

The great Wi-Fi Architecture Debate:

cisco Live !

To Centralize or to Distribute - Where should your Network be?

Simone Arena
Distinguished Engineer, Cisco Wireless

Webex App

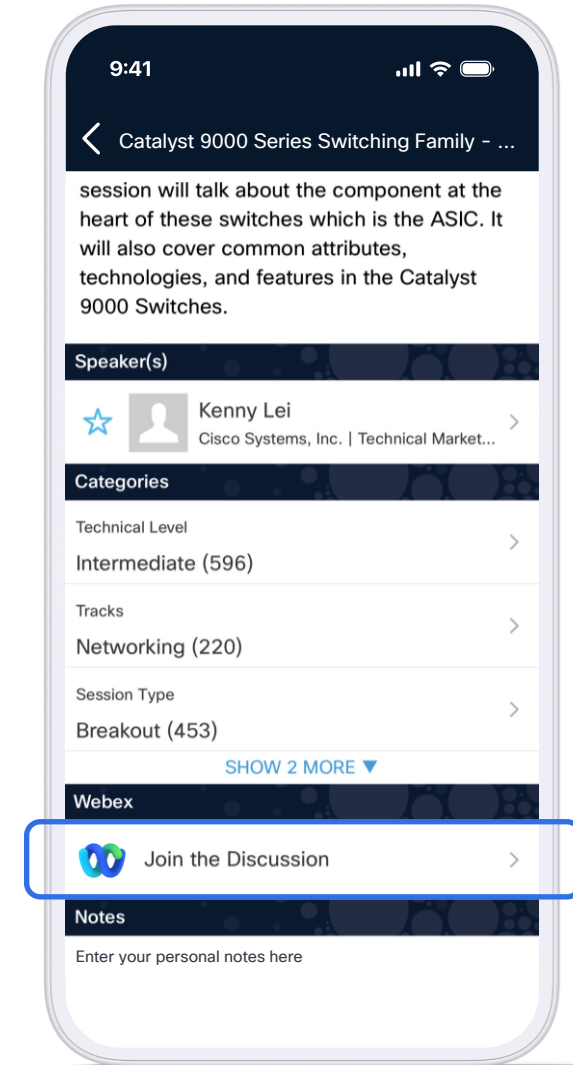
Questions?

Use Webex App to chat with the speaker after the session

How

- 1 Find this session in the Cisco Events App
- 2 Click “Join the Discussion”
- 3 Install the Webex App or go directly to the Webex space
- 4 Enter messages/questions in the Webex space

Webex spaces will be moderated by the speaker until February 27, 2026.



Agenda

- 01 **Intro: Cisco Wireless Strategy**
- 02 **Network Architecture Components**
- 03 **The Debate: Key Design Factors**
- 04 **Conclusions**

Cisco Wireless Strategy

Meeting our customers where they are



On-prem Managed

Use cases require
on-prem delivery.
DIY IT model



Cloud Monitoring

Need to retain control on
prem, cloud assisted. Use
cloud tools to help run
their networks



Cloud Managed

Prefer cloud-enabled
delivery for simplicity.
SaaS IT model

Deliver simplified outcomes and unified experience
to all customers: on cloud, on-prem or hybrid

Cisco Wireless Unified Platform

PROOF POINTS

Unify the Services

AI Services
(ex. AI RRM)


Policy


Analytics


APIs


User/Application Experience


AI RRM
ISE
DEVICE PROFILING

Unify the Management

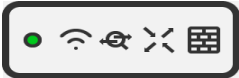


On prem Dashboard

Cloud Dashboard


UX ENGINE
DASHBOARD COLOR ☺
CatC GLOBAL VIEW

Unify the Infrastructure



Common Edge platform









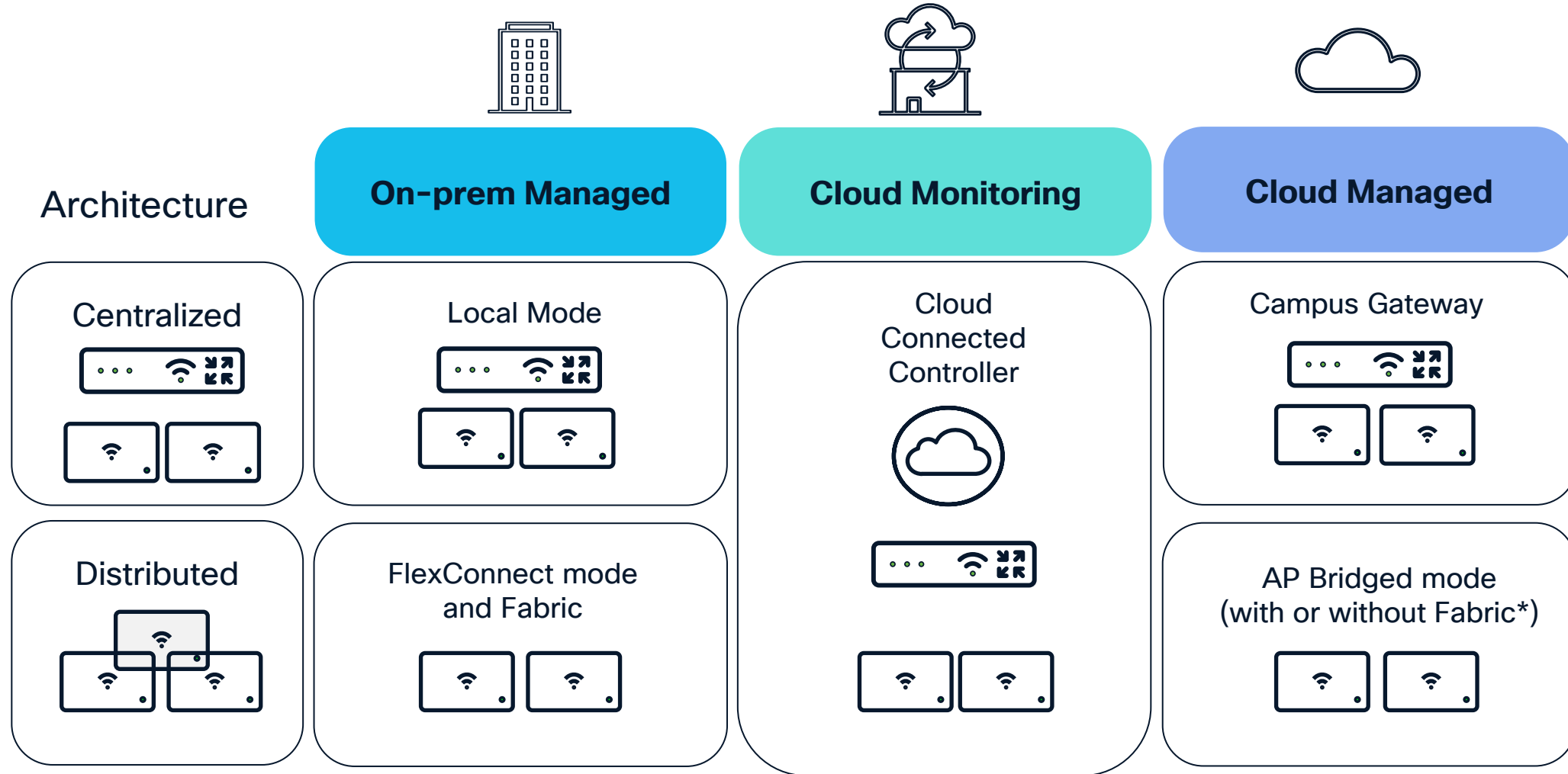
Common AP platform

ONE
HARDWARE

ONE LICENCE

ONE SUPPORT

Cisco Wireless Network Architecture Flexibility



* Announced at Cisco Live EMEA

Network Architecture Flexibility

=

You (as Network Designer/Architect) can now focus on the **requirements** and choose the right **Network Architecture**, independently of the Management platform, on-prem (Catalyst) or Cloud (Meraki)

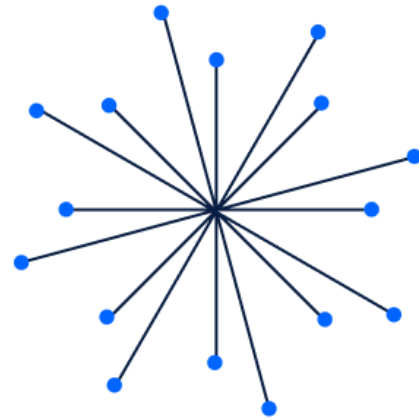
Centralize or Distribute? THAT is the Network Architecture Design question



Lafayette - Photo - London.

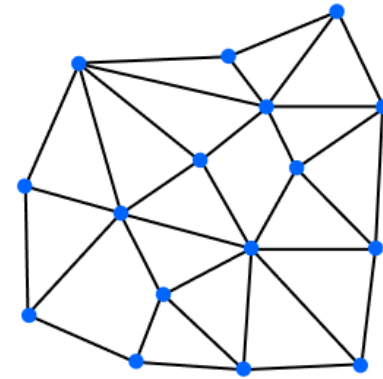
SARAH-BERNHARDT (HAMLET.)

What do you think about when you hear Distributed vs. Centralized Network Architecture?



Centralized

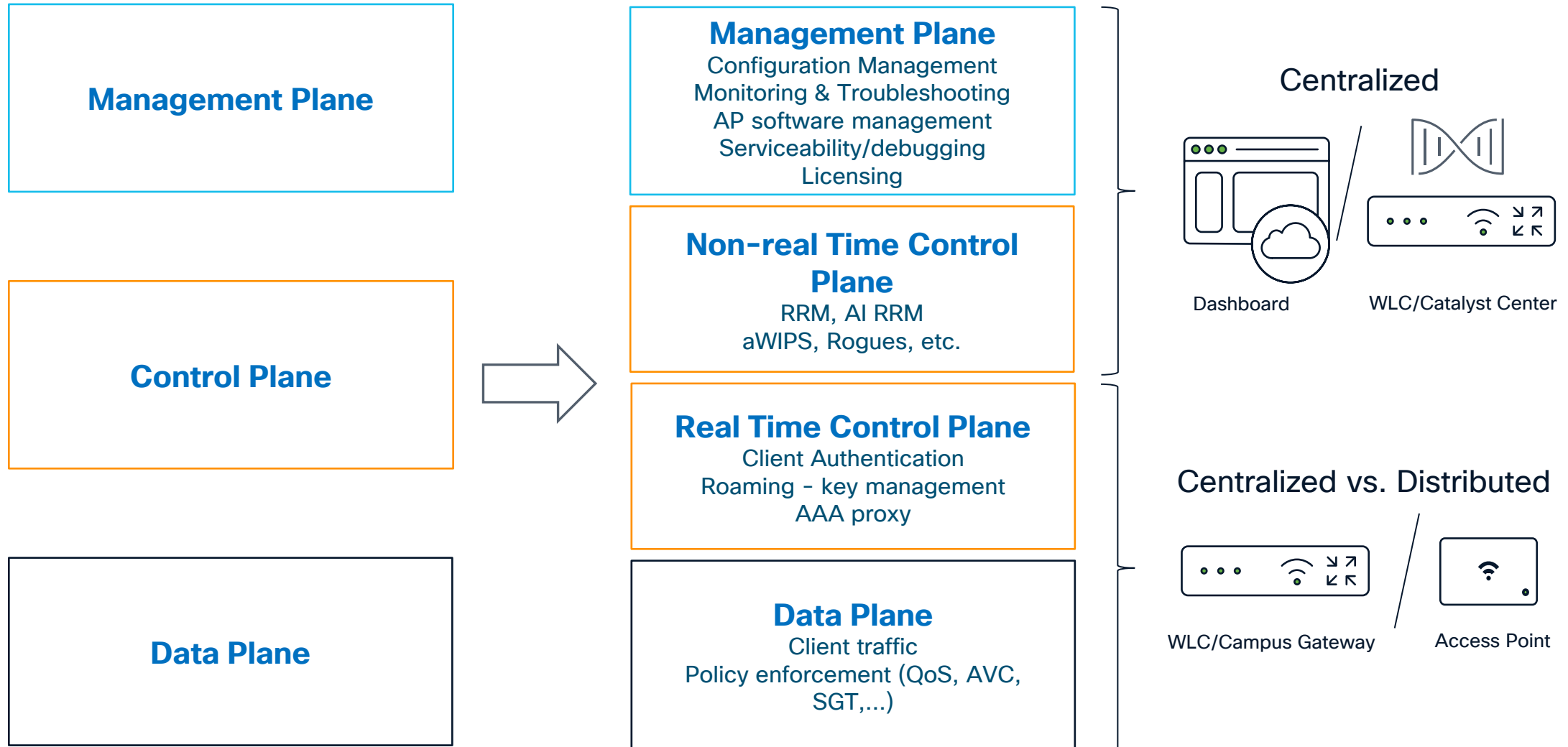
Vs.



