a response from the server before they can continue. Databases or other forms of key/value stores can be highly impacted by increased network latency.

Jitter

Jitter measures how consistent the network latency is, and, like latency, is usually measured in milliseconds (ms).

Application jitter requirements are typically found in real time streaming applications, including video and voice delivery. These applications tend to require their data flow to be at a consistent rate and delay, with small buffers to correct for small amounts of jitter.

Packet loss

Packet loss is the measurement of what percentage of network traffic is not delivered. All networks have some degree of packet loss at times due to high traffic bursts, capacity reductions, network equipment failures, and other reasons. Thus, applications must have some tolerance of packet loss, however, how much they can tolerate can vary from application to application.

Applications that use TCP to transport their traffic have the ability to correct for packet loss via retransmission. Applications that use UDP or their own protocols on top of IP need to implement their own means of handling packet loss, and may be highly sensitive to it. A voice over IP application may simply insert silence into the part of the call that had the packet loss, as opposed to attempting a retransmit. Some VPN solutions include their own mechanisms for recovering from packet loss on the network they are using to carry traffic.

Capabilities to consider

When predictable latency and throughput are required, AWS Direct Connect is the recommended choice, as it provides deterministic performance. Bandwidth can be selected based on throughput requirements. AWS recommends using AWS Direct Connect when you require a more consistent network experience than internet-based connections can provide. Private VIFs and Transit VIFs support jumbo frames, which can reduce the number of packets through the network and can improve throughput due to reduced overhead. AWS Direct Connect [SiteLink](#) allows using the AWS backbone to provide connectivity between your locations and can be enabled on demand. Bandwidth used for SiteLink should be taken into account for your Direct Connect bandwidth selection.

Using a VPN over AWS Direct Connect adds encryption. However, it reduces the MTU size which might reduce throughput. AWS managed Site-to-Site (S2S) VPN capabilities can be found in

the [AWS Site-to-Site VPN documentation](#). Many Direct Connection locations support MACsec if encryption over your connection is the primary encryption requirement. MACsec does not have the same MTU or potential throughput considerations of Site-to-Site VPN connections. AWS Transit Gateway allows customers to horizontally scale the number of VPN connections and raise throughput accordingly with Equal-cost multi-path routing (ECMP). AWS's managed Site-to-Site VPN supports using Direct Connect transit VIFs for private connectivity – see the [Private IP VPN with AWS Direct Connect](#) for details.

Another option is to use an AWS managed Site-to-Site VPN over the internet. It can be an attractive option due to low cost and is widely available. However, keep in mind that performance over the internet is best effort. Internet weather events, congestion, and increased latency periods can be unpredictable. AWS offers a solution with [AWS Accelerated S2S VPN](#), which can mitigate some of the downsides of using an internet path. Accelerated S2S VPN uses AWS Global Accelerator, which allows VPN traffic to enter the AWS network as early and as close as possible to the customer gateway device. This optimizes the network path, using the congestion-free AWS global network, to route traffic to the endpoint that provides the best performance. You can use accelerated VPN connections to avoid network disruptions that can occur when traffic is routed over the public internet.

Cost

Definition

In the cloud, the cost of hybrid connectivity includes the cost of provisioned resources and usage. Cost of provisioned resources is measured in units of time, usually hourly. Usage is for data transfer and processing usually measured to in gigabytes (GB). Other costs include the cost of connectivity to the AWS network point of presence. If your network is within the same colocation facility, it might be as little as the cost of a cross connect. If your network is in a different location, there will be a service provider or APN Direct Connect partner costs involved.

Key questions

- How much data do you anticipate sending into AWS per month from your facility and from the internet?
- How much data do you anticipate sending from AWS per month to your facility and to the internet?
- How often will these amounts change?
- What changes in a failure scenario?

Capabilities to consider

If you have bandwidth-heavy workloads that you wish to run on AWS, AWS Direct Connect can reduce your network costs into and out of AWS in two ways. First, by transferring data to and from AWS directly, you can reduce your bandwidth costs paid to your internet service provider. Second, all data transferred over your dedicated connection is charged at the reduced AWS Direct Connect data transfer rate, rather than internet data transfer rates – see the [Direct Connect pricing page](#) for details.

AWS Direct Connect allows the use of AWS Direct Connect SiteLink to interconnect your sites using the AWS backbone – see [the SiteLink launch blog](#) for more information. Leveraging this capability incurs normal Direct Connect data transfer costs, along with a charge per hour SiteLink is enabled. You can enable and disable SiteLink on-demand, and it may be a good option for failure scenarios involving the internet or private network connectivity.

If you are using a network service provider for connectivity between on-premises and a Direct Connect location, your ability and the time needed to change your bandwidth commitments is based on your contract with the service provider.

The AWS backbone can deliver your traffic to any AWS Region except China from any AWS network point of presence. This capability has many technical benefits over using the internet to access remote AWS Regions, but has a cost – see the [EC2 Data Transfer pricing page](#) for details. If there is an [AWS Transit Gateway](#) in the traffic path, it adds data processing cost per GB, however if using inter-region peering between two Transit Gateways, you are only billed once for the Transit Gateway data processing.

Optimal application design keeps data processing within AWS and minimizes unnecessary data egress charges. Data ingress to AWS is free.

Note

As part of the overall connectivity solution, in addition to the AWS connection cost, you should also consider cost of the end-to-end connectivity including service provider cost, cross connects, racks, and equipment within DX location (if required).

If you are not sure if you should use the internet or a private connection, calculate a breakeven point where AWS Direct Connect becomes less expensive than using the internet. If the volume of

data means that AWS Direct Connect is less expensive, and you require permanent connectivity, AWS Direct Connect is the optimal connectivity choice.

If the connectivity is temporary and the internet meets other requirements, it can be cheaper to use AWS S2S VPN over the internet due to the elasticity of the internet. Note this requires that you have sufficient internet connectivity from your on-premises network.

If you are within a facility which has AWS Direct Connect (the list is [available on the Direct Connect website](#)), you can establish a cross-connect to AWS. This means using dedicated connections at 1,10, or 100Gbps. AWS Direct Connect partners offer more bandwidth options and smaller capacities, which may optimize your connectivity cost. For example, you can start at a 50 Mbps Hosted Connection versus a 1 Gbps Dedicated Connection.

With AWS Transit Gateway, you can share your VPN and Direct Connect connections with many VPCs. While you are charged for the number of connections that you make to the AWS Transit Gateway per hour and the amount of traffic that flows through AWS Transit Gateway, it simplifies management and reduces the number of VPN connections and VIFs required. The benefits and cost savings of lower operational overhead can easily outweigh the additional cost of data processing. Optionally, you can consider a design where AWS Transit Gateway is in the traffic path to most VPCs, but not all. This approach avoids the AWS Transit Gateway data processing fees for use cases where you need to transfer large amounts of data into AWS. Refer to the Connectivity Models section for further details on this design. Another approach is to combine AWS Direct Connect as a primary path with AWS S2S VPN over the internet as backup/failover path. While technically feasible and very cost effective, this solution has technical downsides (discussed in the Reliability section of this whitepaper) and can be more difficult to manage. AWS [doesn't recommend this for highly critical or critical workloads](#).

The final approach is a customer-managed VPN or SD-WAN deployed in Amazon EC2 instance(s). This can be cheaper at scale if there are tens to hundreds of site when compared to AWS S2S VPN. However, there is management overhead, licensing costs, and EC2 resource cost for each virtual appliance to consider.

Decision matrix

Table 3 – Example Corp. Automotive connectivity design inputs

Category	Customer-managed VPN or SD-WAN	AWS S2S VPN	AWS Accelerated S2S VPN	AWS Direct Connect Hosted Connection	AWS Direct Connect Dedicated Connection
Requires internet connection	Yes	Yes	Yes	No	No
Provisioned resources cost	EC2 instance and software licensing	AWS S2S VPN	AWS S2S VPN and AWS Global Accelerator	Applicable capacity slice of port cost	Dedicated port cost
Data transfer cost	Internet rate	Internet rate or Direct Connect rate	Internet with data transfer premium	Direct Connect rate	Direct Connect rate
Transit Gateway	Optional	Optional	Required	Optional	Optional
AWS Data processing cost	N/A	Only with AWS Transit Gateway	Yes	Only with AWS Transit Gateway	Only with AWS Transit Gateway
Can be used over AWS Direct Connect?	Yes	Yes	No	N/A	N/A

Connectivity design selection

This section of the whitepaper covers the considerations which affect your connectivity design selection. Connectivity design includes the logical aspects as well as how to design and optimize your hybrid connectivity reliability.

The following considerations will be covered: scalability, connectivity models, reliability, and customer-managed VPN and SD-WAN.

Considerations

- [Scalability](#)
- [Connectivity models](#)
- [Reliability](#)
- [Customer-managed VPN and SD-WAN](#)

Scalability

Definition

Scalability refers to the ability of your connectivity solution to grow and evolve over time as your requirements change.

When designing a solution, you need to consider the current size, as well as the anticipated growth. This growth can be organic growth, or might be related to rapid expansion, such as in merger and acquisition type of scenarios.

Note: depending on the targeted solution architecture, not all the preceding elements might need to be taken into consideration. However, they can serve as the foundational elements to identify the scalability requirements of most common hybrid network solutions. This whitepaper focuses on the hybrid connectivity selection and design. It is recommended that you also consider the scale of hybrid connectivity with respect to the VPC networking architecture. For more information, see the [Building a Scalable and Secure Multi-VPC AWS Network Infrastructure](#) whitepaper.

Key questions

- What is the current and anticipated number of VPCs which require connectivity to on-premises site or sites?
- Are VPCs deployed in a single AWS Region or multiple Regions?
- How many on-premises sites need to be connected to AWS?
- How many customer gateway devices (typically routers or firewalls) do you have per site that need to connect to AWS?
- How many routes are expected to be advertised to Amazon VPCs and what is the number of expected routes to be received from the AWS side?

- Is there a requirement to increase bandwidth to AWS over time?

Capabilities to consider

Scale is an important factor in hybrid connectivity design. To that point, the subsequent section will incorporate scale as a part of the targeted connectivity model design.

The following are recommended best practices to minimize scale complexity of hybrid network connectivity design:

- Route summarization should be used to reduce the number of routes advertised to and received from AWS. Thus, the IP addressing scheme needs designed to maximize the use of route summarization. Traffic engineering is a key overall consideration. For more information about traffic engineering, refer to the Traffic engineering subsection in the [Reliability](#) section.
- Minimize your number of BGP peering sessions by using DXGW with VGW or AWS Transit Gateway, where a single BGP session can provide connectivity to multiple VPCs.
- Consider Cloud WAN when multiple AWS Regions and on-premises sites need to be connected together.

Connectivity models

Definition

The connectivity model refers to the communication pattern between on-premises network(s) and the cloud resources in AWS. You can deploy cloud resources within an Amazon VPC within a single AWS Region or multiple VPCs across multiple Regions, as well as AWS services which have a public endpoint in a single or multiple AWS Regions, such as Amazon S3 and DynamoDB.

Key questions

- Is there a requirement for inter-VPC communication within a Region and across Regions?
- Is there any requirement to access AWS public endpoints directly from on-premises?
- Is there a requirement to access AWS services using VPC endpoints from on-premises?

Capabilities to consider

The following are some of the most common connectivity model scenarios. Each connectivity model covers requirements, attributes, and considerations.

Note: as highlighted earlier, this whitepaper is focused on the hybrid connectivity between on-premises networks and AWS. For further details on the design to interconnect VPCs, refer to the [Building a Scalable and Secure Multi-VPC AWS Network Infrastructure](#) whitepaper.