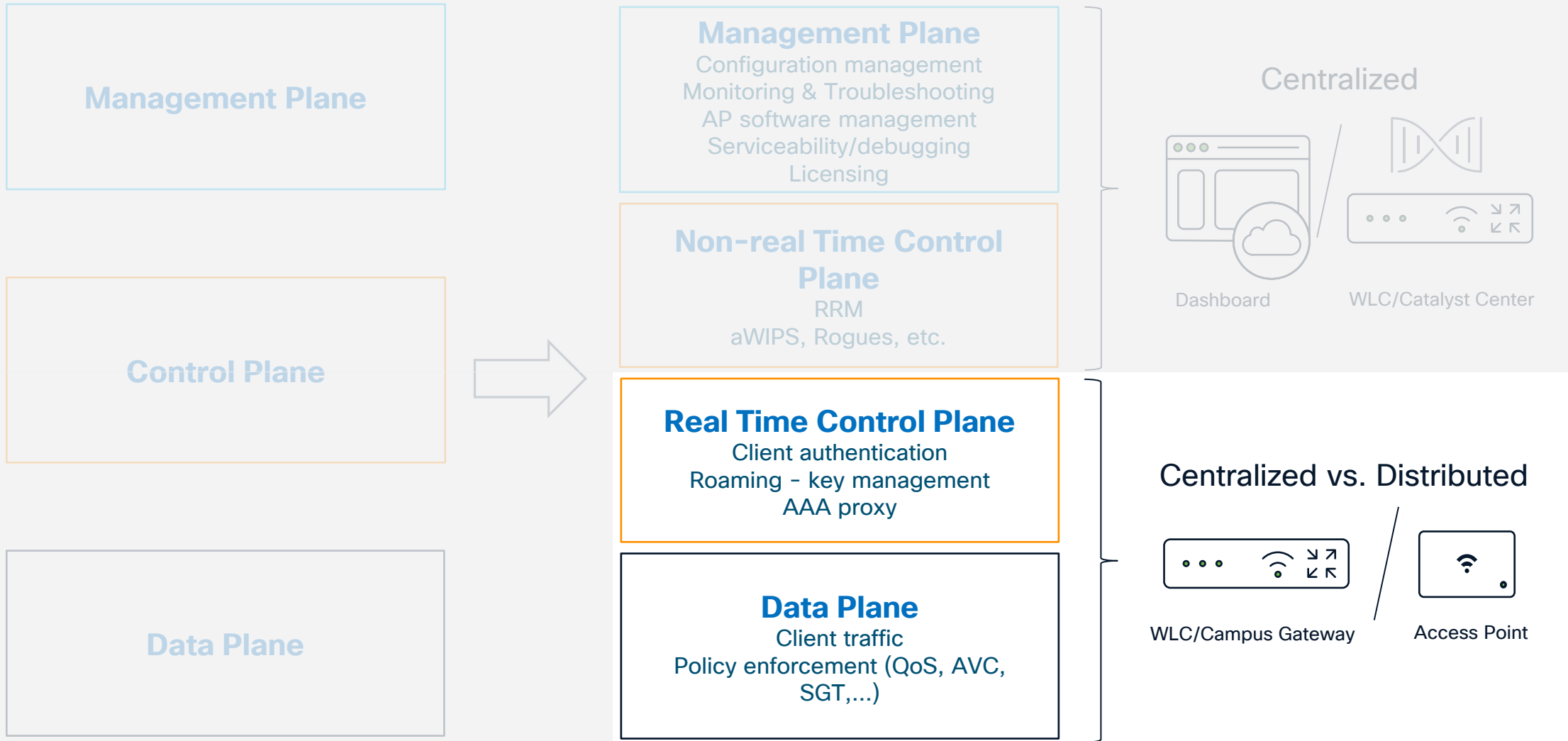
Distributed

**Do you think about where is the brain of your network?
The client traffic (data plane)? Or something else?**

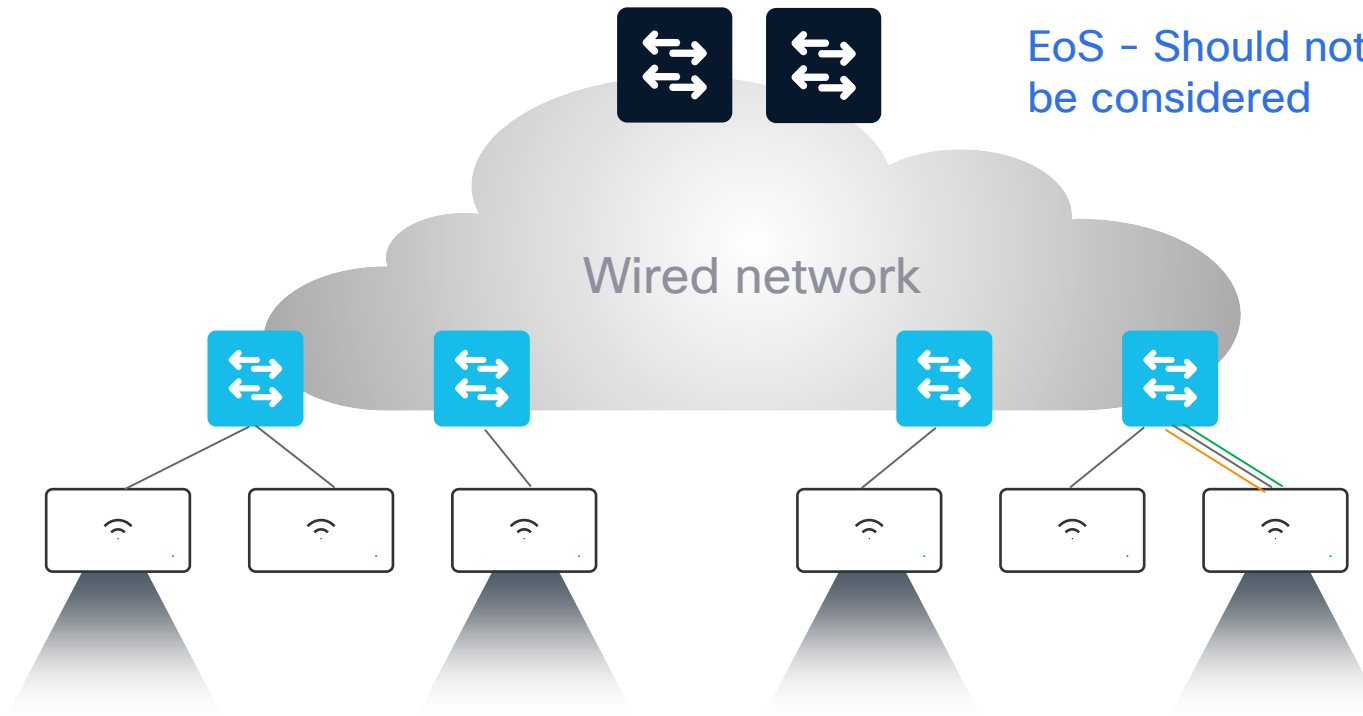
Network Architecture components



Network Architecture components > This session's focus



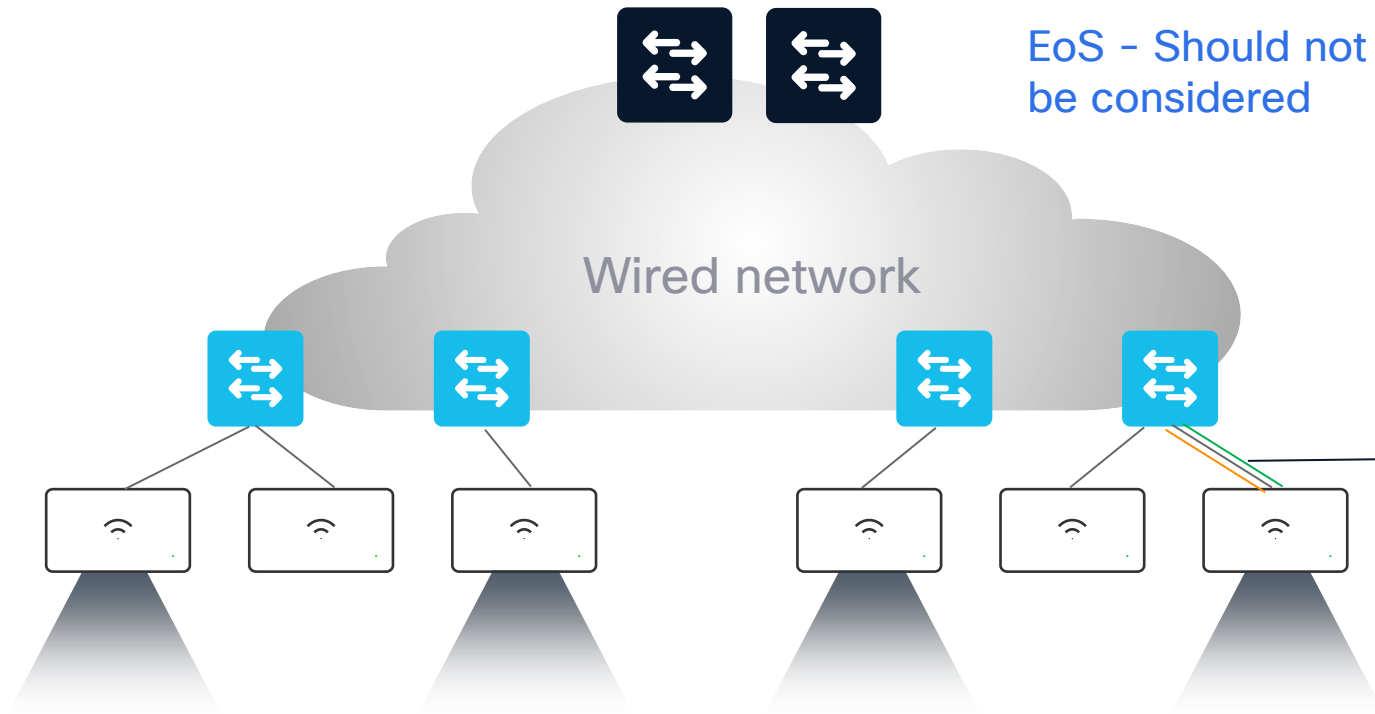
Deployment modes – Distributed



Wireless Deployment modes:

- Catalyst FlexConnect Local Switching
- Catalyst Fabric mode (SDA)
- ~~Embedded Wireless Controller on AP~~
- Meraki Bridge mode
- Meraki Fabric

Deployment modes – Distributed



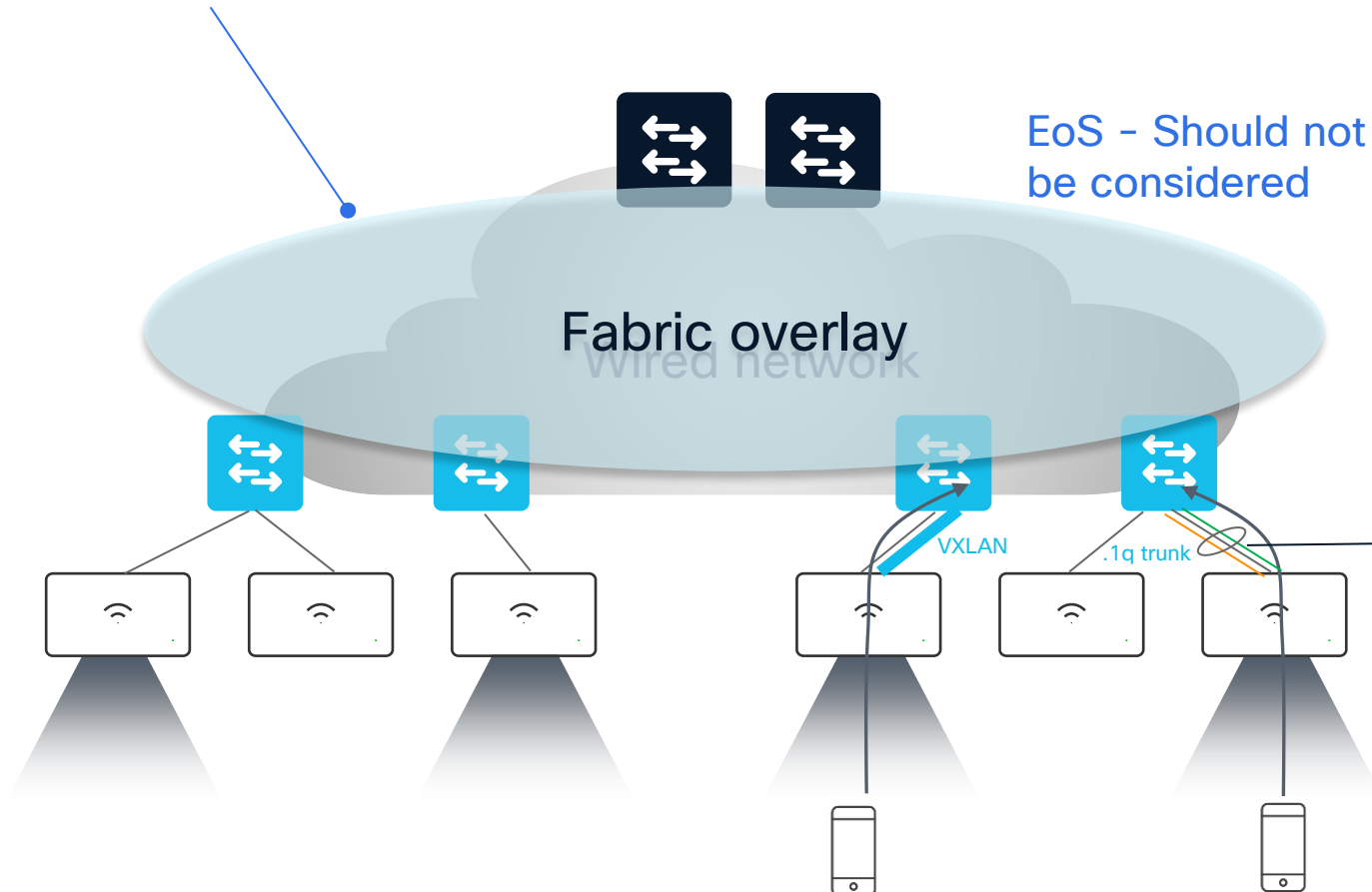
Wireless Deployment modes:

- Catalyst FlexConnect Local Switching
- Catalyst Fabric mode (SDA)
- ~~Embedded Wireless Controller on AP~~
- Meraki Bridge mode
- Meraki Fabric

- Client traffic is bridged at the AP
- SSID/User mapped to a VLAN at the AP

Deployment modes – Distributed

A Fabric overlay can be leveraged with or without wireless integration



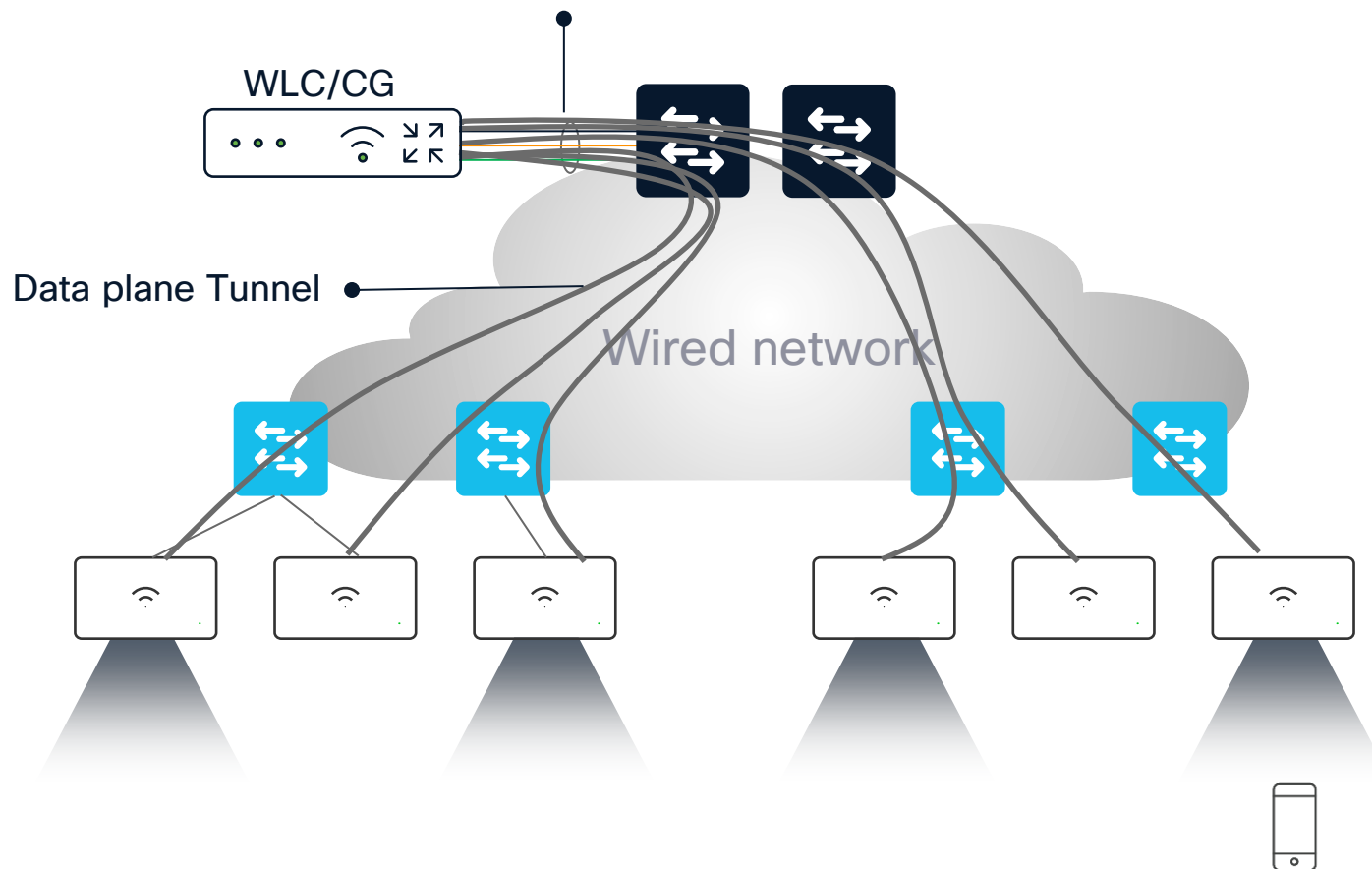
Wireless Deployment modes:

- Catalyst FlexConnect Local Switching
- Catalyst Fabric mode (SDA)
- ~~Embedded Wireless Controller on AP~~
- Meraki Bridge mode
- Meraki Fabric

- Client traffic is bridged at the AP
- SSID/User mapped to a VLAN at the AP
- VLAN (and client policy) carried to switch via uplink 802.1q trunk or VXLAN tunnel

Deployment modes – Centralized

- Client traffic is tunneled to a centralized aggregator where it is bridged on to the wired network
- SSIDs/users mapped to centralized VLANs
- WLC/Campus Gateway uplinks configured as .1q trunk

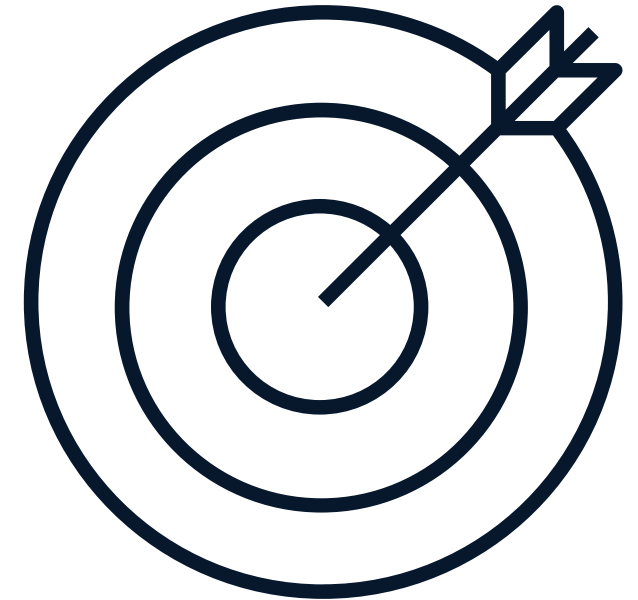


Wireless Deployment Modes:

- Catalyst Local mode
- Catalyst FlexConnect Central Switching (CAPWAP or EoGRE)
- Meraki Tunnel mode (Campus Gateway, ~~MX~~, EoGRE)
- Note: MX is not recommended for centralized design at scale, use Campus Gateway

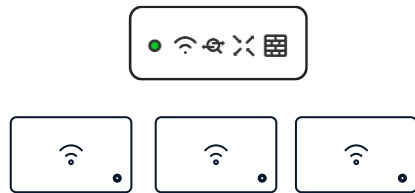
Goal of this session

Identify the factors that may drive the **Network Architecture** choice (distribute vs. centralize) and analyze the pros and cons of the options



Sometimes the choice is easy...

Large, Medium Campus
> Centralized architecture



km^2/mi^2

E.g., University Campus



Distribute Enterprise: Branch
> Distributed architecture



m^2/ft^2

e.g., Retail





**A centralized data plane model allow
me to have one point of ingress into
my wired network for all wireless
traffic, simplifying my deployment**

US based University
Director of IT



I don't like hair-pinning client traffic, the WLC might become a bottleneck and a single point of failure.

Italy based University
Director of IT

So how to go about it?

...Key design factors

Key Design Factors



Features



Services



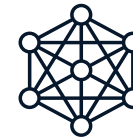
Wired Network Design



Scale/Performances



Policy



Use Cases

**Wait...What about
Management?**

Management/IT Operations

Cloud Managed or On-prem managed?

Is the Management responsibility of one Team? Or multiple regions have different different teams?

Are wireless and wired networks managed by separated IT teams?

Management/IT Operations

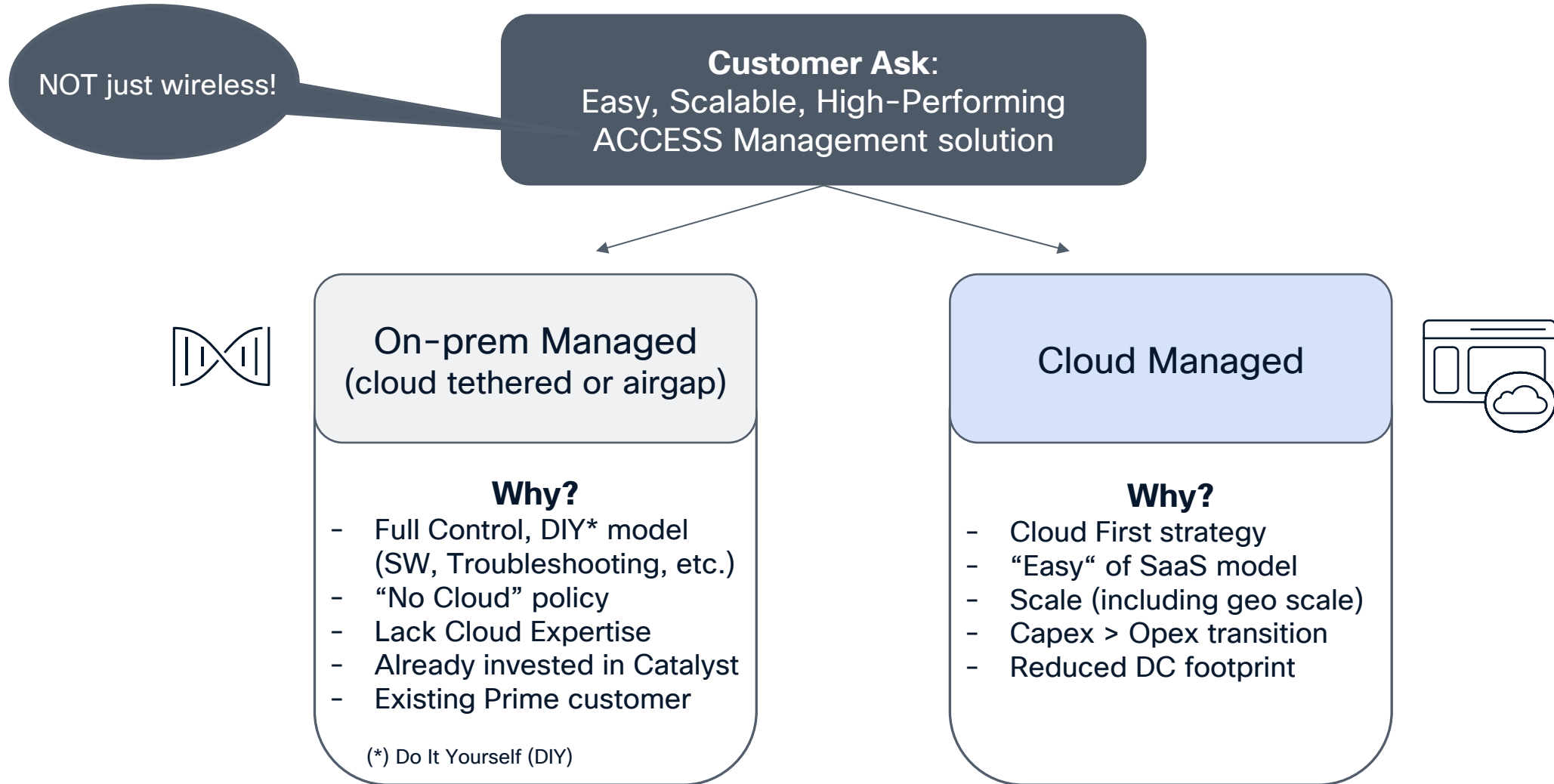
Cloud Managed or On-prem managed?

Is the Management responsibility of one Team? Or multiple regions have different different teams?

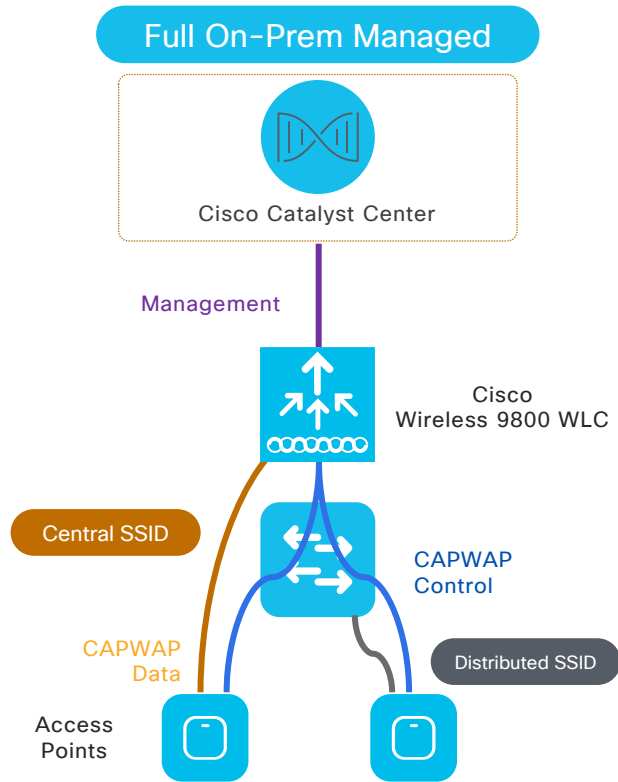
Are wireless and wired networks managed by separated IT teams?

Do we agree the Management choice between On-prem vs Cloud, should be orthogonal to the Network Architecture choice?

Management/IT Operations

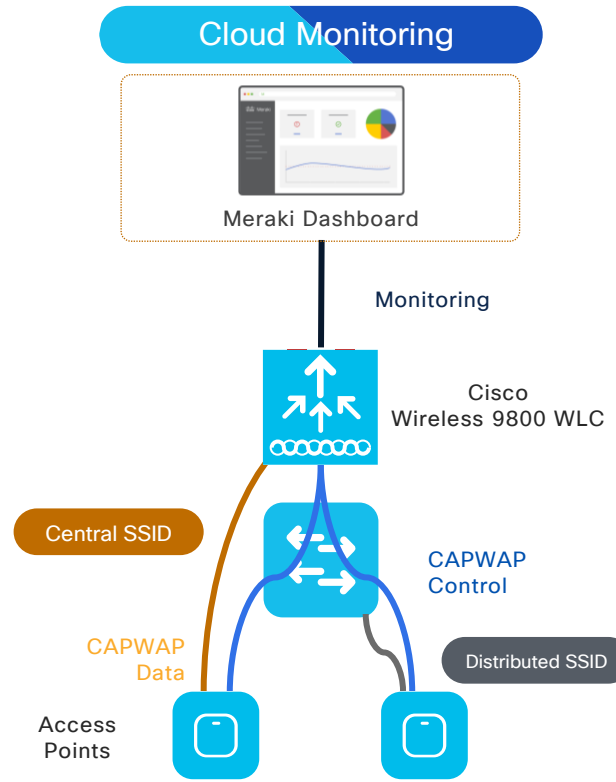


Management/IT Operations



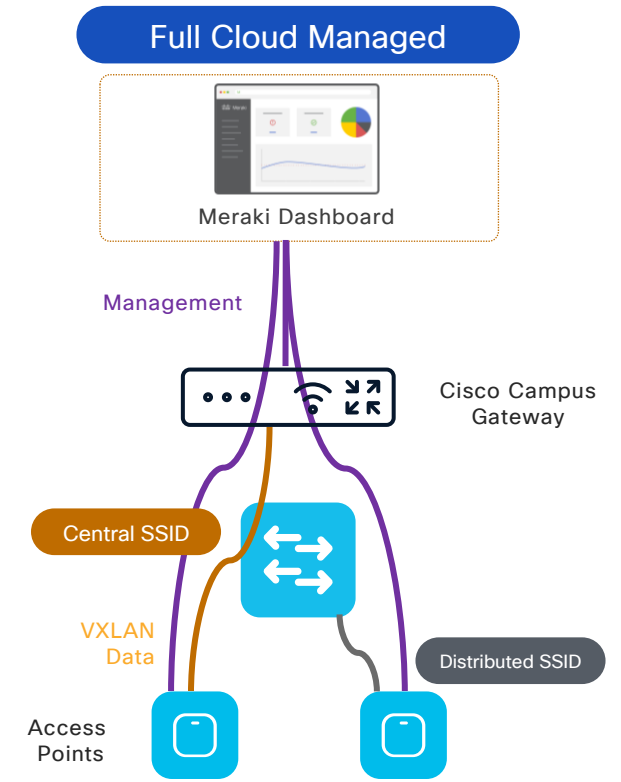
On-Prem Management

- **DIY model** – **Customer controls everything**, from configuration, to software upgrades, to troubleshooting
- Customer may have data protection/retention policies that prevent them to adopt a public cloud model
- All Cisco 18xx/28xx/38xx/CW91xx/CW917x APs supported



Cloud Monitoring

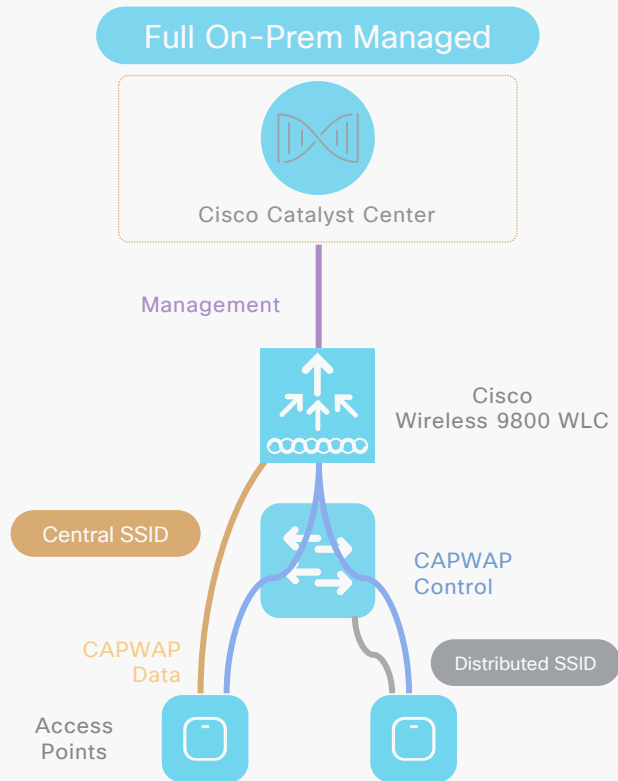
- **Customer wants to adopt cloud but not ready for full management yet** (e.g., older APs or not ready for SaaS).
- Customer wants benefits of DIY but want centralized management to monitor their deployment
- Customer does not have CatC
- Cisco 28xx/38xx/CW91xx/CW917x APs supported