Models

- [AWS Accelerated Site-to-Site VPN – AWS Transit Gateway, Single AWS Region](#)
- [AWS DX – DXGW with VGW, Single Region](#)
- [AWS DX – DXGW with VGW, Multi-Regions, and AWS Public Peering](#)
- [AWS DX – DXGW with AWS Transit Gateway, Multi-Regions, and AWS Public Peering](#)
- [AWS DX – DXGW with AWS Transit Gateway, Multi-Regions \(more than 3\)](#)

AWS Accelerated Site-to-Site VPN – AWS Transit Gateway, Single AWS Region

This model is constructed of:

- Single AWS Region.
- AWS Managed Site-to-Site VPN connection with AWS Transit Gateway.
- Accelerated VPN enabled.

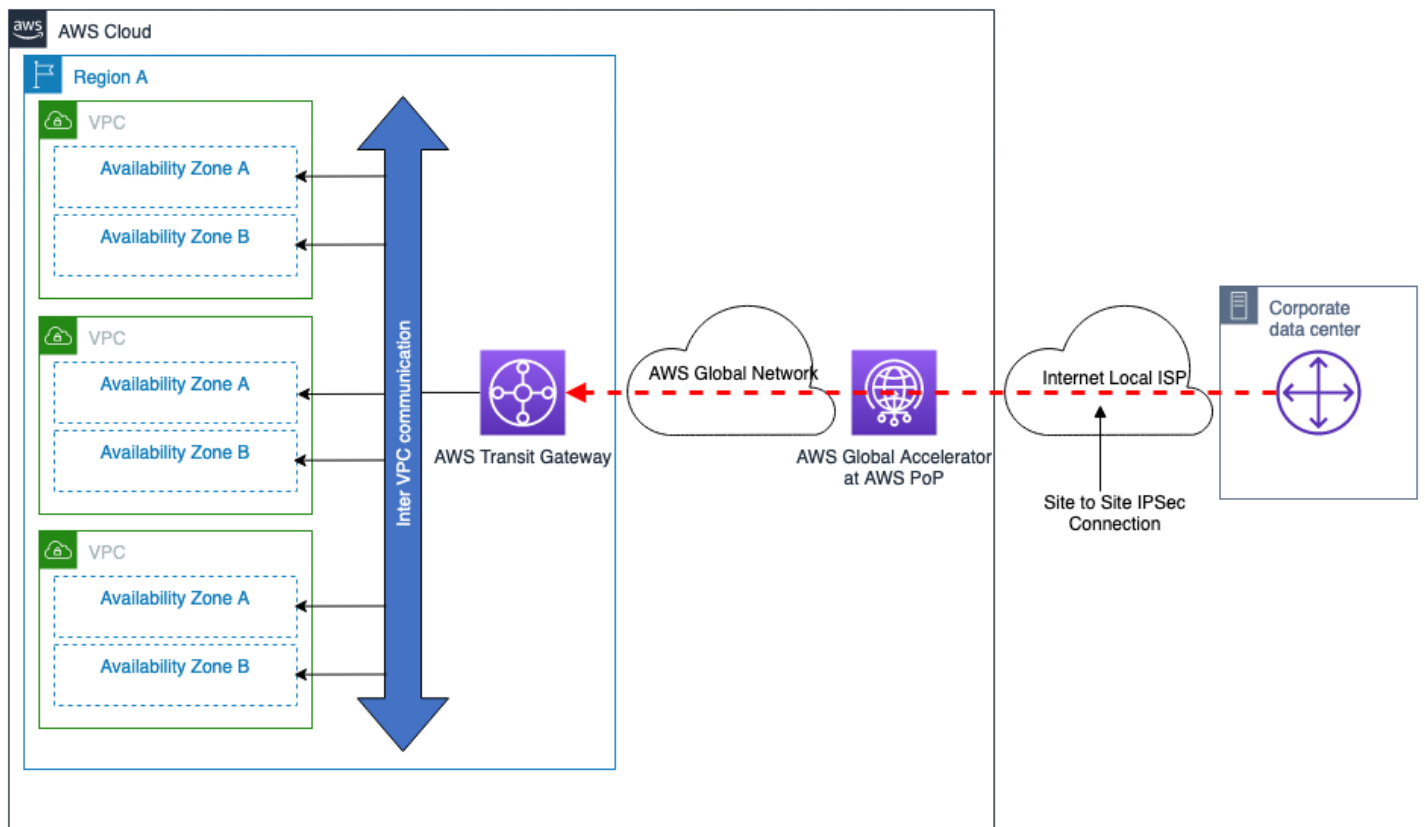


Figure 4 – AWS Managed VPN – AWS Transit Gateway, Single AWS Region

Connectivity model attributes:

- Provide the ability to establish optimized VPN connections over the public internet by using [AWS Accelerated Site-to-Site VPN connections](#).
- Provide the ability to achieve higher VPN connection bandwidth by configuring multiple VPN tunnels with ECMP.
- Can be used for connection from multiple of remote sites.
- Offers automated failover with dynamic routing (BGP).
- With AWS Transit Gateway connected to VPCs, all the connected VPCs can use the same VPN connections. You can also control the desired communication model among the VPCs, for more information refer to [How Transit Gateways Work](#).
- Offers flexible design options to integrate third-party security and SD-WAN virtual appliances with AWS Transit Gateway. See [Centralized network security for VPC-to-VPC and on-premises to VPC traffic](#).

Scale considerations:

- Up 50 Gbps of bandwidth with multiple IPsec tunnels and ECMP configured (each traffic flow will be limited to the maximum bandwidth per VPN tunnel).
- [Thousands](#) of VPCs can be connected per AWS Transit Gateway.
- Refer to the [Site-to-Site VPN quotas](#) for other scale limits, such as number of routes.

Other considerations:

- The additional AWS Transit Gateway processing costs for data transfer between the on-premises data center and AWS.
- Security groups of a remote VPC cannot be referenced in AWS Transit Gateway – this is supported by VPC peering, however.

AWS DX – DXGW with VGW, Single Region**This model is constructed of:**

- Single AWS Region.
- Dual AWS Direct Connect Connections to independent DX locations.
- AWS DXGW directly attached to the VPCs using VGW.
- Optional usage of AWS Transit Gateway for Inter-VPC communication.

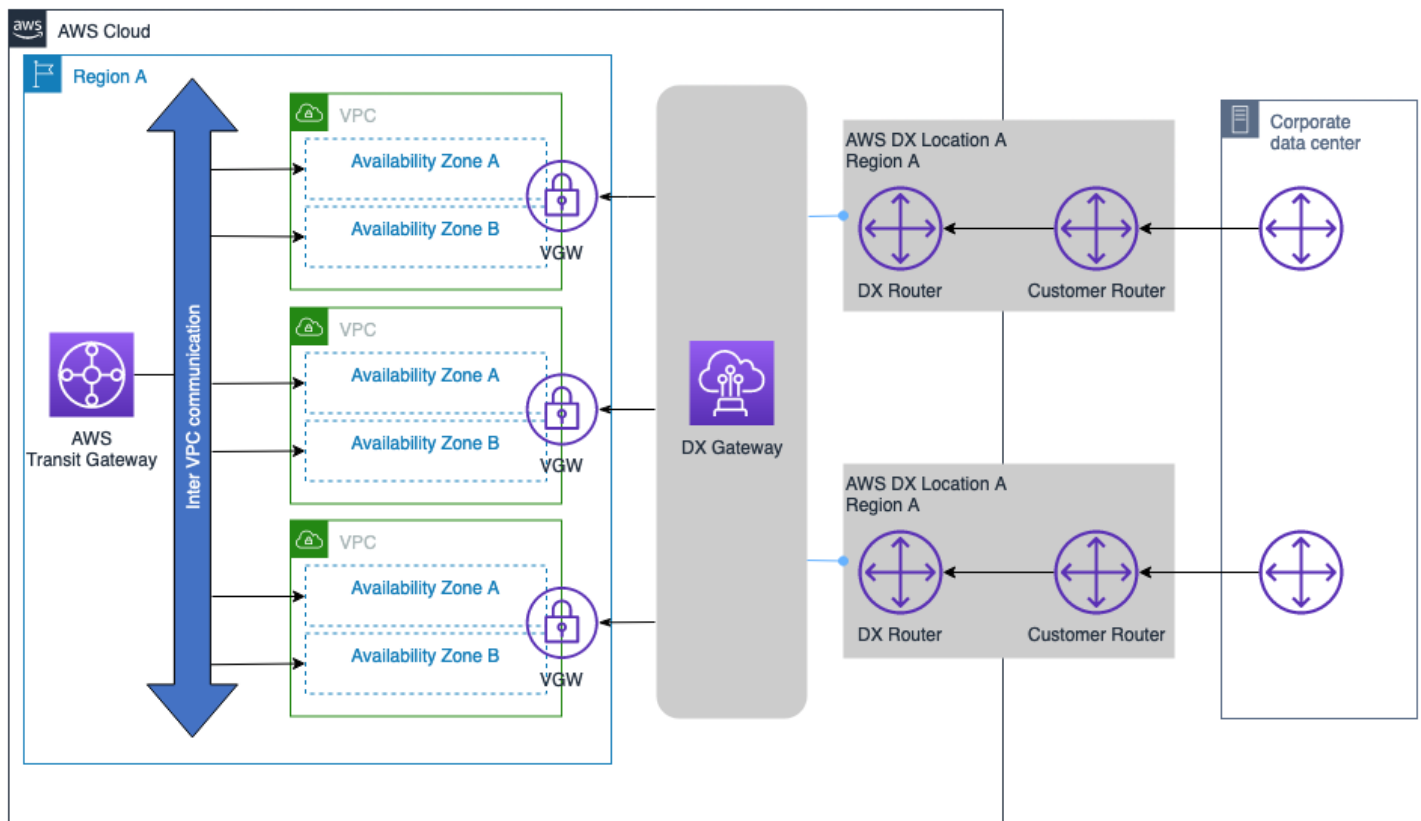


Figure 5 – AWS DX – DXGW with VGW, Single AWS Region

Connectivity model attributes:

- Provides the ability to connect to VPCs and DX connections in other Regions in the future.
- Offers automated failover with dynamic routing (BGP).
- With AWS Transit Gateway you can control the desired communication model among the VPCs. For more information, refer to [How transit gateways work](#).

Scale considerations:

Reference [AWS Direct Connect quotas](#) for more information about other scale limits, such as such number of supported prefixes, number of VIFs per DX connection type (Dedicated, hosted). Some key considerations:

- The BGP session for a private VIF may advertise up to 100 routes each for IPv4 and IPv6.

- Up to 20 VPCs can be connected per DXGW over a single BGP session. If more than 20 VPCs are needed, additional DXGWs can be added to facilitate the connectivity at scale, or consider using Transit Gateway integration.
- Additional AWS Direct Connects can be added as desired.

Other considerations:

- Does not incur AWS Transit Gateway related processing cost for data transfer between AWS and on-premises networks.
- Security groups of a remote VPC cannot be referenced over AWS Transit Gateway (need VPC peering).
- VPC peering can be used instead of AWS Transit Gateway to facilitate the communication between the VPCs, however, this adds operational complexity to build and manage large number VPC point-to-point peering at scale.
- If Inter-VPC communication is not required, neither AWS Transit Gateway nor VPC peering is required in this connectivity model.

AWS DX – DXGW with VGW, Multi-Regions, and AWS Public Peering**This model is constructed of:**

- Multiple on-premises data centers with dual connections to AWS.
- Dual AWS Direct Connect Connections to independent DX locations.
- AWS DXGW directly attached to more than 10 VPCs using VGW, up to 20 VPCs using VGW.
- Optional usage of AWS Transit Gateway for Inter-VPC and Inter-Region communication.

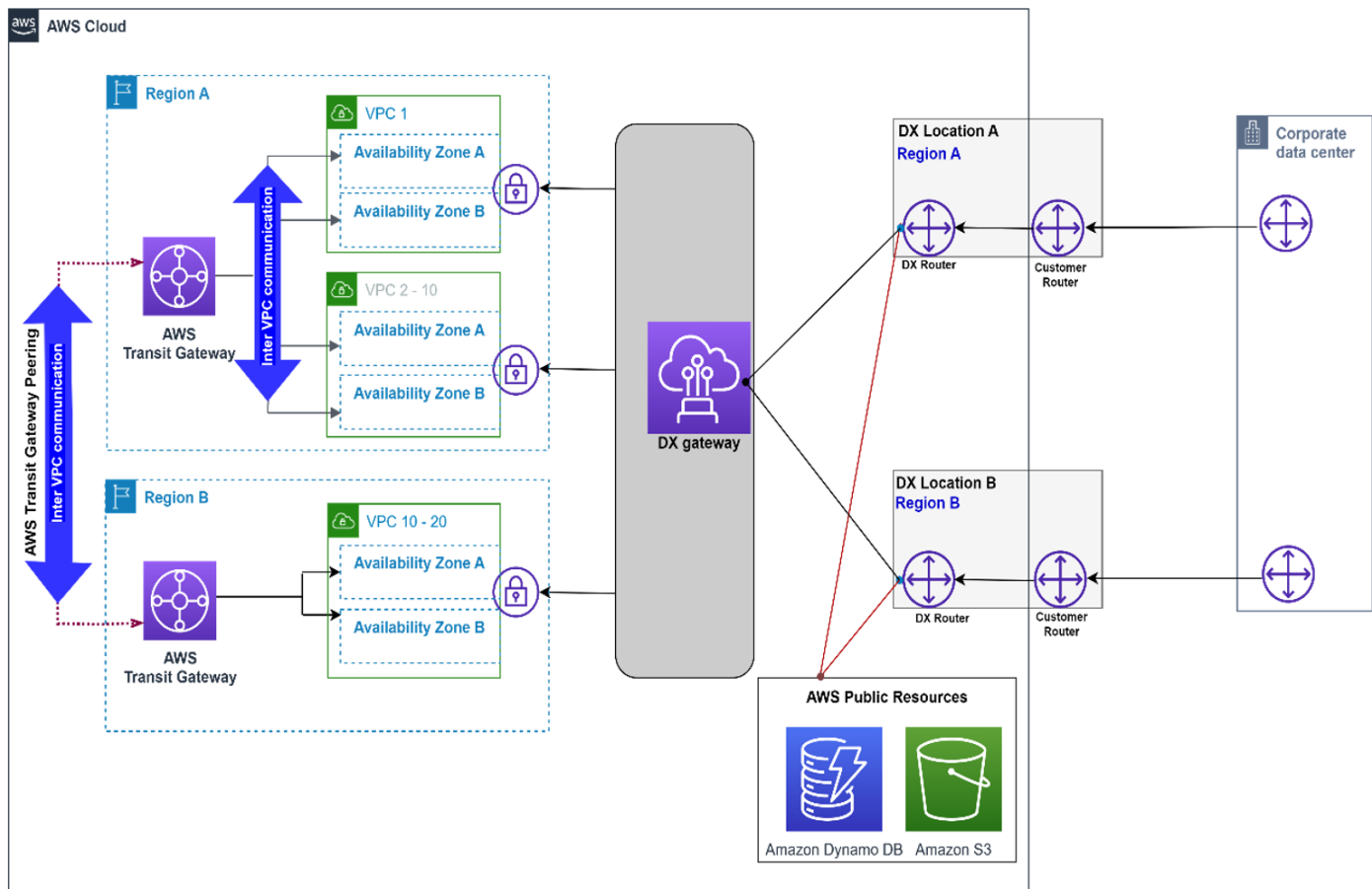


Figure 6 – AWS DX – DXGW with VGW, Multi-Regions, and Public VIF

Connectivity model attributes:

- AWS DXGW directly attached to more than 10 VPCs using VGW up to 20 VPCs using VGW.
- AWS DX public VIF is used to access AWS public services, such as Amazon S3, directly over the AWS DX connections.
- Provide the ability to connect to VPCs and DX connections in other Regions in the future.
- Inter-VPC and Inter-Region VPC communication facilitated by AWS Transit Gateway and Transit Gateway peering.

Scale considerations:

Reference [AWS Direct Connect quotas](#) for more information about other scale limits, such as such number of supported prefixes, number of VIFs per DX connection type (dedicated, hosted). Some key considerations:

- The BGP session for a private VIF can advertise up to 100 routes each for IPv4 and IPv6.
- Up to 20 VPCs can be connected per DXGW over a single BGP session on each private VIF, up to 30 private VIFs per DXGW.
- Additional AWS Direct Connects can be added as desired.

Other considerations:

- Does not incur AWS Transit Gateway related processing cost for data transfer between AWS and on-premises networks.
- Security groups of a remote VPC cannot be referenced by AWS Transit Gateway (need VPC peering).
- VPC peering can be used instead of AWS Transit Gateway to facilitate the communication between the VPCs, however, this will add operational complexity to build and manage large number VPC point-to-point peering at scale.
- If Inter-VPC communication is not required, neither AWS Transit Gateway nor VPC peering is required in this connectivity model.

AWS DX – DXGW with AWS Transit Gateway, Multi-Regions, and AWS Public Peering**This model is constructed of:**

- Multiple AWS Regions.
- Dual AWS Direct Connect Connections to independent DX locations.
- Single on-premises data center with dual connections to AWS.
- AWS DXGW with AWS Transit Gateway.
- High scale of VPCs per Region.

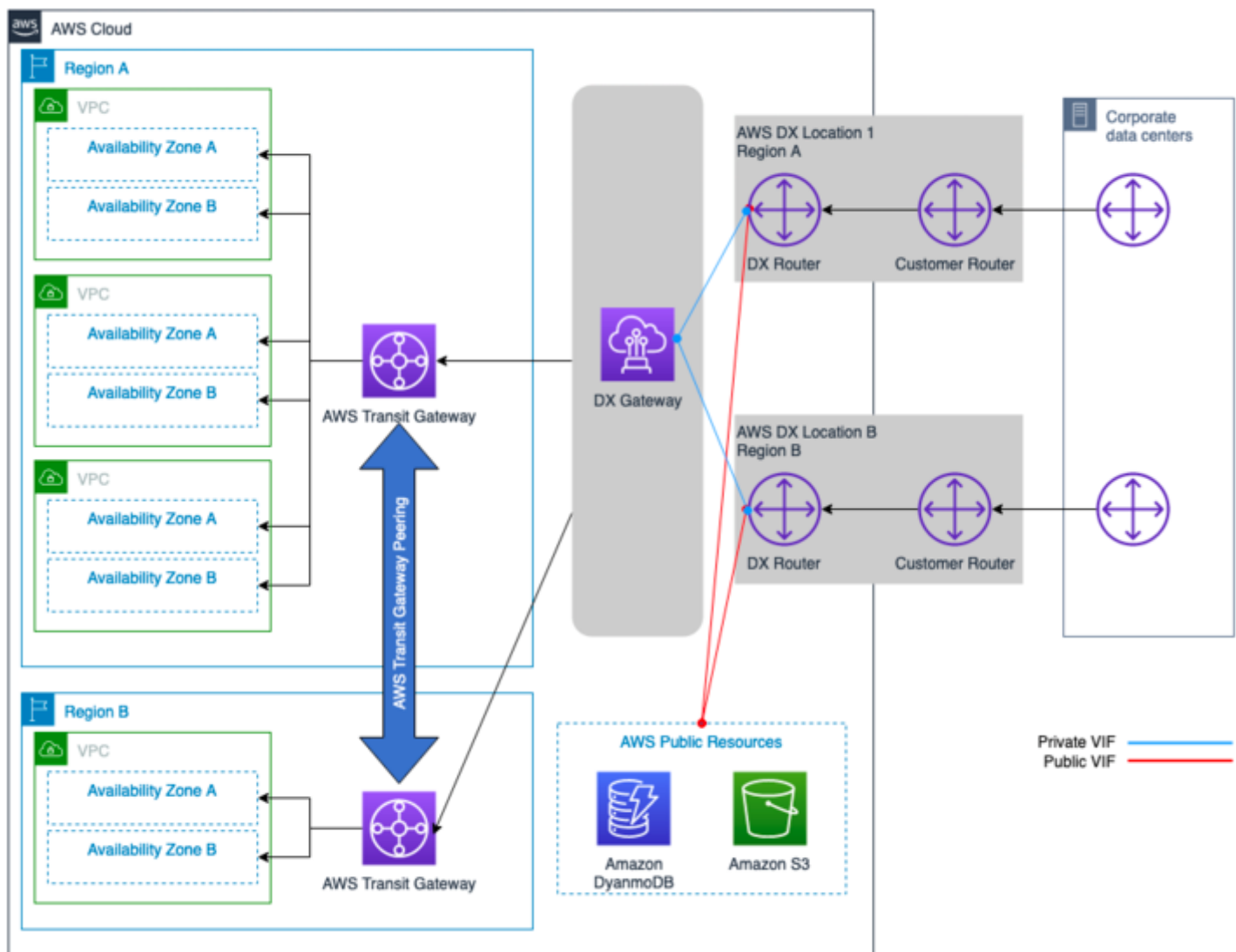


Figure 7 – AWS DX – DXGW with AWS Transit Gateway, Multi-Regions, and AWS Public VIF

Connectivity model attributes:

- AWS DX public VIF is used to access AWS public resources such as S3 directly over the AWS DX connections.
- Provide the ability to connect to VPCs and/or DX connections in other Regions in the future.
- With AWS Transit Gateway connected to VPCs, full or partial mesh connectivity can be achieved between the VPCs.
- Inter-VPC and Inter-Region VPC communication facilitated by AWS Transit Gateway peering.

- Offers flexible design options to integrate third-party security and SDWAN virtual appliances with AWS Transit Gateway. See: [Centralized network security for VPC-to-VPC and on-premises to VPC traffic](#).

Scale considerations:

- The number of routes to and from AWS Transit Gateway is limited to the maximum supported number of routes over a Transit VIF (inbound and outbound numbers vary). Refer to the [AWS Direct Connect quotas](#) for more information about the scale limits and supported number of routes and VIFs.
- Scale up to thousands of VPCs per AWS Transit Gateway over a single BGP session.
- Single Transit VIF per AWS DX.
- Additional AWS DX connections can be added as desired.

Other considerations:

- Incurs additional AWS Transit Gateway processing costs for data transfer between AWS and on-premises site.
- Security groups of a remote VPC cannot be referenced by AWS Transit Gateway (need VPC peering).
- VPC peering can be use instead of AWS Transit Gateway to facilitate the communication between the VPCs, however, this will add operational complexity to build and manage large number VPC point-to-point peering at scale.

AWS DX – DXGW with AWS Transit Gateway, Multi-Regions (more than 3)

This model is constructed of:

- Multiple AWS Regions (more than 3).
- Dual on-premises data centers.
- Dual AWS Direct Connect Connections across to independent DX locations per Region.
- AWS DXGW with AWS Transit Gateway.
- High scale of VPCs per Region.
- Full mesh of peering between AWS Transit Gateways.