Cloud Management

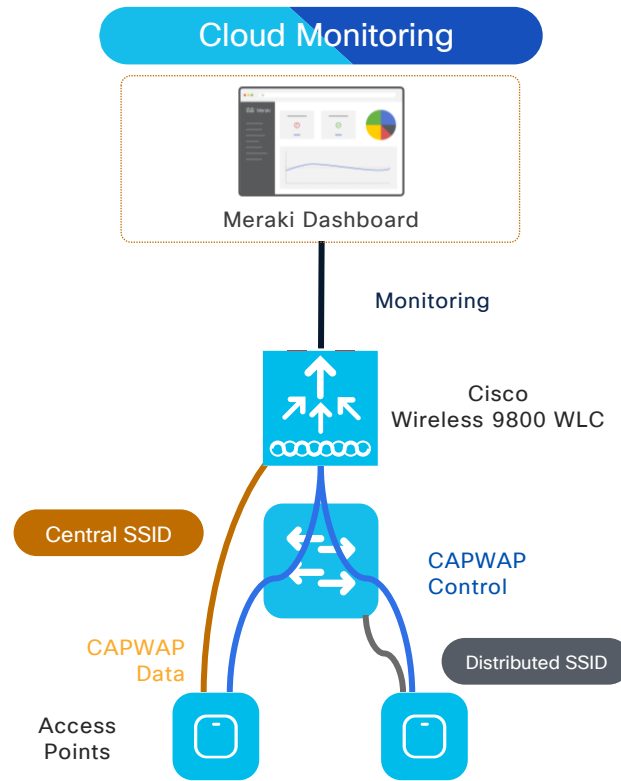
- **Cloud First customer**
- Customer wants a SaaS model
- Customer is on-prem but ready to move to cloud (e.g., has new Access Points or ready to refresh)
- Cisco CW916x/CW917x APs supported

Management/IT Operations



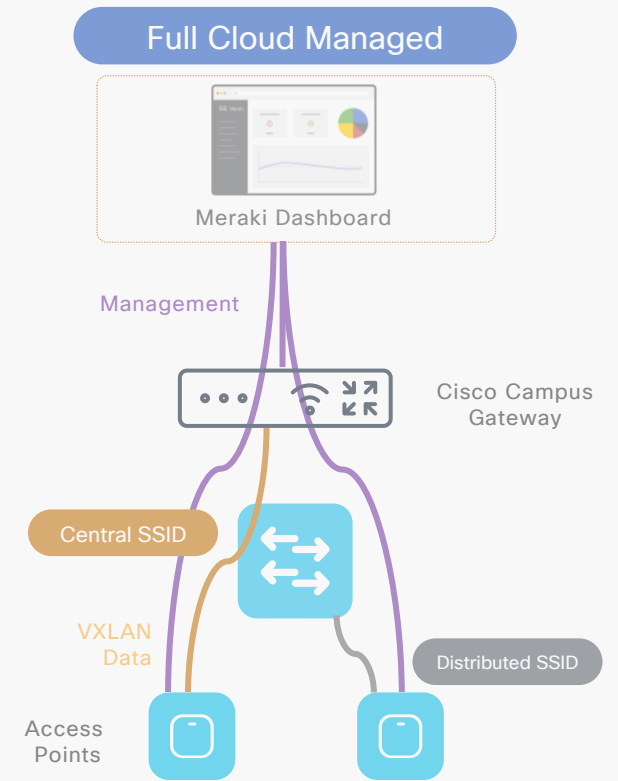
On-Prem Management

- **DIY model** – **Customer controls everything**, from configuration, to software upgrades, to troubleshooting
- Customer may have data protection/retention policies that prevent them to adopt a public cloud model
- All Cisco 18xx/28xx/38xx/CW91xx/CW917x APs supported



Cloud Monitoring

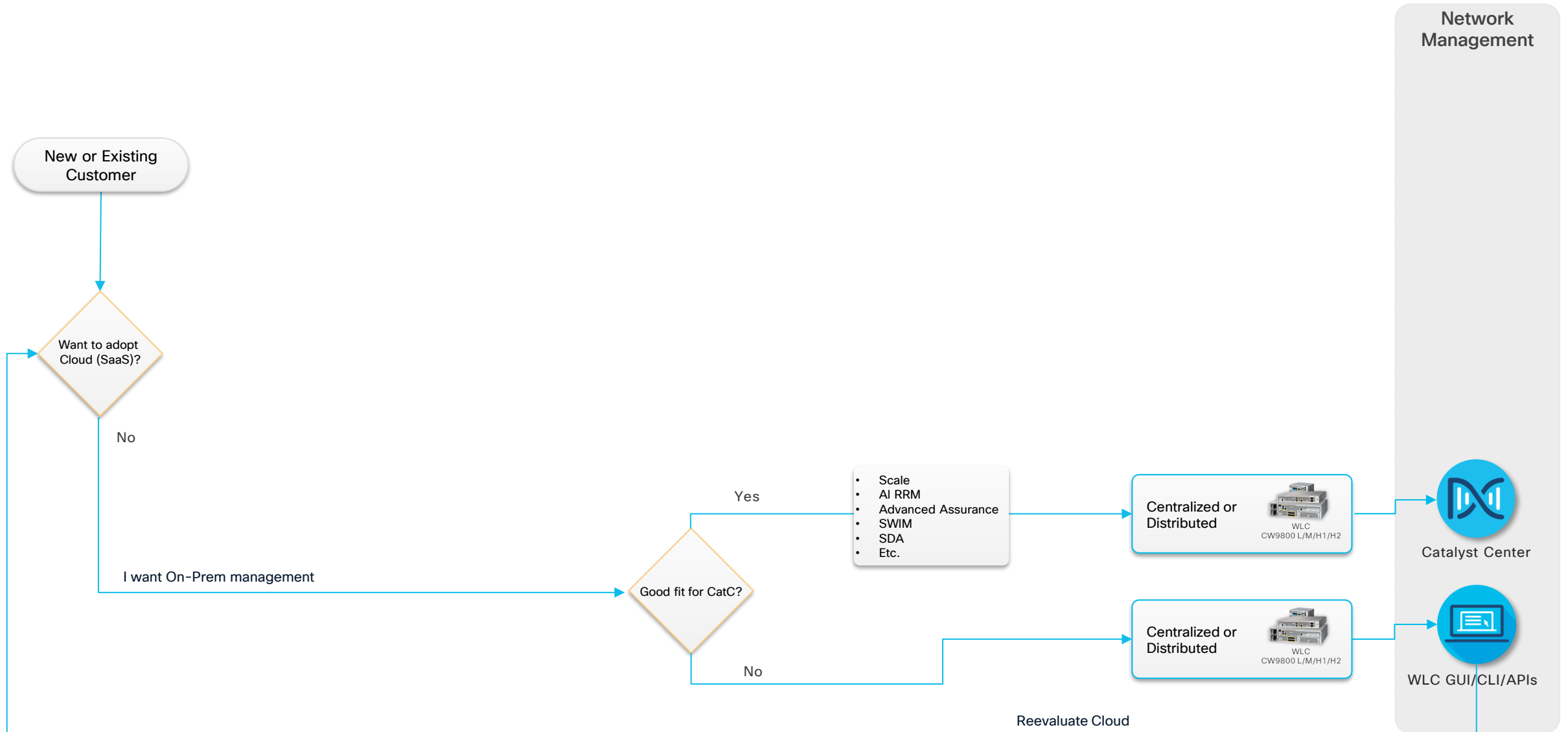
- **Customer wants to adopt cloud but not ready for full management yet** (e.g., older APs or not ready for SaaS).
- Customer wants benefits of DIY but want centralized management to monitor their deployment
- Customer does not have CatC
- Cisco 28xx/38xx/CW91xx/CW917x APs supported



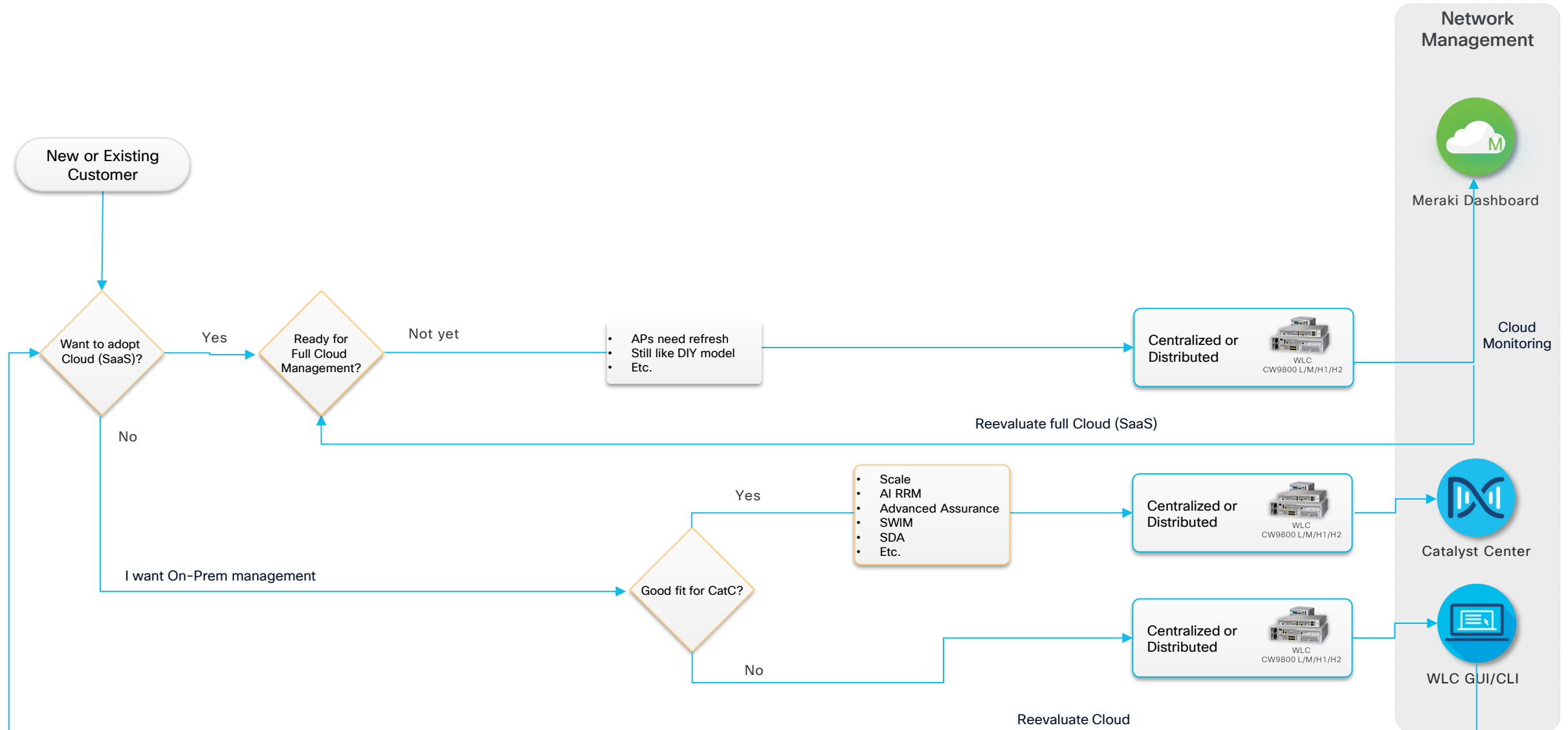
Cloud Management

- **Cloud First customer**
- Customer wants a SaaS model
- Customer is on-prem but ready to move to cloud (e.g., has new Access Points or ready to refresh)
- Cisco CW916x/CW917x APs supported

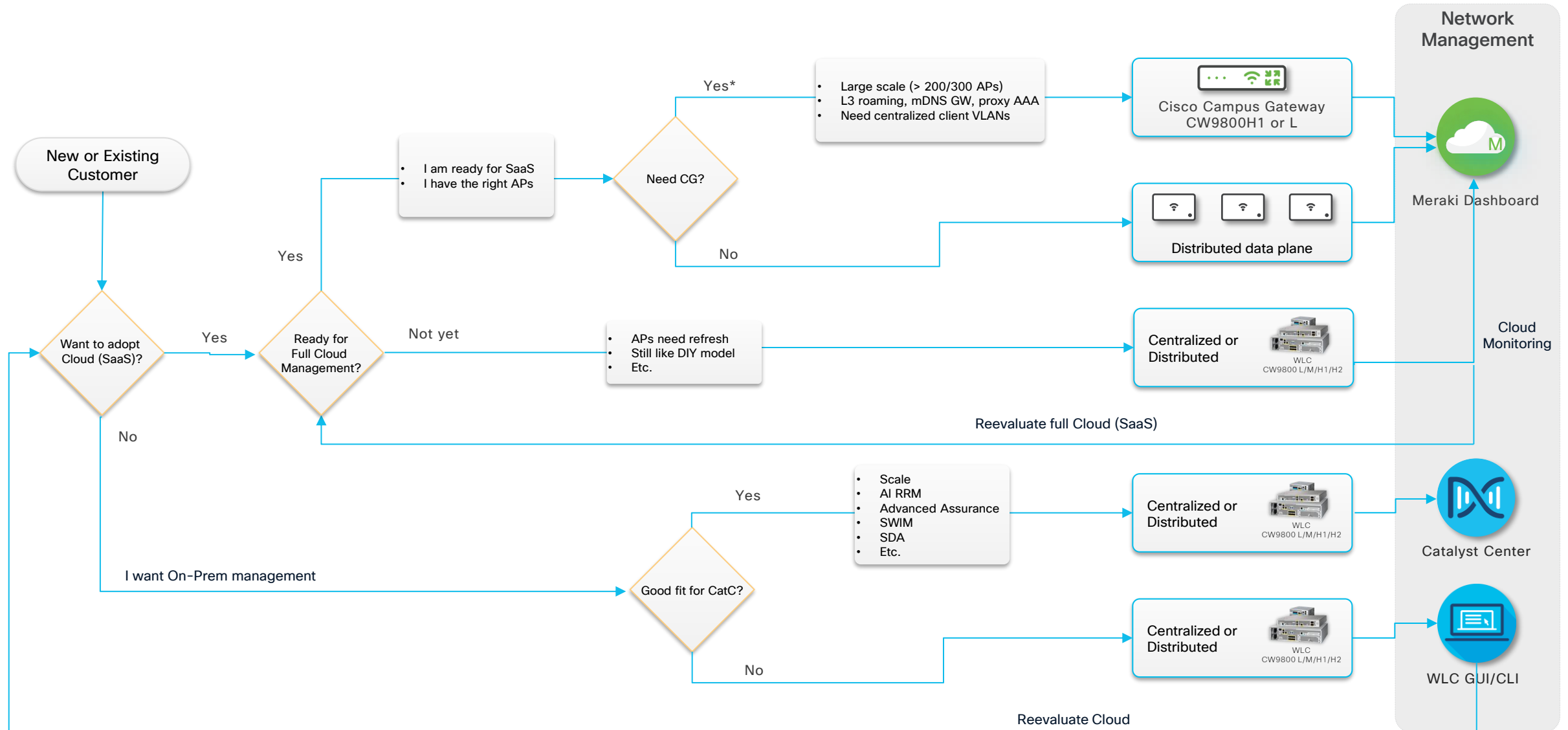
Management/IT Operations: Your options



Management/IT Operations: Your options



Management/IT Operations: Your options



* Customers with CW9800 will be given Campus Gateway device as a swap to move to cloud

BRKEWN-2002

...back to the Design Factors



Features



Services



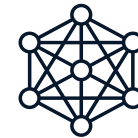
Wired Network Design



Scale/Performances



Policy



Use Cases

Services

Services – Roaming, DHCP, VLAN/Subnetting, mDNS

Seamless Roaming domain's size,

How to manage DHCP with distributed networks?