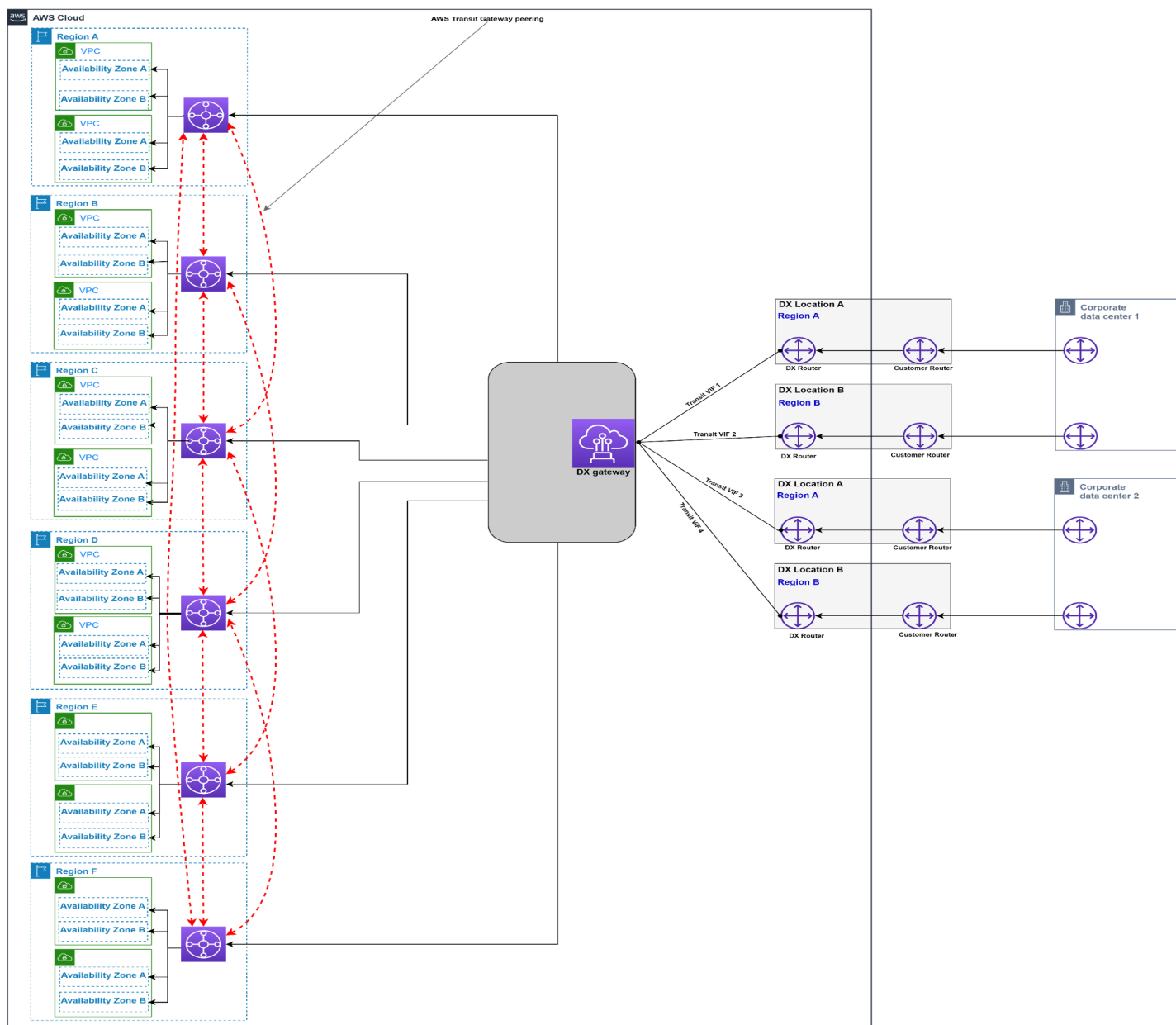


Figure 8 – AWS DX – DXGW with AWS Transit Gateway, Multi-Regions (more than three)

Connectivity model attributes:

- Lowest operational overhead.
- AWS DX public VIF is used to access AWS public resources, such as S3, directly over the AWS DX connections.
- Provide the ability to connect to VPCs and DX connections in other Regions in the future.

- With AWS Transit Gateway connected to VPCs, full or partial mesh connectivity can be achieved between the VPCs.
- Inter-Region VPC communication is facilitated by AWS Transit Gateway peering.
- Offers flexible design options to integrate third-party security and SDWAN virtual appliances with AWS Transit Gateway. See: [Centralized network security for VPC-to-VPC and on-premises to VPC traffic](#).

Scale considerations:

- The number of routes to and from AWS Transit Gateway is limited to the maximum supported number of routes over a Transit VIF (inbound and outbound numbers vary). Refer to the [AWS Direct Connect quotas](#) for more information about the scale limits. Consider route summarization if needed to reduce the number of routes.
- Scale up to thousands of VPCs per AWS Transit Gateway over a single BGP session per DXGW (assuming the provided performance by the provisioned AWS DX connections is sufficient).
- Up to six AWS Transit Gateways can be connected per DXGW.
- If more than three Regions need to be connected using AWS Transit Gateway, then additional DXGWs are required.
- Single Transit VIF per AWS DX.
- Additional AWS DX connections can be added as desired.

Other considerations:

- Incurs additional AWS Transit Gateway processing cost for data transfer between the on-premises site and AWS.
- Security groups of a remote VPC cannot be referenced by AWS Transit Gateway (need VPC peering).
- VPC peering can be used instead of AWS Transit Gateway to facilitate the communication between the VPCs, however, this will add operational complexity to build and manage large number VPC point-to-point peering at scale.

The following decision tree covers the scalability and communication model considerations:

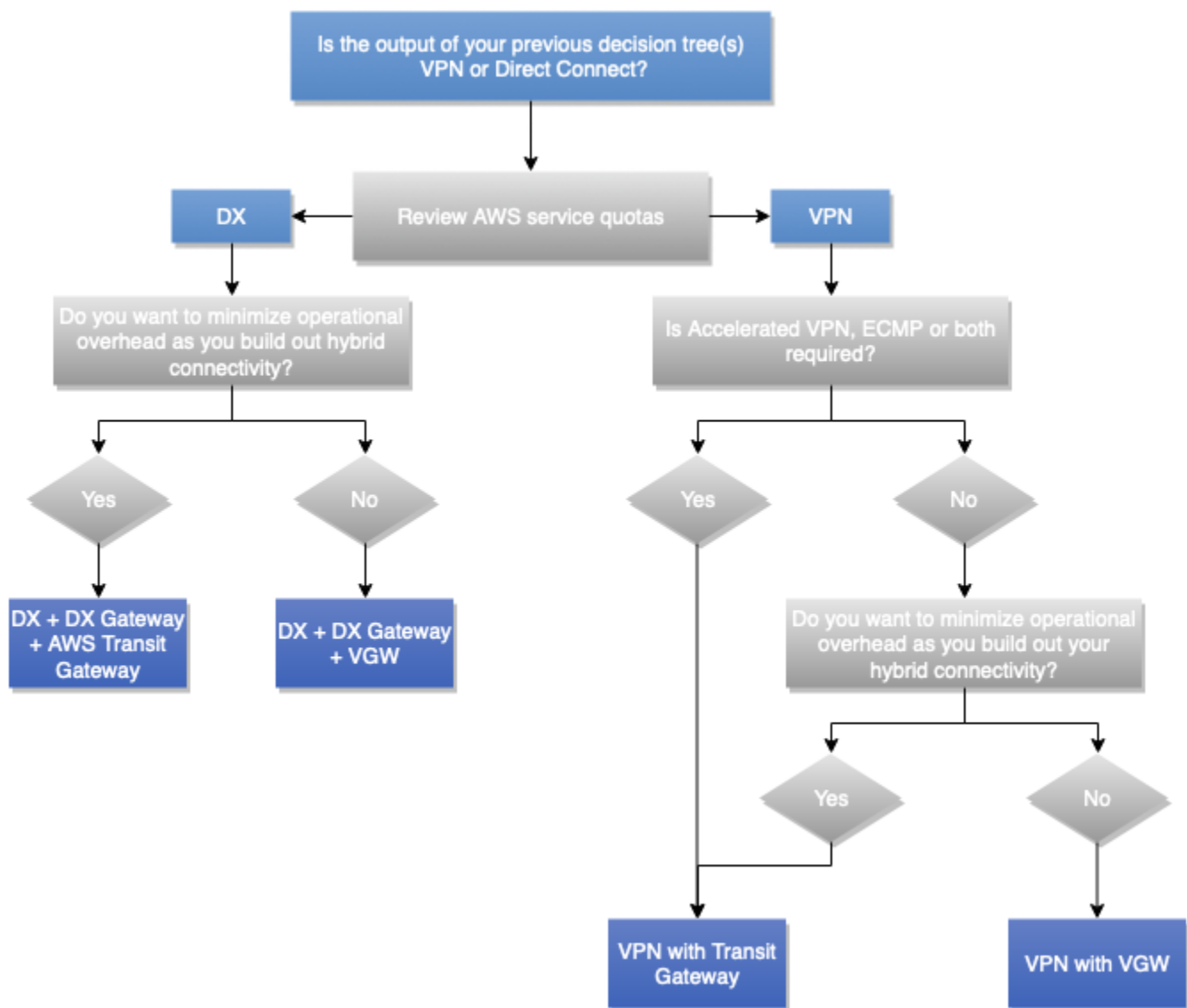


Figure 9 – Scalability and communication model decision tree

Note

If the selected connection type is VPN, typically at the performance consideration, the decision should be made whether the VPN termination point is AWS VGW or AWS Transit Gateway AWS S2S VPN connection. If not made yet, then you can consider the required communication model between the VPC along with the number of required VPC to be connected to the VPN connection(s) to help you make the decision.

Reliability

Definition

Reliability refers to the ability of a service or system to perform its expected function when required. The reliability of a system can be measured by the level of its operational quality within a given timeframe. Contrast this to resiliency, which refers to the ability of a system to recover from infrastructure or service disruptions, dynamically and reliably.

For more details of how availability and resiliency are used to measure reliability, refer to the [Reliability Pillar](#) of the AWS Well-Architected Framework.

Key questions

Availability

Availability is the percentage of time that a workload is available for use. Common targets include 99% (3.65 days of downtime allowed per year), 99.9% (8.77 hours), and 99.99% (52.6 minutes), with a shorthand of the number of nines in the percentage ("two nines" for 99%, "three nines" for 99.9%, and so on). The availability of the networking solution between AWS and the on-premises data center may be different than overall solution or application availability.

Key questions for the availability of a networking solution include:

- Can my AWS resources continue to operate if they cannot communicate to my on-premises resources? Vice versa?
- Should I consider scheduled downtime for planned maintenance as included or excluded from the availability metric?
- How will I measure the availability of the networking layer, separate from overall application health?

The [Availability section](#) of the Well-Architected Framework Reliability Pillar has suggestions and formulas for calculation availability.

Resiliency

Resiliency is the ability of a workload to recover from infrastructure or service disruptions, dynamically acquire computing resources to meet demand, and mitigate disruptions, such as misconfigurations or transient network issues. If a redundant network component (link, network

devices, and so on) does not have sufficient availability to provide the expected function on its own, then it has low resiliency to failures. The consequence is a poor and degraded user experience.

Key questions for resiliency of a networking solution include:

- How many simultaneous, discrete failures should I allow for?
- How can I reduce single points of failure with both the connectivity solutions and my internal network?
- What is my vulnerability to distributed denial of service (DDoS) events?