How to manage mDNS at scale?

Seamless Roaming: Distribute or Centralize?

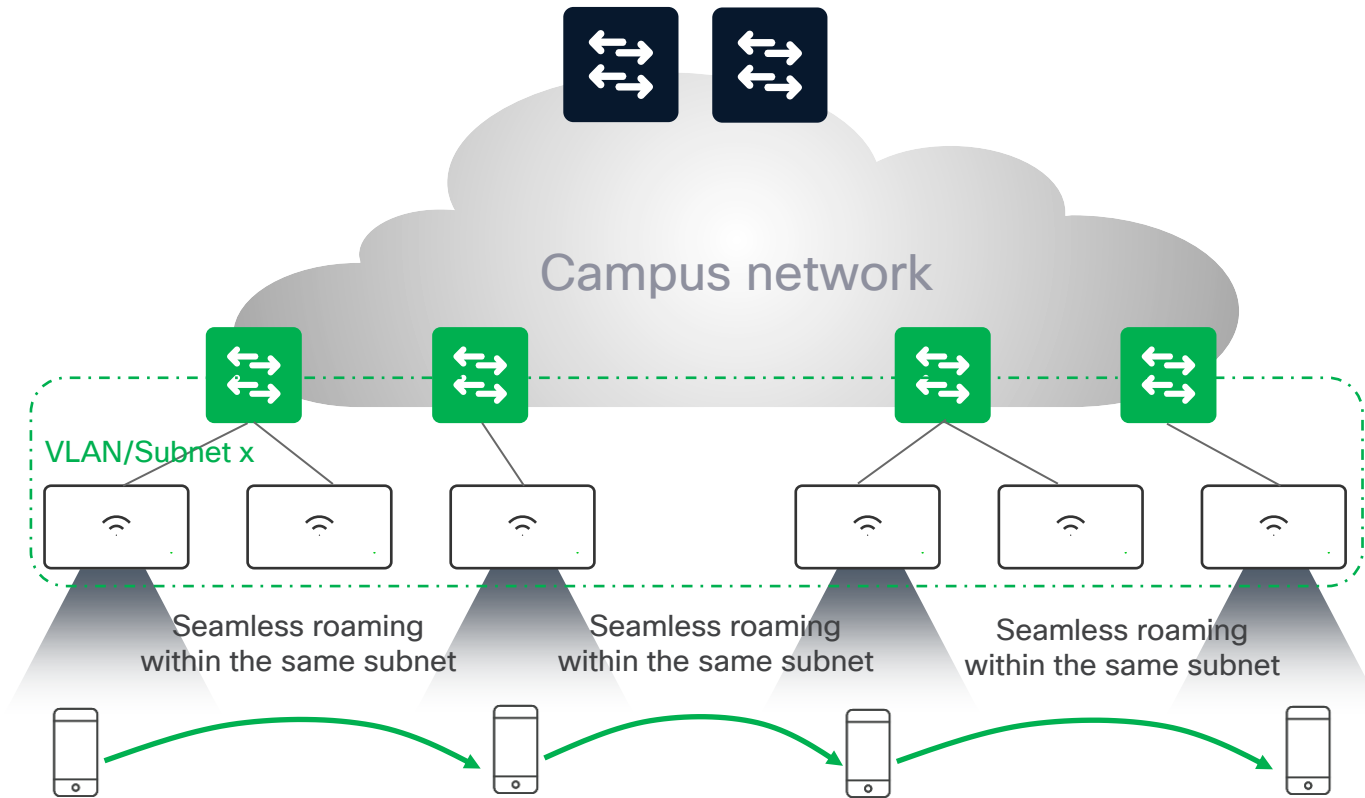
Seamless roaming domain: keep same IP and policy as client roams > same VLAN/L2 broadcast domain

How can you create a L2 roaming domain?

- Spanning the VLAN everywhere

Q: How big can be the L2 roaming domain?

- Need to consider the type of layer 2 and Layer 3 switches and their MAC/ARP tables size, the impact of spanning tree (SPT), the DHCP scope design, etc.



How big can be the L2 roaming domain?

Let's start with client considerations:



A **Single Dual Stack Host** will have **1 x IPv4** address, and
at least 3* x IPv6 Addresses
(IPv4 Unicast, IPv6 Link Local, IPv6 Unique Local, IPv6 Global Unicast)

Windows 11: up to 16 IPv6 IP addresses (!!)

* Not a guideline, just a realistic assumption

How big can be the L2 roaming domain?

Let's identify a roaming domain size in terms the number of APs:

Layer 2 switch = Catalyst 9200L (or MS equivalent)

Table



MAC addresses table limit: 16,000

- Each wireless device will take a MAC address entry
- If we consider Random MAC, this number can be higher
- If we assume 40* clients per AP > $16,000/40 = \text{max } 400 \text{ APs per L2 roaming domain}$

* Not a guideline, just a realistic assumption

How big can be the L2 roaming domain?

Let's identify a roaming domain size in terms the number of APs:

Layer 3 switch = Catalyst 9300 (or MS equivalent)

Table

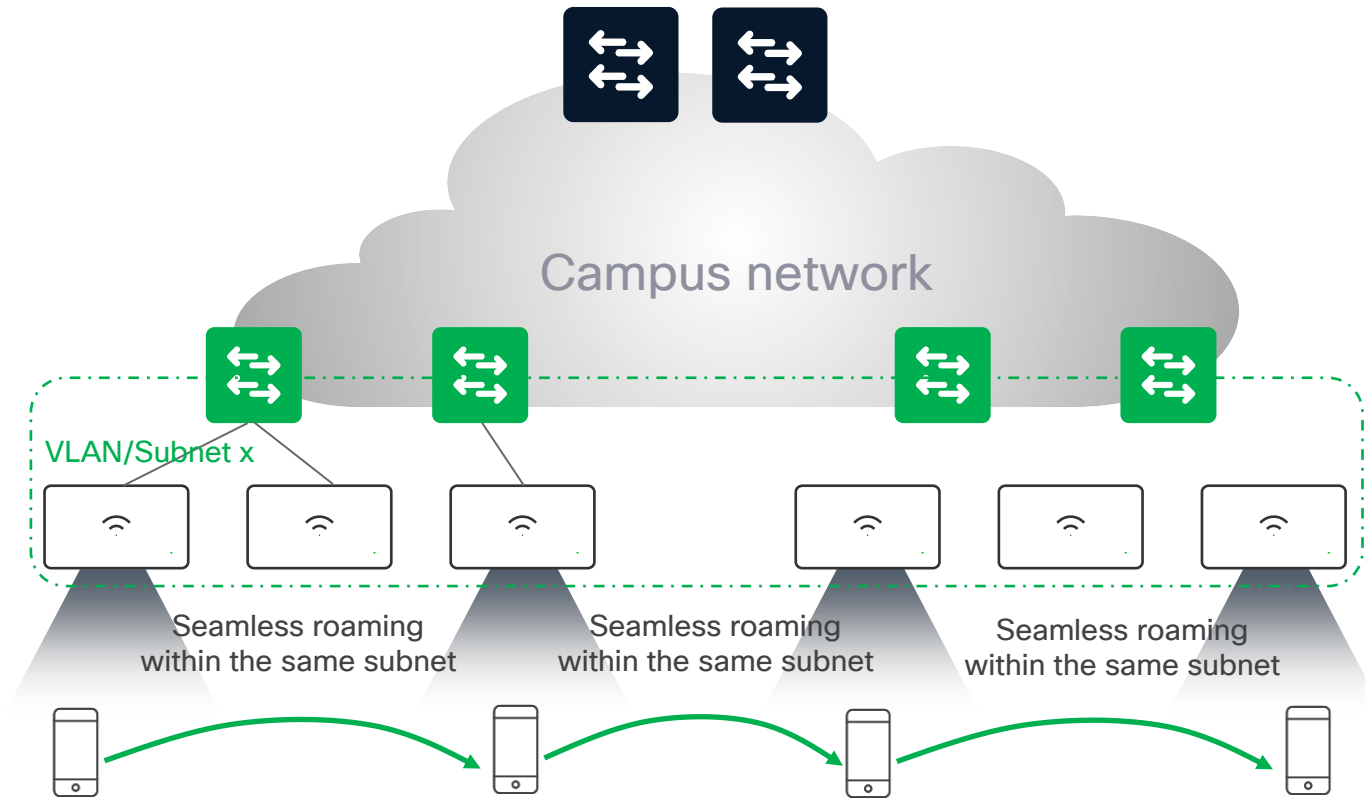


ARP entries: 32,000

- For dual stack clients, the scale numbers are divided by at least 4 (one entry for IPV4 and three entries for IPv6) > For C9300 the max number of clients is $32k/4 = 8k$
- If we assume 40* clients per AP > $8,000/40 = \text{max } 200 \text{ APs L2 roaming domain}$

* Not a guideline, just a realistic assumption

Seamless Roaming: Distribute or Centralize?



Seamless roaming domain: keep same IP and policy as client roams > same VLAN/L2 broadcast domain

How can you create a L2 roaming domain?

- Spanning VLAN everywhere

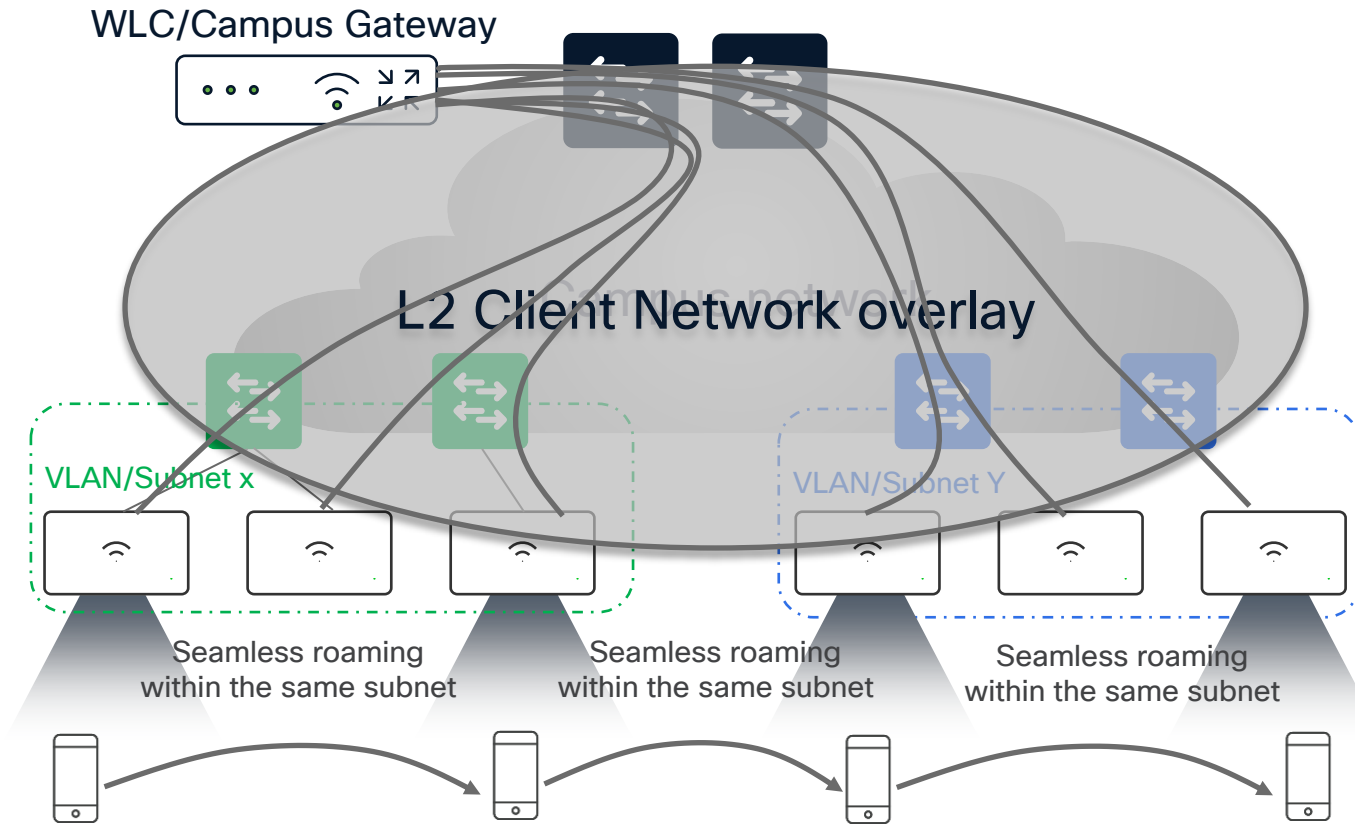
Q: How big can be the L2 roaming domain?

- Recommendation: 200 APs, 8k clients. Anyway, there is a limit...
- If you consider Fast roaming, and Flex Local switching: max 300 APs in a roaming domain (site tag)

Q: How to scale more?

- You need seamless roaming across L3 boundaries > create L2 overlay

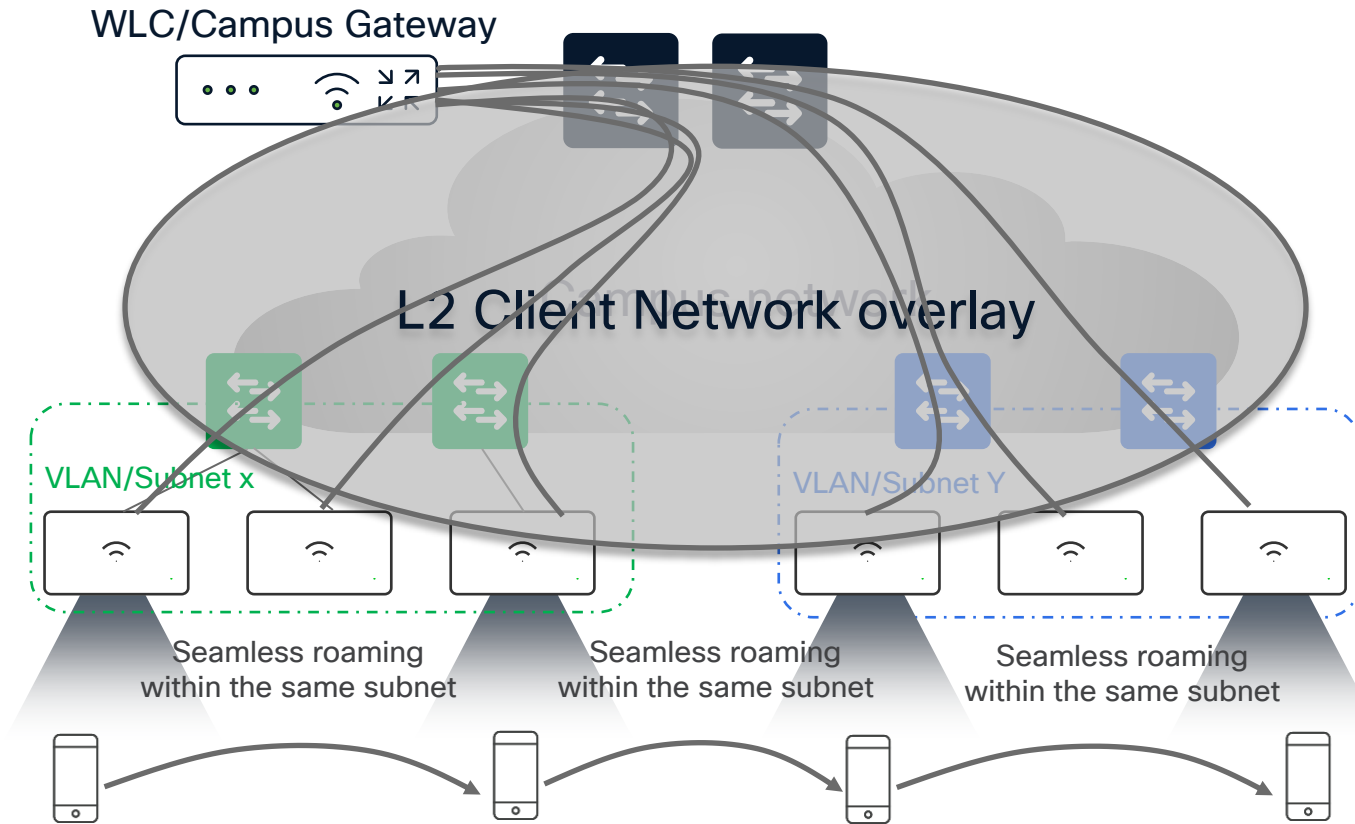
Seamless Roaming: Distribute or Centralize?