Technical solution

First, it is important to note that not every hybrid network connectivity solution requires a high level of reliability, and that increasing levels of reliability have a corresponding increase in cost. In some scenarios, a primary site may require reliable (redundant and resilient) connections as the downtime has a higher impact on the business, while regional sites, may not require the same level of reliability due to the lower impact on the business in case of a failure event. It is recommended to refer to the [AWS Direct Connect Resiliency Recommendations](#) as it explains the AWS best practices for ensuring high resiliency with AWS Direct Connect design.

To achieve a reliable hybrid network connectivity solution in the context of resiliency, the design needs to take into consideration the following aspects:

- **Redundancy:** Aim to eliminate any single point of failure in the hybrid network connectivity path, including but not limited to network connections, edge network devices, redundancy across Availability Zones, AWS Regions, and DX locations, and device power sources, fiber paths, and operating systems. For the purpose and scope of this whitepaper, redundancy focuses on the network connections, edge devices (for example, customer gateway devices), AWS DX location, and AWS Regions (for multi-Region architectures).
- **Reliable failover components:** In some scenarios, a system might be functional, but not performing its functions at the required level. Such a situation is common during a single failure event where it is discovered that planned redundant components were operating non-redundantly - their networking load has no other place to go to due to usage, which results in insufficient capacity for the entire solution.
- **Failover time:** Failover time is the time it takes for a secondary component to fully take over the role of the primary component. Failover time has multiple factors – how long it takes to detect the failure, how long to enable secondary connectivity, and how long to notify the remainder of

the network of the change. Failure detection can be improved using Dead Peer Detection (DPD) for VPN links, and Bidirectional Forwarding Detection (BFD) for AWS Direct Connect links. The time to enable secondary connectivity can be very low (if these connections are always active), may be a short time window (if a pre-configured VPN connection needs to be enabled), or longer (if physical resources need to be moved or new resources configured). Notifying the remainder of the network usually occurs via routing protocols inside the customer's network, each of which has different convergence times and options for configuration – the configuration of these is outside the scope of this whitepaper.

- **Traffic Engineering:** Traffic engineering in the context of resilient hybrid network connectivity design aims to address how traffic should flow over multiple available connections in normal and failure scenarios. It is recommended to follow the concept of *design for failure*, where you need to look at how the solution will operate in different failure scenarios and whether it will be acceptable by the business or not. This section discusses some of the common traffic engineering uses case that aims to enhance the overall resiliency level of the hybrid network connectivity solution. The [AWS Direct Connect section on routing and BGP](#) talks about several traffic engineering options for influencing traffic flow (communities, BGP local preference, AS Path length). To design an effective traffic engineering solution, you need to have a good understanding of how each of the AWS networking components handle IP routing in terms of route evaluation and selection, as well as the possible mechanisms to influence the route selection. The details of this are outside the scope of this document. For more information, see [Transit Gateway Route Evaluation Order](#), [Site-to-Site VPN Route Priority](#), and [Direct Connect Routing and BGP](#) documentation as needed.

Note

In the VPC route table, you might reference a prefix list which has additional route selection rules. For more information about this use case, refer to [route priority for prefix lists](#). AWS Transit Gateway route tables also support prefix lists, but once applied they get expanded to specific route entries.

Dual Site-to-Site VPN connections with more specific routes example

This scenario is based on a small on-premises site connecting to a single AWS Region over redundant VPN connections via the internet to AWS Transit Gateway. The traffic engineering design

depicted in Figure 10 shows that with traffic engineering you can influence the path selection that increases the hybrid connectivity solution reliability by:

- **Resilient hybrid connectivity:** Redundant VPN connections each provide the same performance capacity, support automated failover by using dynamic routing protocol (BGP), and speed up connection failure detection by using VPN dead peer detection.
- **Performance efficiency:** Configuring ECMP across both VPN connections to AWS Transit Gateway helps to maximize the overall VPN connection bandwidth. Alternatively, by advertising different, more specific, routes along with the site summary route, load can be managed independency across the two VPN connections

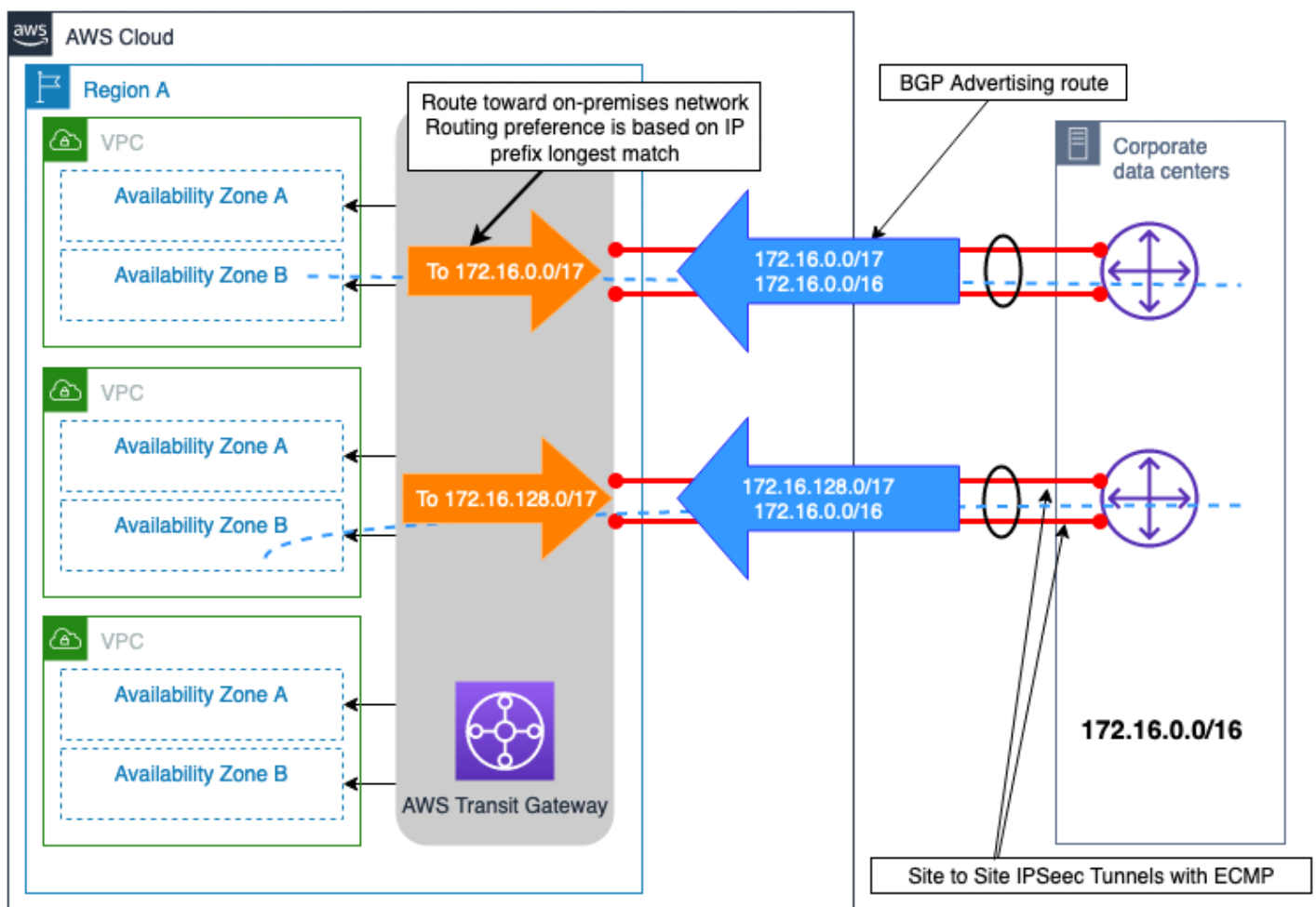


Figure 10 – Dual Site-to-Site VPN connections with more specific routes example

Dual on-premises sites with multiple DX connections example

The scenario illustrated in Figure 11 shows two on-premises data center sites located in different geographical Regions, and connected to AWS using the Maximum Resiliency connectivity model (described in the [AWS Direct Connect Resiliency Recommendations](#)) using AWS Direct Connect with DXGW and VGW. These two on-premises sites are interconnected to each other over a data center interconnect (DCI) link. The on-premises IP prefixes (192.168.0.0/16) that belongs to remote branch sites are advertised from both on-premises data center sites. The primary path for this prefix should be data center 1. Traffic to and from the remote branch sites will failover to data center 2 in a failure event of data center 1 or both DX locations. Also, there is a site-specific IP prefix for each data center. These prefixes need to be reached directly, and via the other data center site in case of both DX locations failure.

By associating BGP Community attributes with the routes advertised to AWS DXGW, you can influence the egress path selection from AWS DXGW side. These community attributes control AWS's BGP Local Preference attribute assigned to the advertised route. For more information, refer to AWS DX [Routing policies and BGP communities](#).

To maximize the reliability of the connectivity at the AWS Region level, each pair of AWS DX connections configures ECMP so that both can be utilized at the same time for data transfer between each on-premises site and AWS.

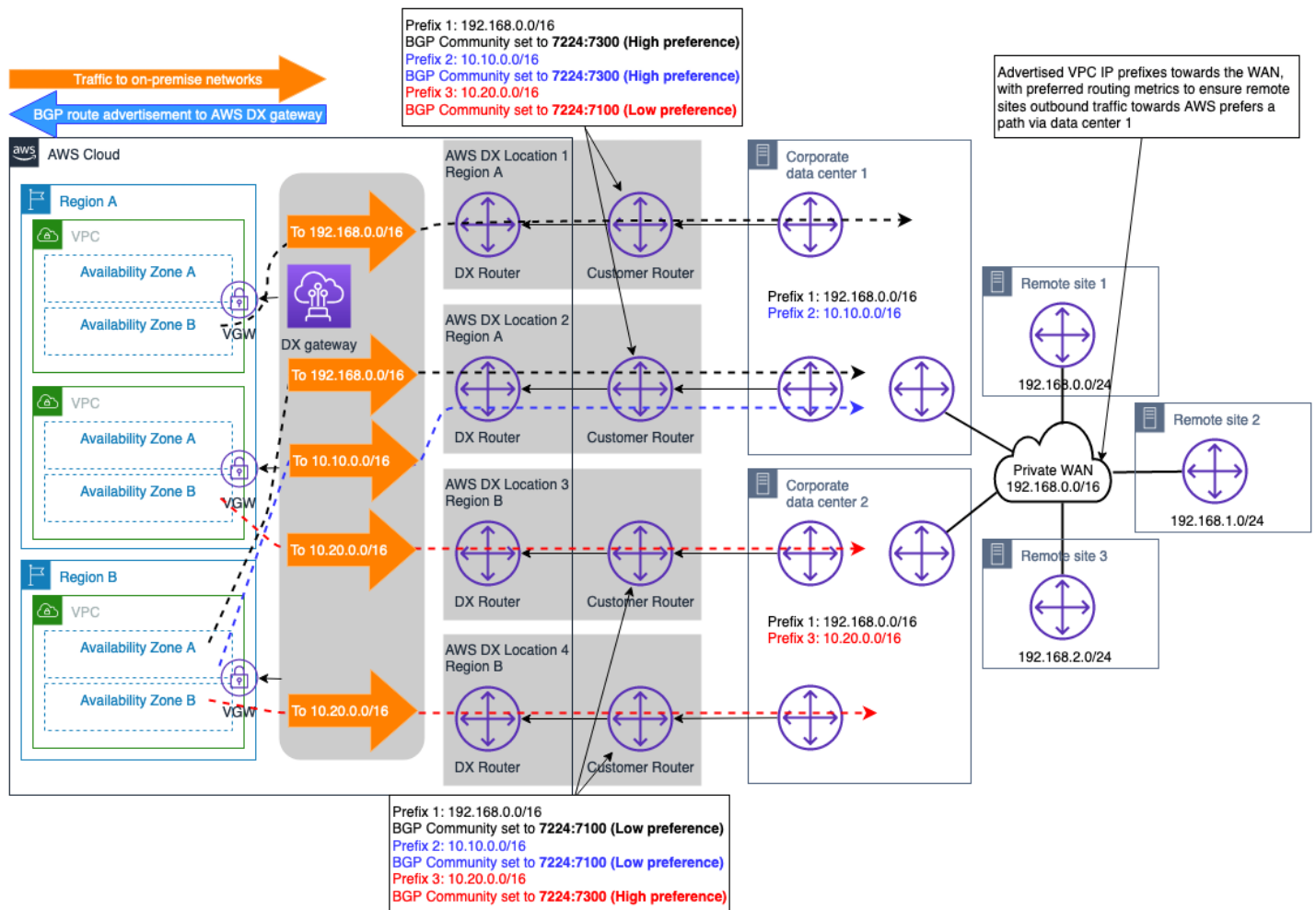


Figure 11 – Dual on-premises sites with multiple DX connections example

With this design, the traffic flows destined to the on-premises networks (with the same advertised prefix length and BGP community) will be distributed across the dual DX connections per site using ECMP. However, if ECMP is not required across the DX connection, the same concept discussed earlier and described in the [Routing policies and BGP communities](#) documentation can be used to further engineer the path selection at a DX connection level.