How can you create the L2 overlay?

- Via Tunneling (VXLAN/CAPWAP) to a Centralized Aggregator/Edge/Concentrator

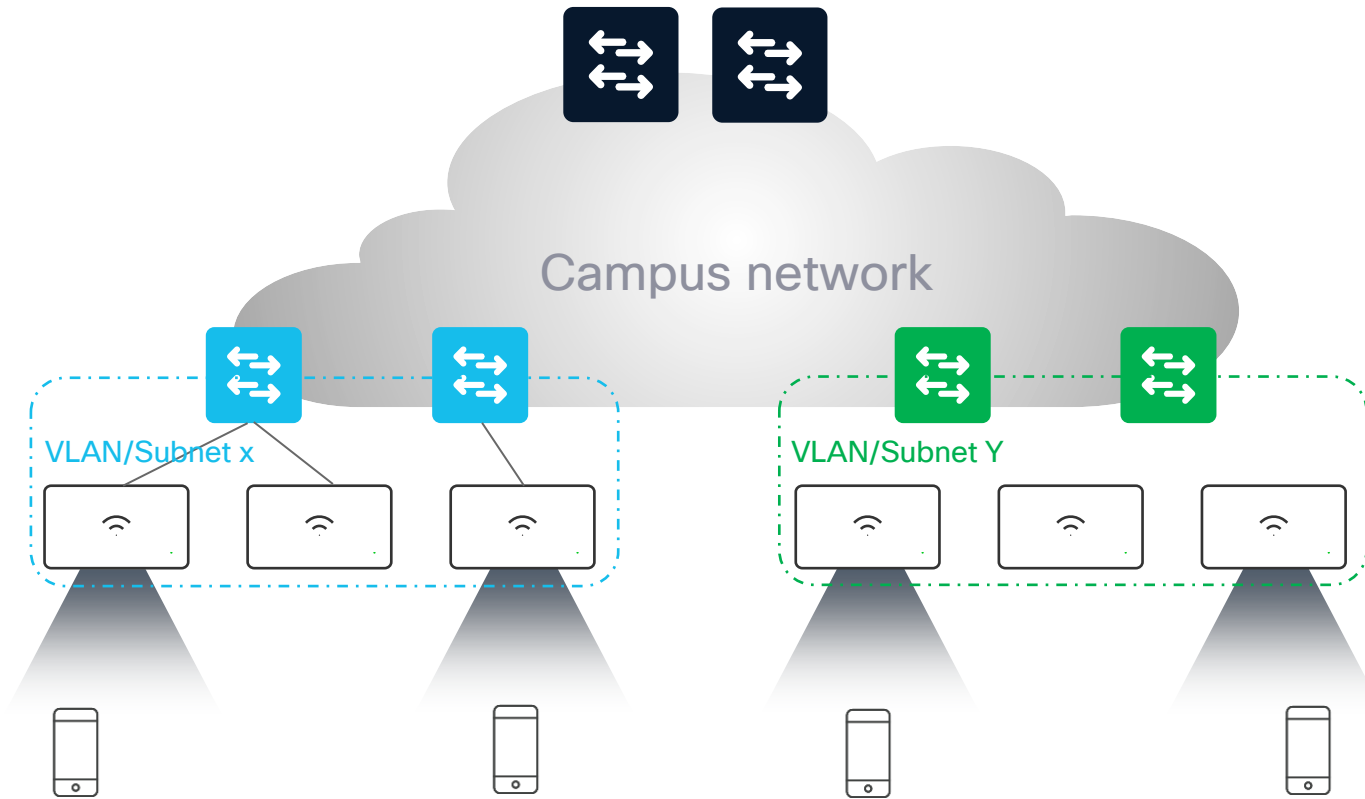
Seamless Roaming: Distribute or Centralize?



How can you create the L2 overlay?

- Via Tunneling (VXLAN/CAPWAP) to a Centralized Aggregator/Edge/Concentrator
- Traffic is centralized, client VLANs exist only at the distribution/core switch
- No need to touch the existing underlay network > great for faster brownfield adoption
- Sizing and scale of the distribution/core switches needs to be considered

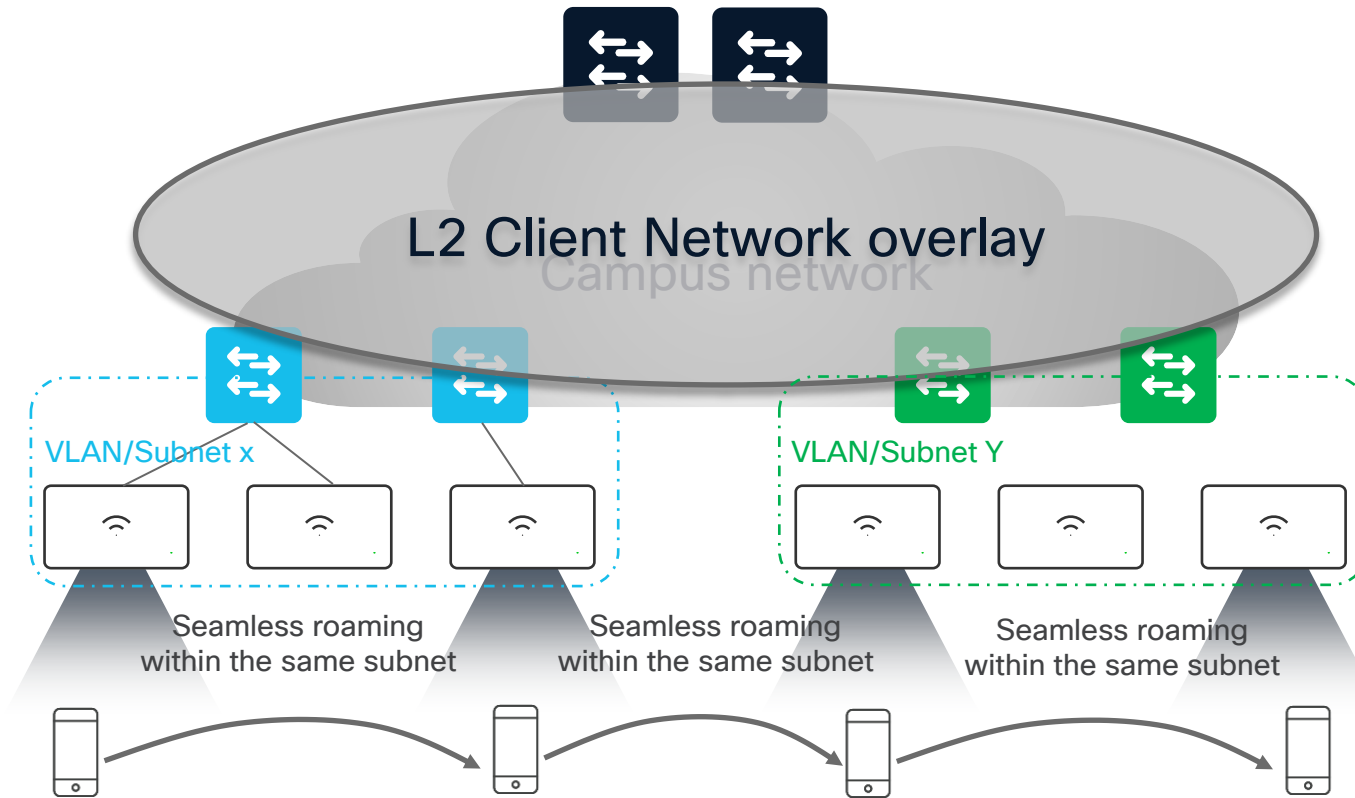
Seamless Roaming: Distribute or Centralize?



How can you create the L2 overlay?

- Via Tunneling (VXLAN/CAPWAP) to a Centralized Aggregator/Edge/Concentrator
- Wired Fabric

Seamless Roaming: Distribute or Centralize?



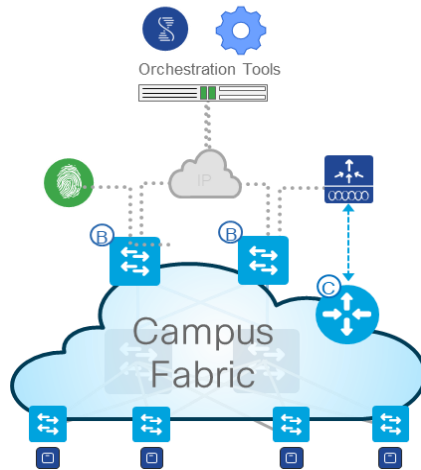
How can you create the L2 overlay?

- Via Tunneling (VXLAN/CAPWAP) to a Centralized Aggregator/Edge/Concentrator
- Wired Fabric
 - Need to create a L2 overlay on top of the underlay campus network
 - APs can be in bridge/flex mode with distributed data plane
- How? On-prem SDA Fabric, or the new Cloud Managed Campus Fabric

Campus Fabrics

The Strategic Foundation for Cisco Enterprise Networks

What is a Fabric



Campus

Abstracts the physical network into a unified, secure overlay delivering identity, segmentation, policy at scale.

Automation

Simplifies operations, ensuring consistency, speed, resilience.

Fabric Value



Operational simplicity

Network simplification, Unified wired & wireless



Secure networking

Identity-based segmentation, multi-domain policy



AI-driven Assurance

Network insights & Analytics

IOS XE Infrastructure | Catalyst Center / Meraki Dashboard | Identity Services Engine / Access Manager

Cisco Cloud-Managed Fabric

Multi-Network, Cloud-Managed EVPN-VXLAN Fabric



Simplified management as a single logical entity



Focused on **brownfield migration**



Flexible deployment and staging



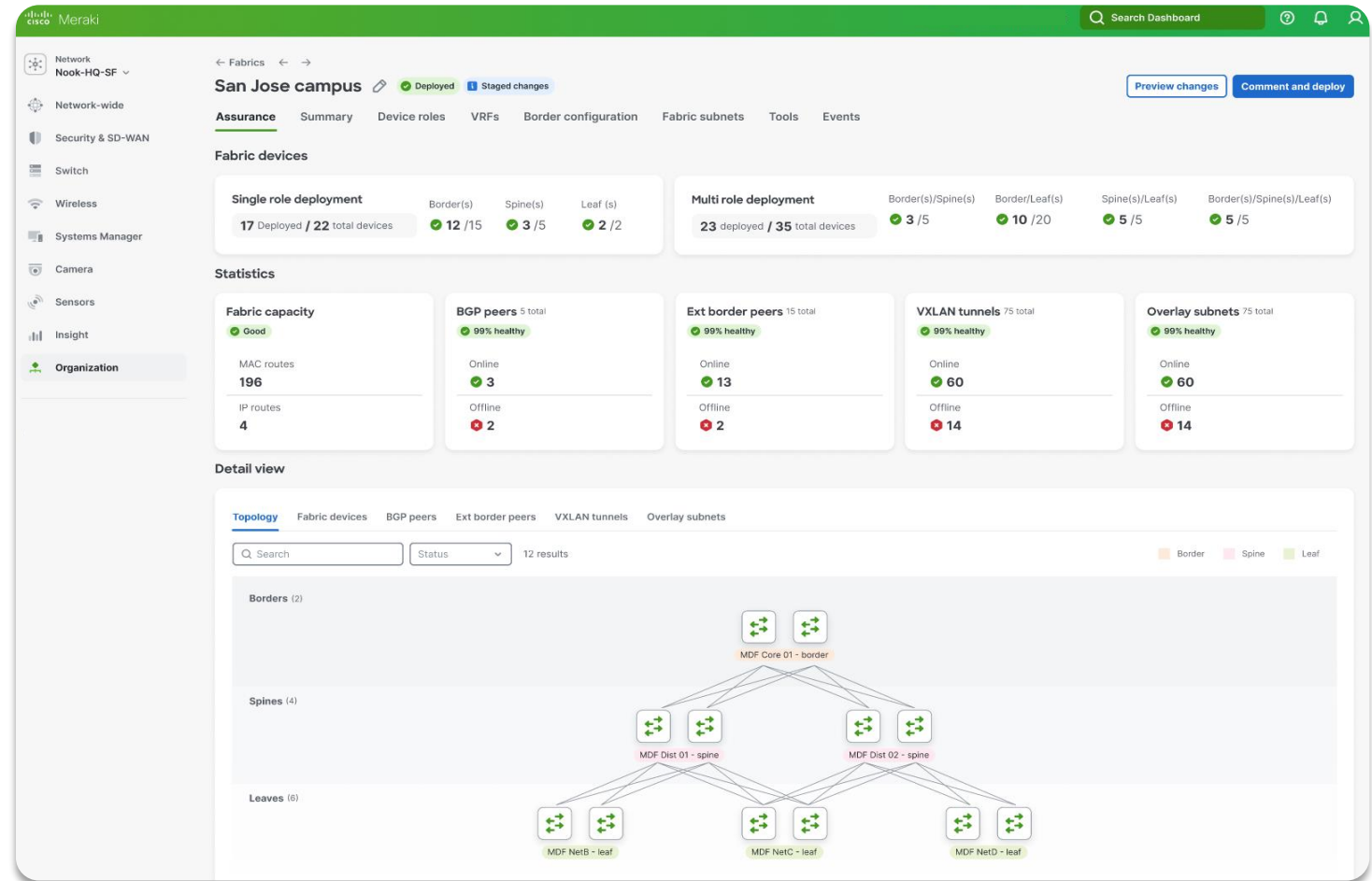
Securing the network with Adaptive Policy



Optimized L2 extension

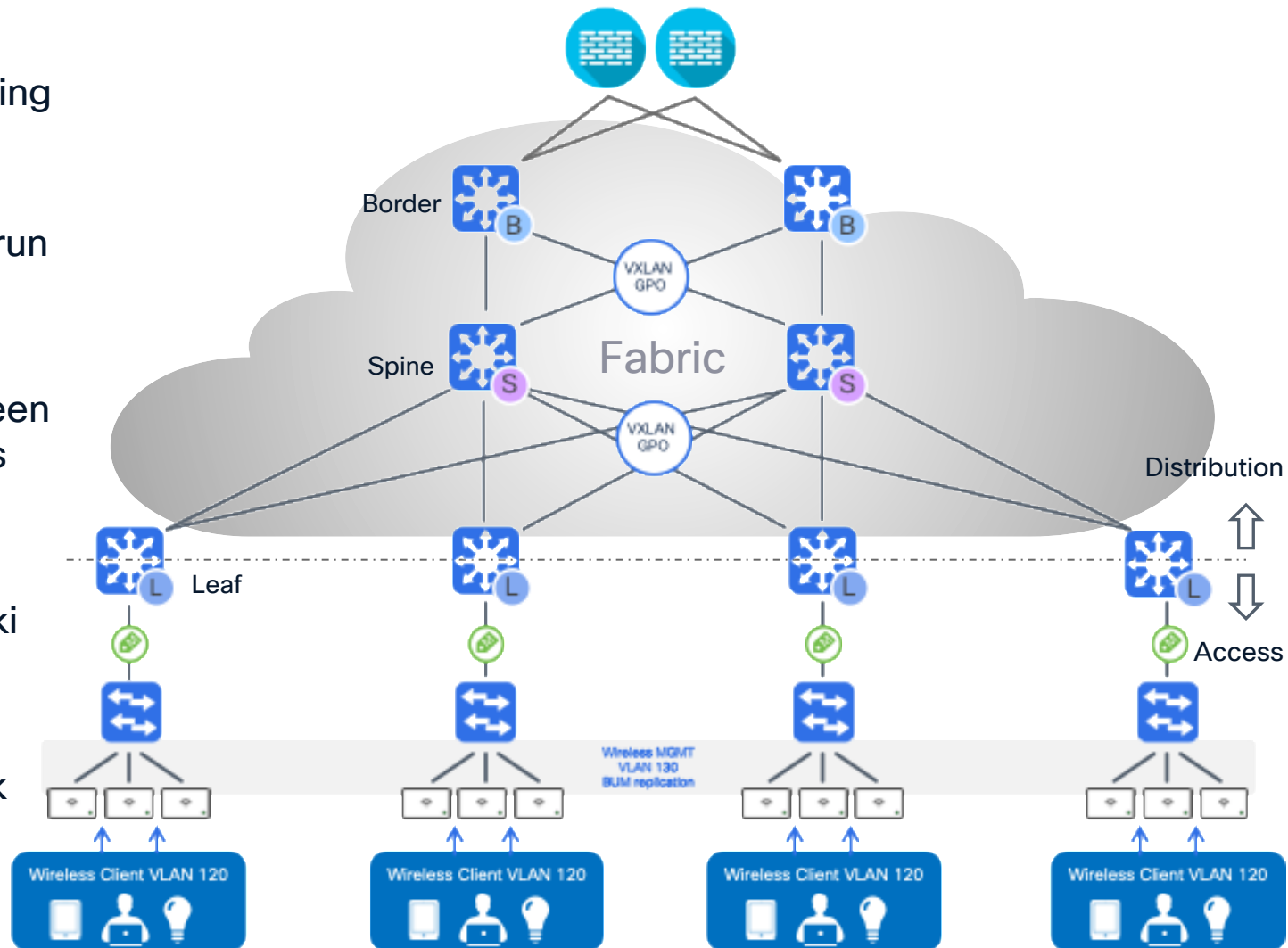


Day-2 Operational Assurance



Cisco Cloud-Managed Fabric – Wireless integration

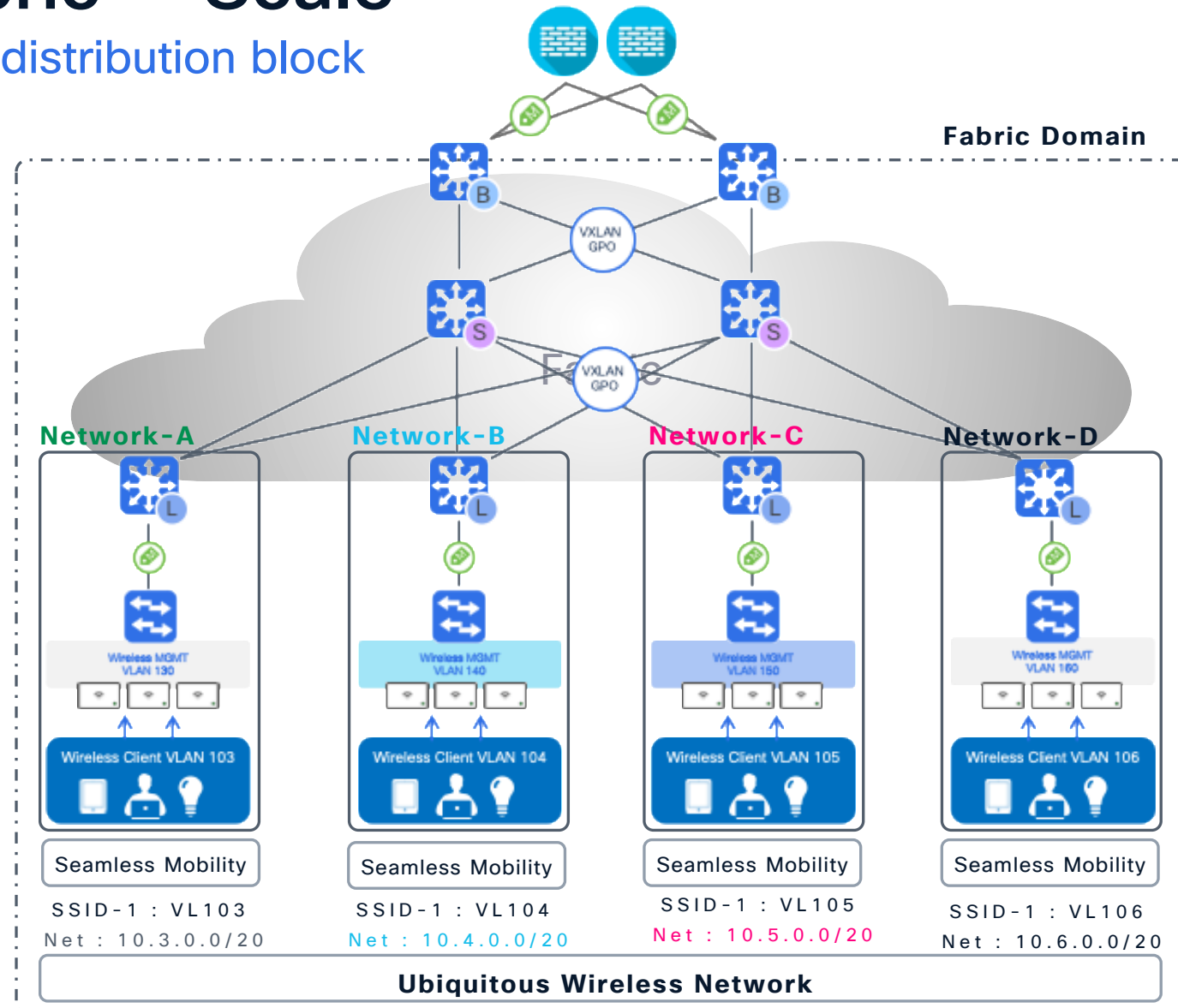
- Fabric starts at the distribution block (E.g., building level). The distribution switch is the leaf of the Fabric.
- Access switches are not part of Fabric and still run spanning tree.
- Client traffic is bridged at the AP
- Support for fast roaming with key sharing between APs and magic ARP for proactive fabric updates
- Adaptive Policy (micro-segmentation) and VRF (macro-segmentation) are supported
- Cloud Fabric is supported across multiple Meraki networks for switching and wireless
- Supported on switches on 17.18.2 and above
- **Wireless scale:** up to 10k IPv4 clients (across 1k APs) per seamless roaming domain



Cisco Cloud-Managed Fabric - Scale

Scenario 1: Seamless roaming domain per distribution block

- **Scenario 1:** seamless roaming domain is limited to the building(s) (or to the distribution block), no outdoor RF coverage and no seamless roaming across building(s).
- No Anycast gateway functionality required in the Fabric
- Client and AP **subnets** are in **Routed mode**
- **Wireless scale: 10k IPv4 clients** (across 1k APs) per building/distribution block.
- **Fabric domain scale:** Four Fabric leaves are supported across the Campus > up to 40k clients (4k APs)
- Routed Overlay Wireless Network: This is the **recommended** design > higher Fabric scale



Scenario 2: Seamless roaming domain across buildings

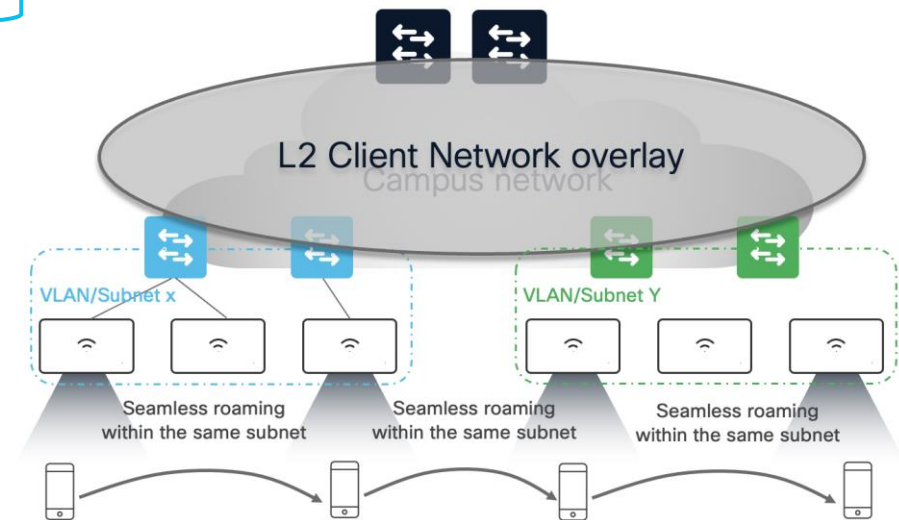
-
- s buildings
- Fabric Domain**
- Network-A**
- Seamless Mobility**
- SSID - 1 : VL120
Net : 10.2.0.0/20
- Ubiquitous Wireless Network**

Wireless Integrated in Fabric - Deployment options

There are multiple ways to integrate wireless in a Fabric network:

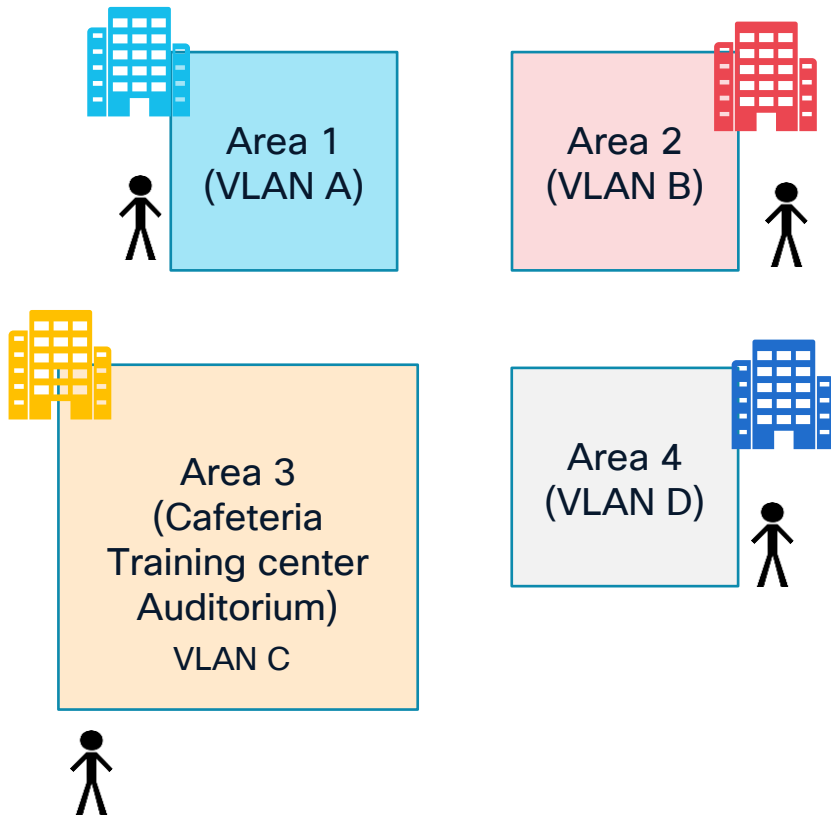
1. **Over the Top (OTT) Distributed** > client traffic is bridged at the AP and enters the Fabric at the access switch. This would be with Flex local switching or Meraki bridge mode
2. **Over the Top (OTT) Centralized** > client traffic is tunneled to a centralized aggregator (WLC/Campus Gateway), Fabric is just a transport. Traffic can enter the Fabric at the aggregator
3. **Integrated (Fabric Enabled)** > Wireless and Fabric Control plane is integrated; client traffic is tunneled from AP to the Fabric edge. This is available only with on-prem SDA (LISP), today.

Wireless network and Fabric network are two ships in the night. In these modes, Fabric reconverge on client roaming needs to be considered.



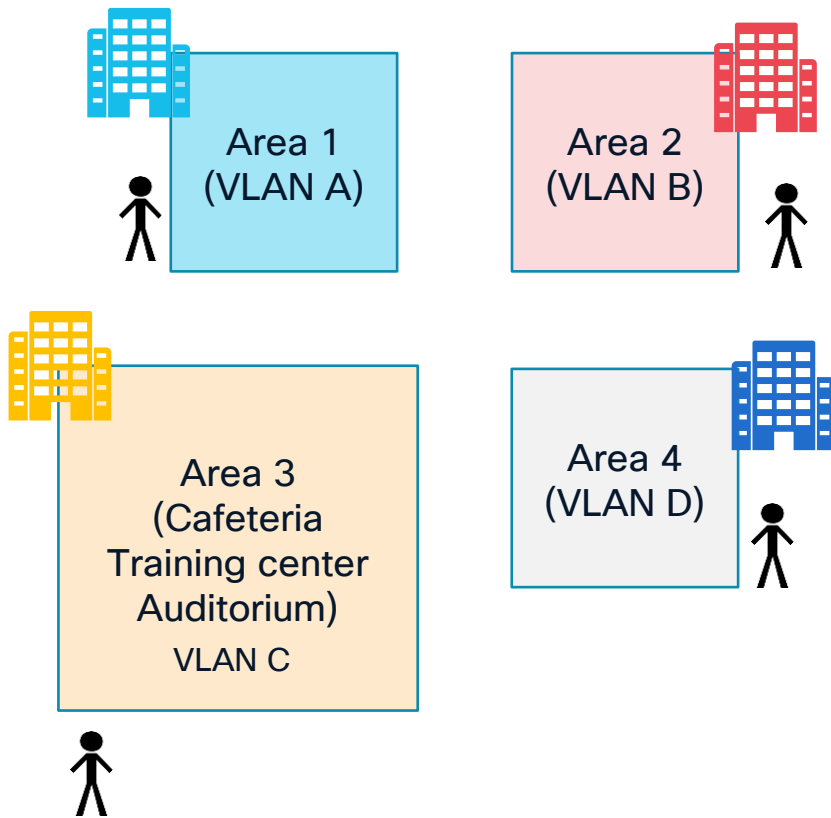
What about other Services?