Note: If there are security devices in the path within the on-premises data centers, these devices need to be configured to allow traffic flows leaving over one DX link and coming from another DX link (both links utilized with ECMP) within the same data center site.

VPN connection as a backup to AWS DX connection example

VPN can be selected to provide a backup network connection to an AWS Direct Connect connection. Typically, this type of connectivity model is driven by cost, because it provides a lower level of reliability to the overall hybrid connectivity solution due to indeterministic performance over the internet, and there is no SLA that can be obtained for a connection over the public internet. It is a valid and cost-effective connectivity model, and should be used when cost is the top priority consideration and there is a limited budget, or possibly as an interim solution until a secondary DX can be provisioned. Figure 12 illustrates the design of this connectivity model. One key consideration with this design, where both the VPN and DX connections are terminating at the AWS Transit Gateway, is that the VPN connection can advertise higher number of routes compared to the ones that can be advertised over a DX connection connected to AWS Transit Gateway. This may cause a suboptimal routing situation. An option to resolve this issue is to configure route filtering at the customer gateway device (CGW) for the routes received from the VPN connection, allowing only the summary routes to be accepted.

Note: To create the summary route on the AWS Transit Gateway, you need to specify a static route to an arbitrary attachment in the AWS Transit Gateway route table so that the summary is sent along the more specific route.

From the AWS Transit Gateway routing table's point of view, the routes for the on-premises prefix are received both from the AWS DX connection (via DXGW) and from VPN, with the same prefix length. Following the [route priority logic of AWS Transit Gateway](#), routes received over Direct Connect have a higher preference than the ones received over Site-to-Site VPN, and thus the path over the AWS Direct Connect will be the preferred to reach the on-premises network(s).

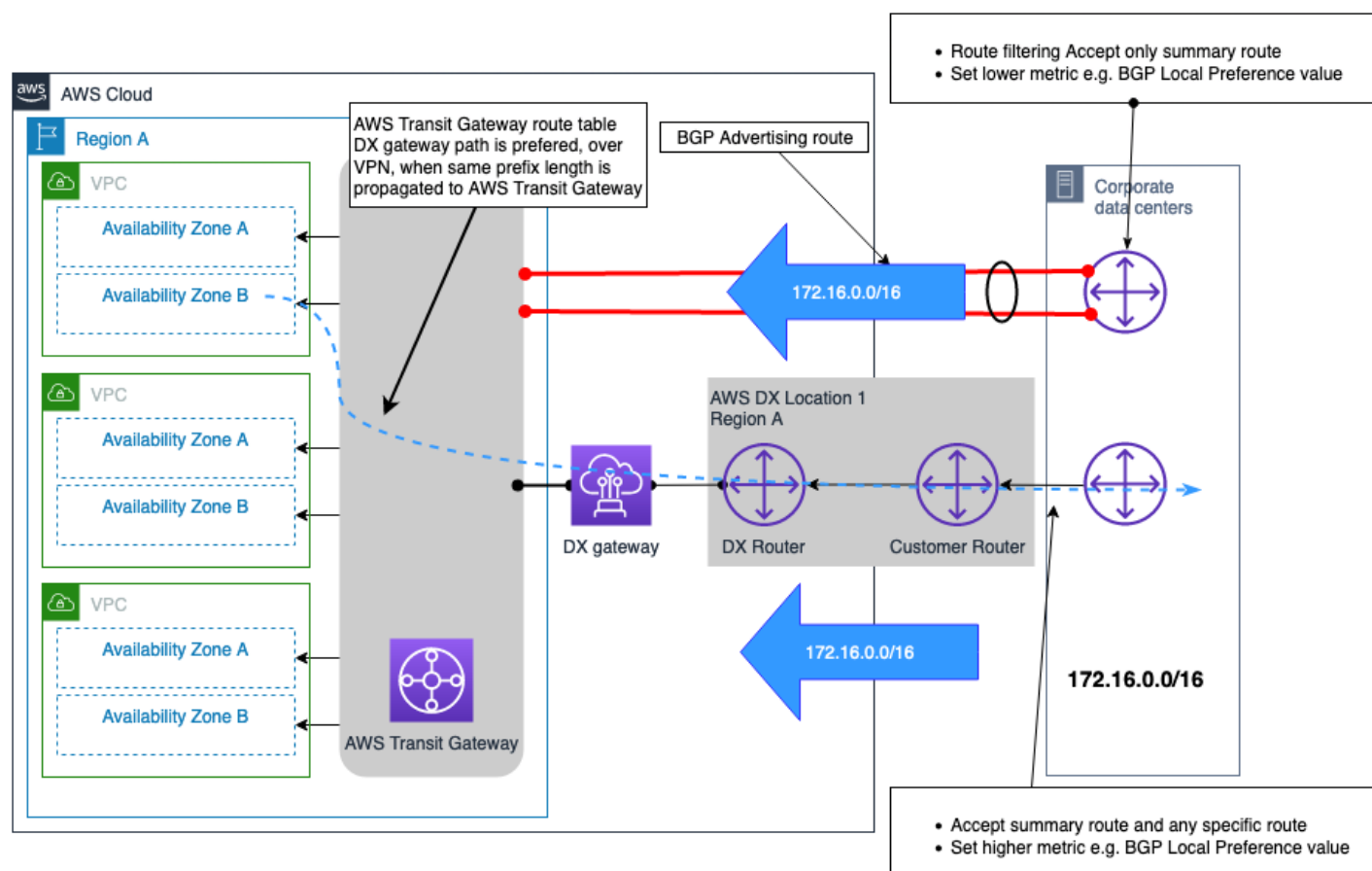


Figure 12 – VPN connection as a backup to AWS DX connection example

The following decision tree guides you through making the desired decision for achieving a resilient (which will result in a reliably) hybrid network connectivity. For more information, refer to [AWS Direct Connect Resiliency Toolkit](#).

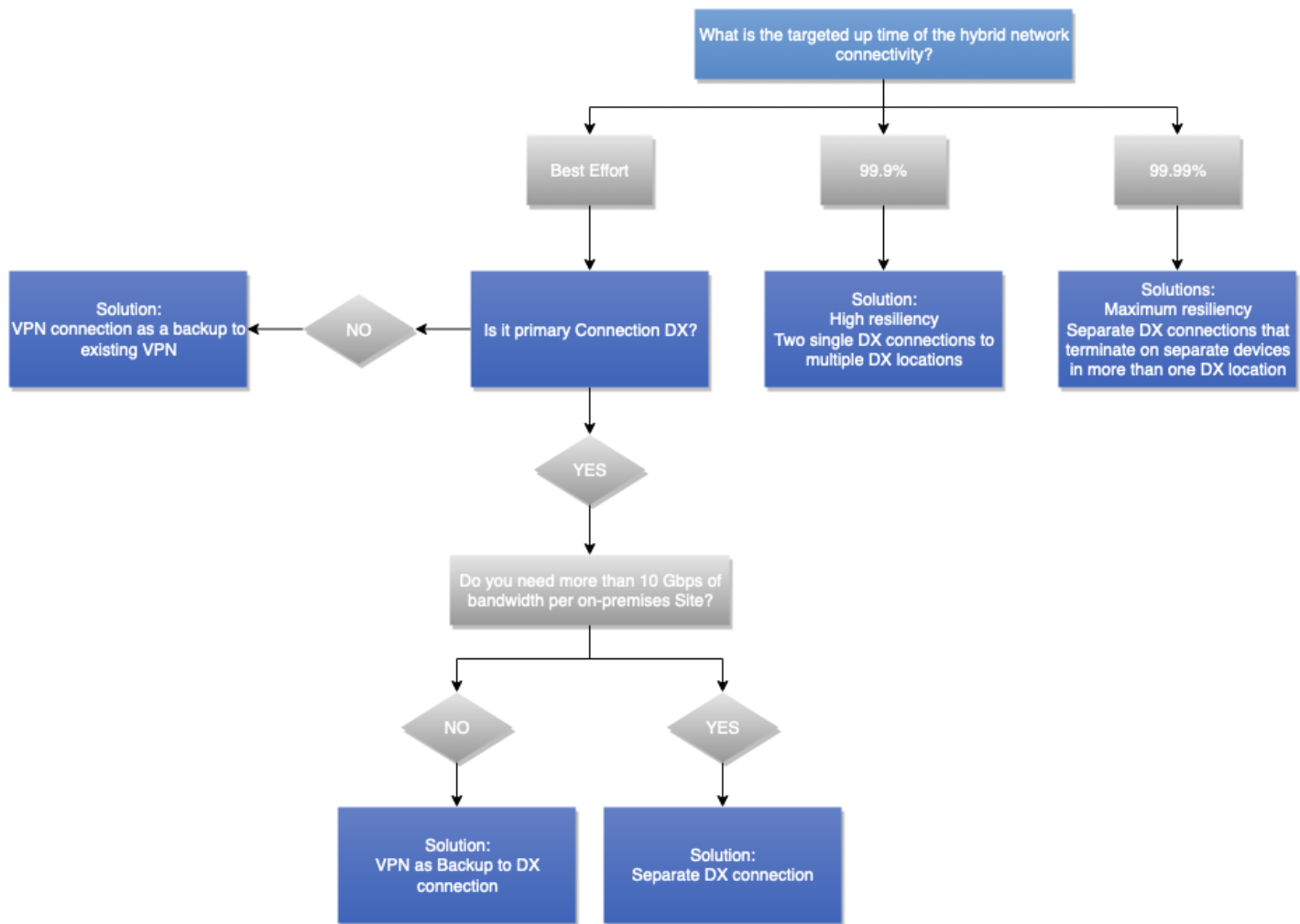


Figure 13 – Reliability decision tree

Customer-managed VPN and SD-WAN

Definition

Connectivity to the internet is a commodity and available bandwidth continues to increase every year. Some customers choose to build a virtual WAN on top of the internet instead of building and operating a private WAN. A software-defined wide area network (SD-WAN) allows companies to rapidly provision and manage centrally this virtual WAN through clever use of software. Other customers choose to adopt traditional self-managed site to site VPNs.

Impact on design decisions

SD-WAN and customer-managed VPNs can run over internet or AWS Direct Connect. SD-WAN (or any software VPN overlay) is as reliable as the underlying network transport. Therefore, the

reliability and SLA considerations discussed earlier in this whitepaper are applicable here. For instance, building a SD-WAN overlay over the internet will not offer the same reliability versus if it's built over an AWS Direct Connect.

Requirement definition

- Do you use SD-WAN in your on-premises network?
- Are there specific features you require which are only available on certain virtual appliances used for VPN termination?

Technical solutions

AWS recommends integrating SD-WAN with AWS Transit Gateway, and publishes a list of [the vendors who support AWS Transit Gateway integration](#). AWS can act as a hub for SD-WAN sites or as a spoke site. The AWS backbone can be used to connect different SD-WAN hubs deployed in AWS with a highly reliable and performant network. SD-WAN solutions support automated failover through any available path, additional monitoring, and observability capabilities in a single management pane. Extensive use of auto configuration and automation allows rapid provisioning and visibility compare to traditional WANs. However, the use of tunneling and encryption overheads do not compare to dedicated, high-speed fiber links used in private connectivity.

In some cases, you may choose to use a virtual appliance with VPN capability. Reasons for selecting a self-managed virtual appliance include technical features and compatibility with the rest of your network. When you select a self-managed VPN or an SD-WAN solution which uses a virtual appliance deployed in an EC2 instance, you are responsible for the management of such appliance. You are also responsible for high availability and failover between virtual appliances. Such design increases your operational responsibility; however, it could provide you more flexibility. The features and capabilities of the solution depend on the virtual appliance you select.

AWS Marketplace contains many VPN virtual appliances which customers can deploy on Amazon EC2. AWS recommends starting with AWS managed S2S VPN and look at other options if it doesn't meet your requirements. The management overhead of virtual appliances is the customer responsibility.

Example Corp. Automotive use case

This section of the whitepaper demonstrates how the considerations, requirement definition questions, and decision trees are used to help you decide on the optimal hybrid network design. Identifying and capturing requirements is important since they are used as input to the decision trees. Capturing requirements upfront avoids further design iterations. Halting a project altogether if the design must be revisited and having valuable resources on hold can be minimized and ideally avoided when requirements are understood upfront.

Example Corp. Automotive will be used throughout this section as the illustrative customer. They are looking to initially deploy their first analytics project on AWS. The analytics project is focused on analyzing data from cars manufactured by the company and other datasets that already exist in the company's data centers. Initially, the company's architecture group thinks they will need an AWS account, an Amazon VPC, and few subnets to host production and development environments. The project team is eager to get started, and they requested development environment access as soon as possible. They aim to go in production three months from now.

Example Corp. Automotive also plans to use AWS for several additional projects, such as migrating their ERP systems, Virtual Desktop Infrastructure (VDI), and another 20 applications from on-premises to AWS over the next 6 months. Some requirements for additional projects are still being defined, but it's clear that their AWS Cloud usage is going to grow.

The architecture team decided to leverage the approach outlined in this whitepaper. They used the requirement definition questions outlined under each consideration to capture the inputs to make their design decisions.

They start with requirements related to the connectivity type which are summarized in the following table.

Table 4 – Example Automotive Corp reliability inputs

Connectivity type selection considerations	Requirement definition questions	Answers
Time to Deploy	What is the required timeline for the deployment? Hours, days, weeks, or months?	• Dev/Test: 1 month • Production: 3 months

Connectivity type selection considerations	Requirement definition questions	Answers
Security	Do your security requirements and policies allow the usage of encrypted connections over the internet to connect to AWS or mandate the usage of private network connections?	• Dev/Test: Site-to-Site VPN acceptable • Production: Private network required
	When leveraging private network connections, does the network layer have to provide encryption in transit?	No, application layer encryption will be used.
SLA	Is hybrid connectivity SLA with service credits required?	• Dev/Test: No • Production: Yes
	What is the uptime target?	• Dev/Test: N/A • Production: 99.99%
	Does the entire hybrid network adhere to the uptime target?	• Dev/Test: N/A • Production: Yes
Performance	What is the required throughput?	• Dev/Test: 100 Mbps • Production: 500 Mbps growing to 2 Gbps
	What is the maximum acceptable latency between AWS and on-premises network?	• Dev/Test: No hard requirements • Production: Less than 30 ms

Connectivity type selection considerations	Requirement definition questions	Answers
	What is the maximum acceptable network jitter?	• Dev/Test: No hard requirements • Production: Minimum jitter required
Cost	How much data would you send to AWS per month?	• Dev/Test: 2 TB • Production: 20 TB growing to 50 TB
	How much data would you send from AWS per month?	• Dev/Test: 1 TB • Production: 10 TB growing to 25 TB
	Is this connectivity permanent ?	Yes

Based on requirements received, the architecture team followed the connectivity type decision tree from Figure 9. It allowed the architecture team to decide on the connectivity type for the development and test and production environments. For the production environment, they considered the immediate as well as the upcoming requirements. For development and test Example Corp. Automotive will establish a site-to-site VPN over the internet. For production, they are going to work with a service provider to connect their corporate network with AWS Direct Connect. Example Corp. Automotive initially considered using a Direct Connect Hosted Connection, however due to the requirements for an [AWS provided SLA](#) they selected Direct Connect Dedicated Connections.

After deciding on the connectivity type, the next step is to capture the requirements which impacts the connectivity design selection. This is related with the logical design, such how the connections are configured and which AWS services to use to support business and technical requirements.

To capture the scalability and communication model requirements, the architecture team used the requirement definition questions from the associated sections of this whitepaper. The requirements related to those two considerations are summarized in the following table.

Table 5 – Requirement definition questions

Connectivity design selection considerations	Requirement definition questions	Answers
Scalability	What is the current or anticipated number of VPCs which require connectivity to on-premises sites?	2 initially, growing to 30 in 6 months
	Are these VPCs deployed in a single AWS Region or multiple Regions?	Single Region
	How many on-premises sites need to be connected to AWS?	2 data centers
	How many customer gateway devices do you have, per site, that need to connect to AWS?	2 routers per data center
	How many routes are expected to be advertised to AWS VPCs as well as the number of expected routes to be received from AWS side?	• Routes to be advertised to AWS: 20 routes • Routes to be received from AWS: 1 /16 route
	Is there any plan to consider bandwidth increase of the connection to AWS in the near future?	• Dev/Test: 100 Mbps • Production: 500Mbps growing to 2Gbps.
Connectivity design models	Is there a requirement for inter-VPC communication to be enabled (within a Region and/or across Regions)?	Yes, within an AWS Region

Connectivity design selection considerations	Requirement definition questions	Answers
	Is there a requirement to access AWS public endpoints services directly from on-premises?	Yes
	Is there a requirement to access AWS services using VPC endpoints from on premises?	No

Based on inputs, the architecture team followed the decision tree from the Connectivity Design section. After anticipating that the number of VPCs is going to grow from 2 to 30 in the next 6 months, the architecture team decided to use AWS Transit Gateway as the termination gateway for the connection and for inter-VPC routing. Independent AWS Transit Gateways will terminate the VPN connection used for development and testing, and for the production connectivity with AWS Direct Connect. The usage of separated AWS Transit Gateways makes change management simpler and provides a clear demarcation between dev/test and production environments. For the production, AWS Direct Connect gateway is required because of AWS Transit Gateway. A public VIF will be used for access to AWS public endpoint services. Figure 14 illustrates the path taken on the decision tree based on requirements collected.

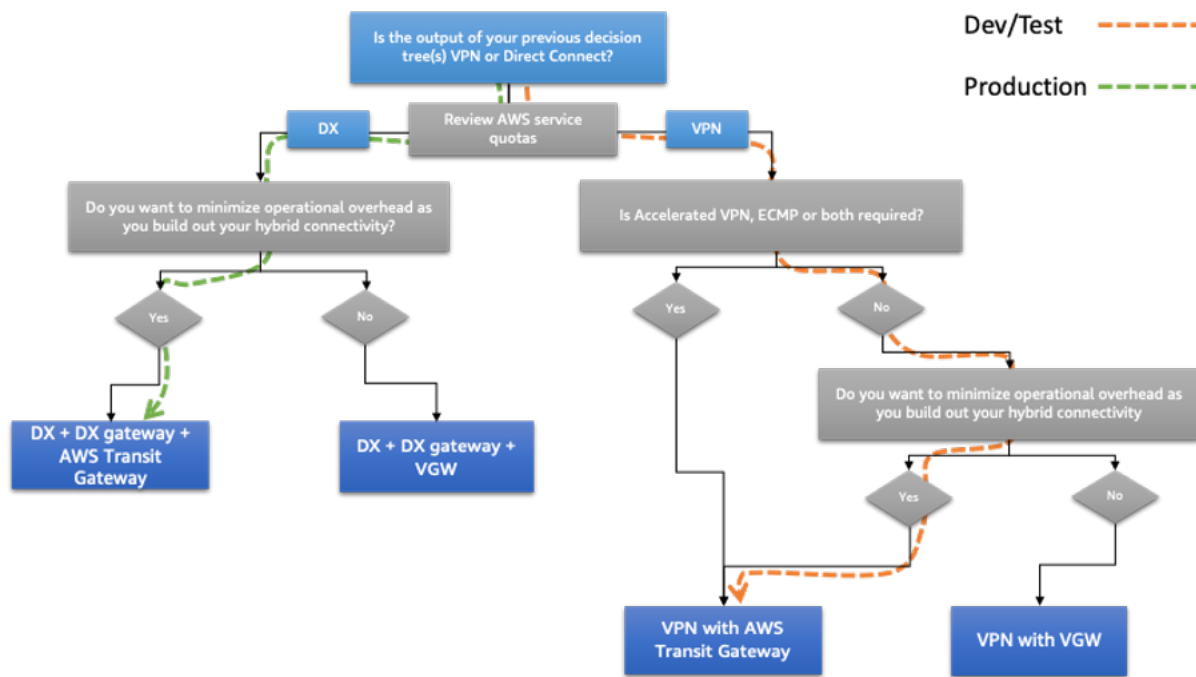


Figure 14 – Example Corp. Automotive connection design decision tree

After deciding on the solution to meet the scalability and communication model requirements, the next step is to capture the requirements associated with reliability. This is related to the required level of availability and resilience.

To capture the reliability requirements, the architecture team used the requirement definition questions from the associated section of this whitepaper. The requirements are summarized in the following table.

Table 6 – Reliability requirement questions

Connectivity design selection considerations	Requirement definition questions	Answers
Reliability	What is the impact magnitude on the business in case of a connectivity failure to AWS?	• Dev/Test: Low • Production: High
	From a business point of view, does the cost of following a connectivity failure to AWS outweigh the cost of	• Dev/Test: No • Production: Yes

Connectivity design selection considerations	Requirement definition questions	Answers
	deploying a highly reliable connectivity model to AWS?	

Based on inputs received, the architecture team followed the decision tree from the reliability considerations sections covered previously on this whitepaper. After considering the uptime target of 99.99% for the production connectivity and the high business impact if there was a service interruption, the architecture team decided to use 2 Direct Connect locations and have 2 links from each on-premises data center to each Direct Connect location (4 links in total). The VPN connectivity used for development and testing will also use two VPN connections for additional redundancy. Using route engineering techniques discussed in the reliability section, connectivity will be configured as follows:

- For development and testing, traffic is going to be load balanced using ECMP over the 2 tunnels going to the primary data center. This allows for higher throughput. The tunnels going to the secondary data center are going to be used in case of failure of the primary tunnels.
- For production, the latency between on-premises and AWS over either of the Direct Connect locations is very similar. In this case, it has been decided to load balance the traffic between AWS and on-premises over the two connections going to the primary data center for the on-premises systems deployed in the primary data center. Similarly, for on-premises systems running in the secondary data center, traffic is going to be load balanced between the two connections to the secondary data center. In case of failure of the connections, BGP will facilitate an automated failover.

Figure 15 illustrates the path taken on the decision tree based on requirements collected.

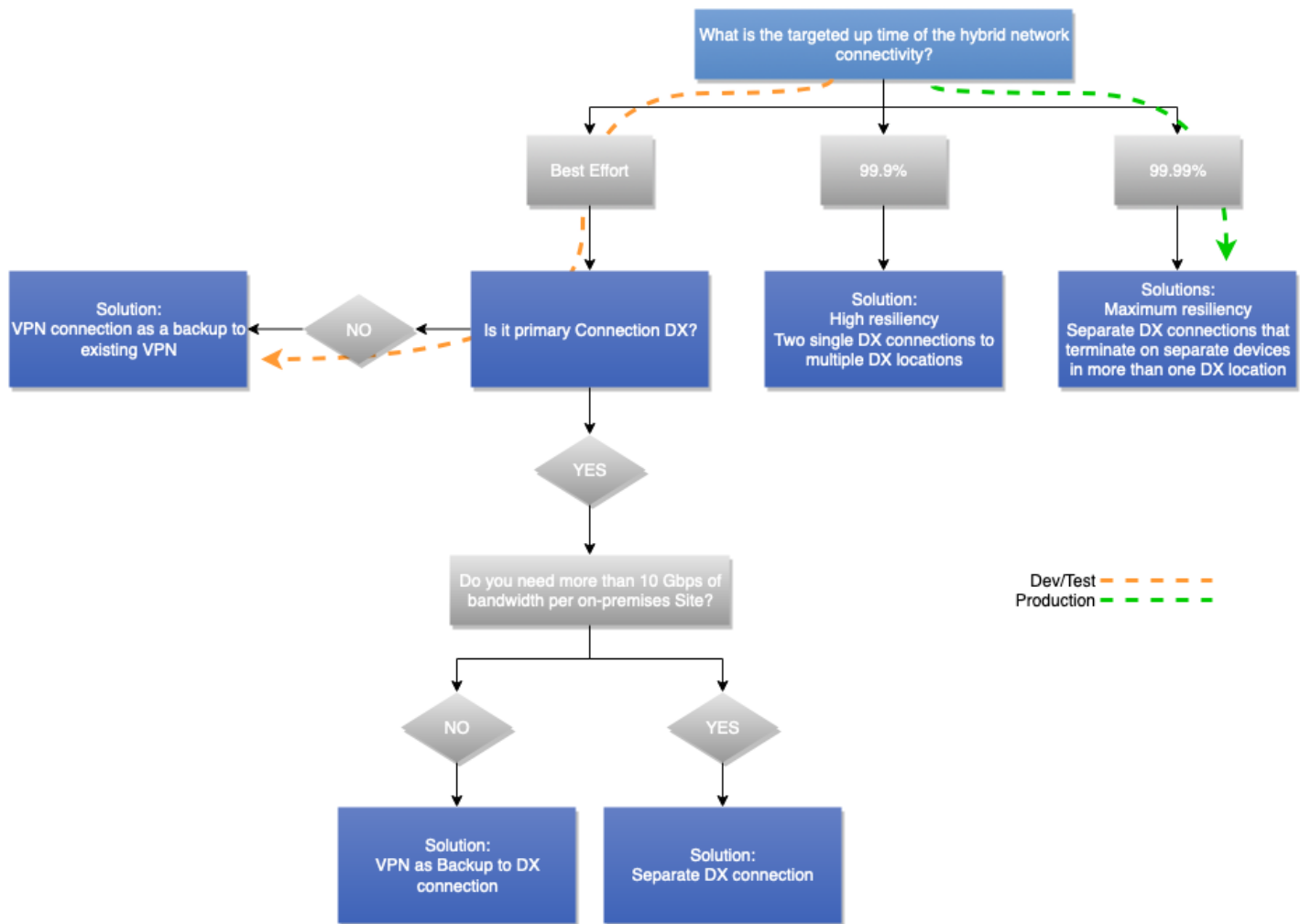


Figure 15 – Example Corp. Automotive reliability decision tree

Architecture selected by Example Corp. Automotive

The following diagram illustrates the architecture selected by Example Corp. Automotive after collecting the requirements and navigating the decision trees covered in the previous sections of this whitepaper.

It uses AWS S2S VPN over the internet terminating on AWS Transit Gateway for development and testing. It then uses AWS Direct Connect with Direct Connect gateway and a second AWS Transit Gateway for the production traffic. AWS Transit Gateway is used for Inter-VPC routing. From a data path perspective, the VPN tunnels for the primary data center are used as primary paths for development and testing, with the tunnels to the secondary data center used as failover paths. For the production traffic, all connections are used simultaneously. Traffic from AWS prefers the most

optional network connection based on the data center in which the on-premises system is located. Example Corp. Automotive uses similar route engineering techniques to prefer the appropriate path when traffic is sent to AWS ensuring symmetric traffic paths are used to minimize the use of the corporate network between on-premises primary and secondary data centers.

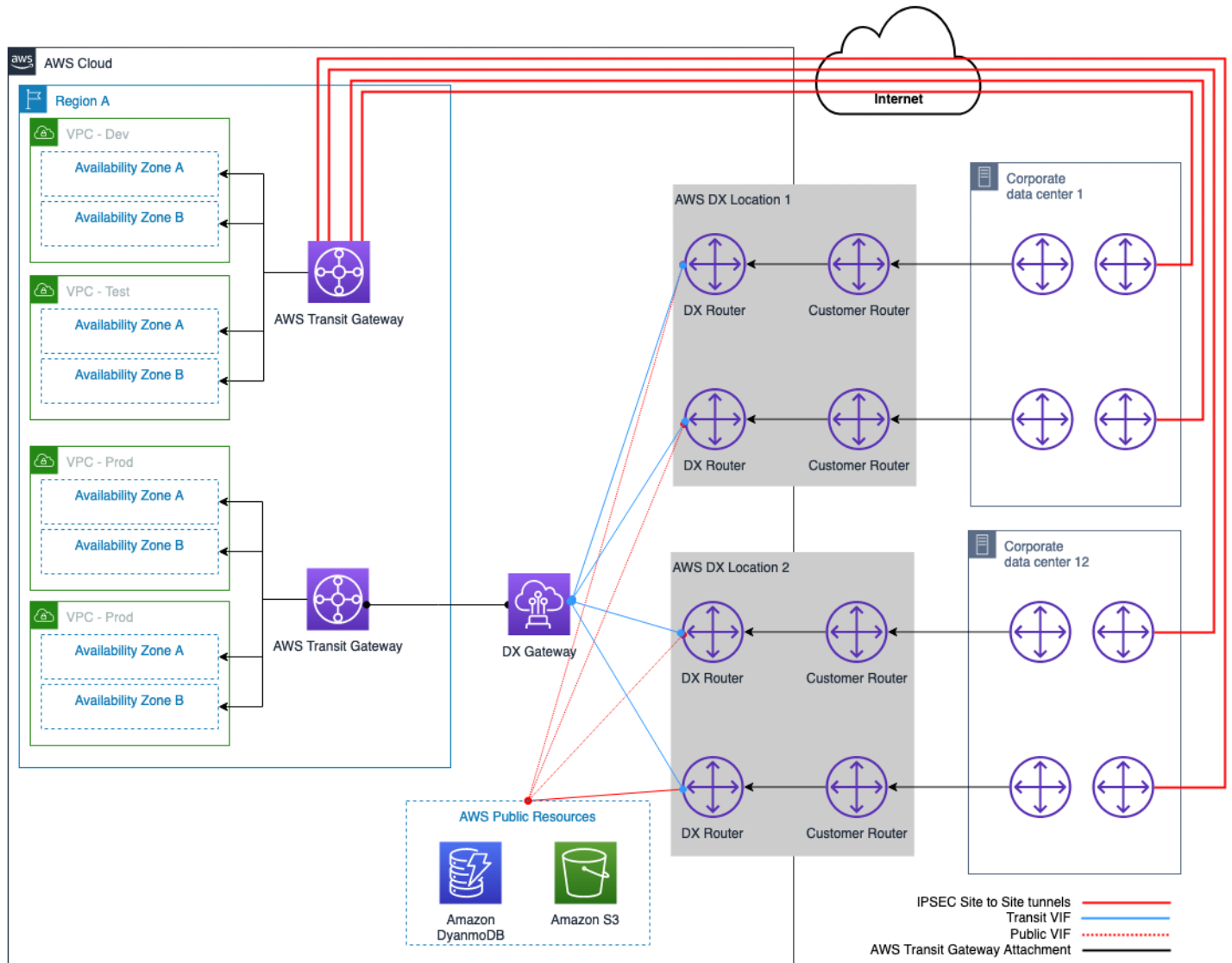


Figure 16 – Example Corp. Automotive selected hybrid connectivity model

Conclusion

A hybrid connectivity model is one of the fundamental starting points for the adoption of cloud computing. A hybrid network can be built with an optimal architecture following the connectivity model selection process outlined in this whitepaper.

The process consists of considerations arranged in a logical order. The order closely resembles a mental model followed by a seasoned network and cloud architects. Within each group of considerations, decision trees allow for rapid connectivity model selection, even with limited input requirements. You may find that some considerations and corresponding impacts point to different solutions. In those cases, as a decision maker, you may need to compromise on some requirements and select the most optimal solution that meets your business and technical requirements.

Contributors

Contributors to this document include:

- James Devine, Principal Solutions Architect, Amazon Web Services
- Andrew Gray, Principal Solutions Architect – Networking, Amazon Web Services
- Maks Khomutskyi, Senior Solutions Architect, Amazon Web Services
- Marwan Al Shawi, Solutions Architect, Amazon Web Services
- Santiago Freitas, Head of Technology, Amazon Web Services
- Evgeny Vaganov, Specialist Solutions Architect – Networking, Amazon Web Services
- Tom Adamski, Specialist Solutions Architect – Networking, Amazon Web Services
- Armstrong Onaiwu, Solutions Architect, Amazon Web Services

Further reading

- [Building a Scalable and Secure Multi-VPC AWS Network Infrastructure](#)
- [Hybrid Cloud DNS Options for Amazon VPC](#)
- [Amazon Virtual Private Cloud Connectivity Options](#)
- [Amazon Virtual Private Cloud Documentation](#)
- [AWS Direct Connect Documentation](#)
- [What's the difference between a hosted virtual interface \(VIF\) and a hosted connection?](#)

Document revisions

To be notified about updates to this whitepaper, subscribe to the RSS feed.

Change	Description	Date
Minor update	Updated to reflect DX quota limit increase.	July 10, 2023
Major update	Updated to incorporate latest best practices, services, and capabilities.	July 6, 2023
Minor update	Updated reference architecture diagrams to reflect changes in DX quota.	June 27, 2023
Minor update	Fixed broken links.	March 22, 2022
Initial publication	Whitepaper first published	September 22, 2020

Notices

Customers are responsible for making their own independent assessment of the information in this document. This document: (a) is for informational purposes only, (b) represents current AWS product offerings and practices, which are subject to change without notice, and (c) does not create any commitments or assurances from AWS and its affiliates, suppliers or licensors. AWS products or services are provided "as is" without warranties, representations, or conditions of any kind, whether express or implied. The responsibilities and liabilities of AWS to its customers are controlled by AWS agreements, and this document is not part of, nor does it modify, any agreement between AWS and its customers.

© 2023 Amazon Web Services, Inc. or its affiliates. All rights reserved.

AWS Glossary

For the latest AWS terminology, see the [AWS glossary](#) in the *AWS Glossary Reference*.