VLAN/DHCP Design: Distribute or Centralize?



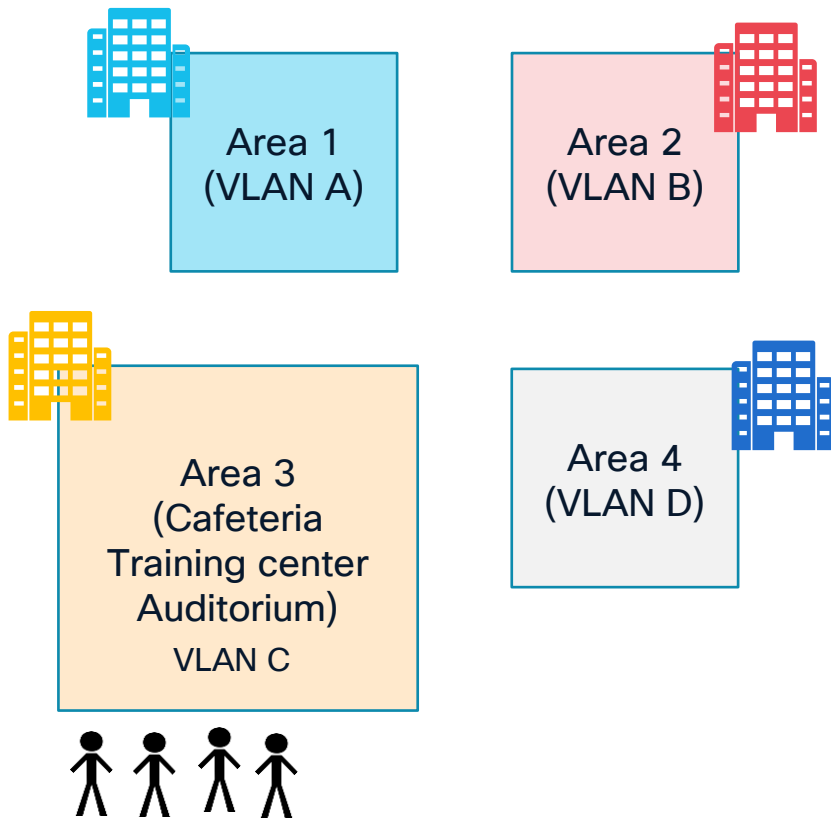
- For **Distributed data plane** deployment consider that any new SSID means new VLAN/Subnet that you need to add on the trunk port to every AP
- Also, additional VLAN/subnet means defining the L3 interface, routing and the related DHCP scope

VLAN/DHCP Design: Distribute or Centralize?



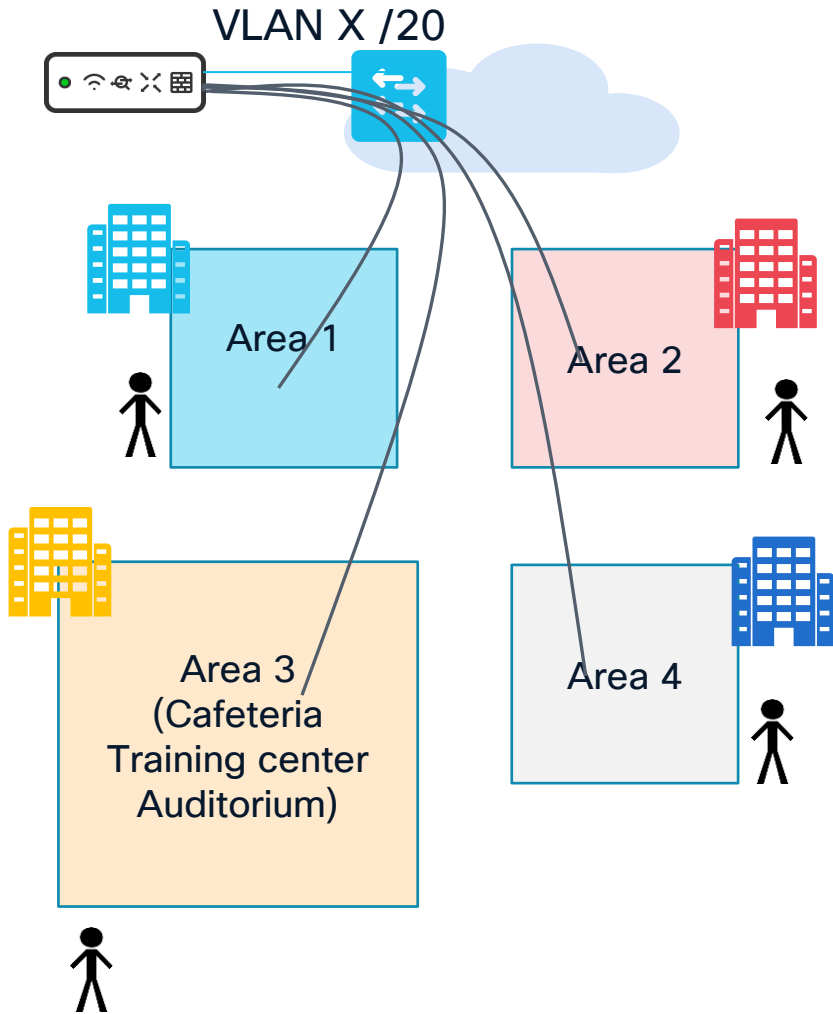
- For **Distributed data plane** deployment consider that any new SSID means new VLAN/Subnet that you need to add on the trunk port to every AP
- Also, additional VLAN/subnet means defining the L3 interface, routing and the related DHCP scope
- Size your **DHCP scope** considering all the possible devices that could join in that area but also roaming devices from other areas. This is important to prevent **DHCP scope starvation**

VLAN/DHCP Design: Distribute or Centralize?



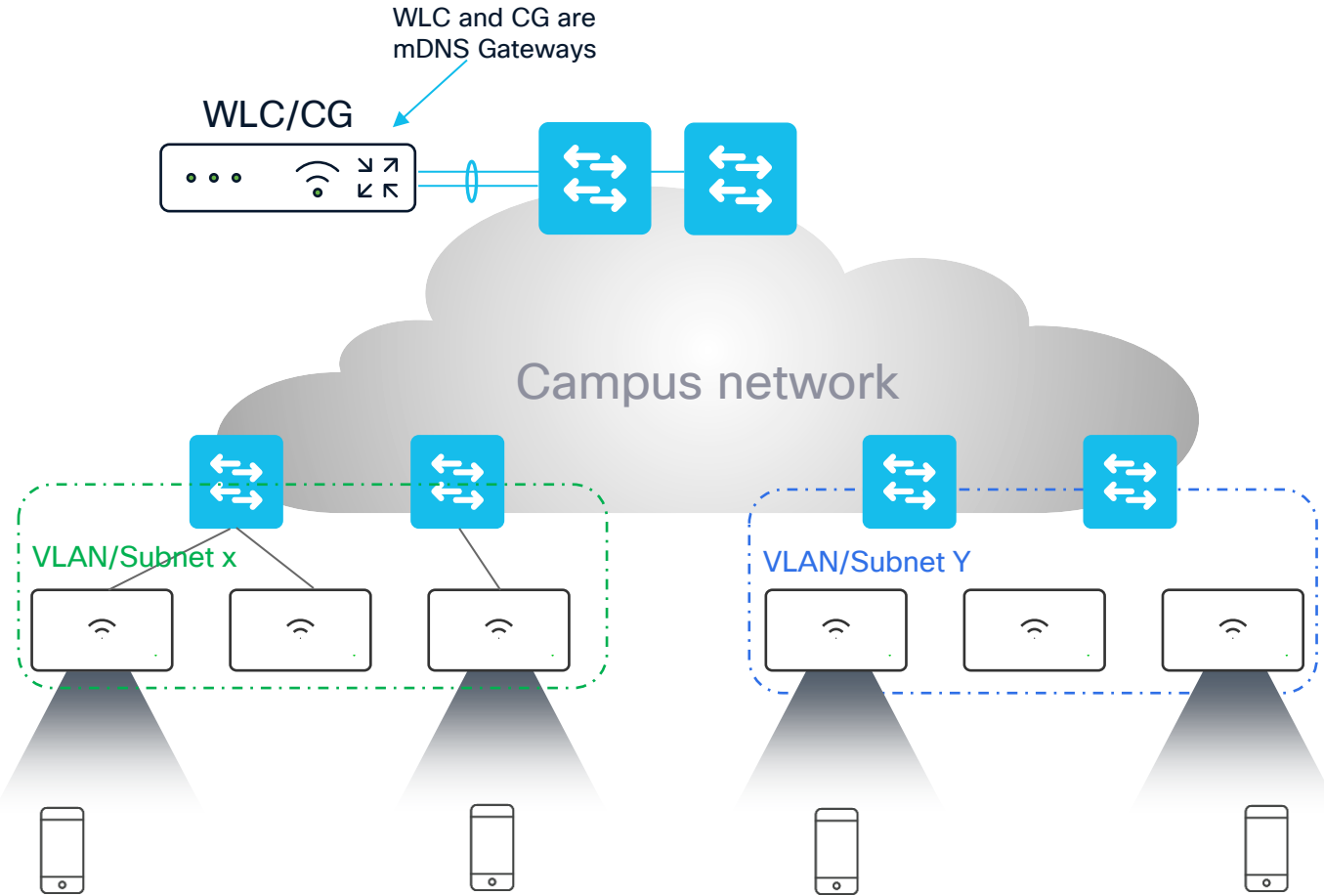
- For **Distributed data plane** deployment consider that any new SSID means new VLAN/Subnet that you need to add on the trunk port to every AP
- Also, additional VLAN/subnet means defining the L3 interface, routing and the related DHCP scope
- Size your **DHCP scope** considering all the possible devices that could join in that area but also roaming devices from other areas. This is important to prevent **DHCP scope starvation**

VLAN/DHCP Design: Distribute or Centralize?



- For **Distributed data plane** deployment consider that any new SSID means new VLAN/Subnet that you need to add on the trunk port to every AP
- Also, additional VLAN/subnet means defining the L3 interface, routing and the related DHCP scope
- Size your **DHCP scope** considering all the possible devices that could join in that area but also roaming devices from other areas. This is important to prevent **DHCP scope starvation**
- For **Centralized data plane** it's easier as as you just need to add the client VLANs in one place, on the link between the edge/aggregator and the distribution switch
- For **DHCP scope**, you can simply have one large subnet (/20, /16 are typical in high scale deployments) and call it the day
- Of course, you need to make sure that the distribution/core switches can scale to the number of MAC and ARP entries for the clients supported by WLC/Campus Gateway.

mDNS: Distribute or Centralize?



mDNS

- If you need mDNS at scale, you need to work across VLANs and you need a **services filtering** function > You need mDNS Gateway
- It's easier to deploy mDNS Gateway as a centralized function (WLC or CG)
- There is mDNS Gateway function in FlexConnect local switching, but not for Meraki Bridged mode

Key Design Factors



Features



Services



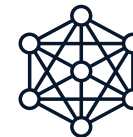
Wired Network Design



Scale



Policy



Use Cases

Wired Network Design

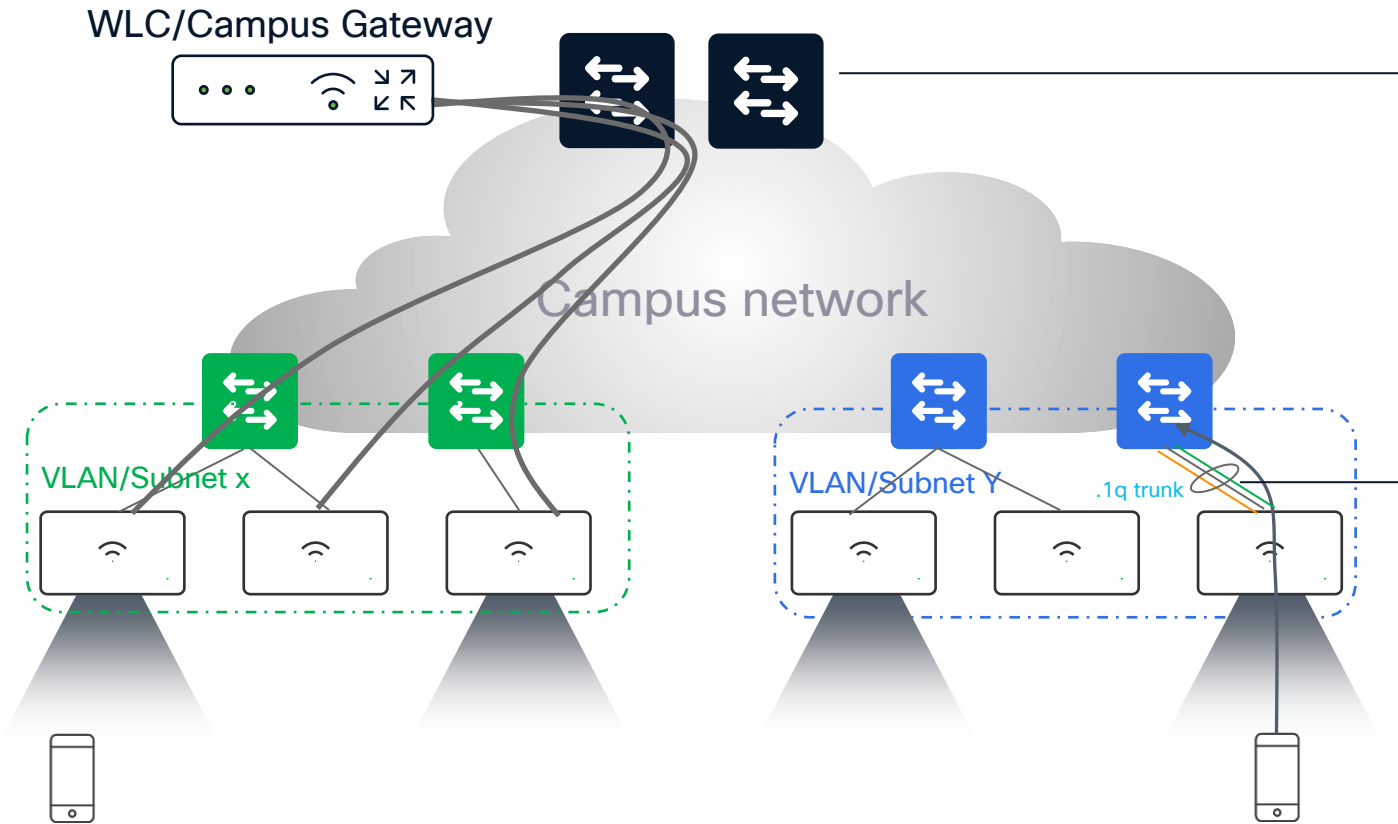
Wireless traffic entry point into the wired network

Distribute: Access Switches configuration overhead?

Centralize: Edge/Aggregator high availability & scale?

Fabric architecture (SDA/BGP EVPN)?

Wired Network Design



Centralized

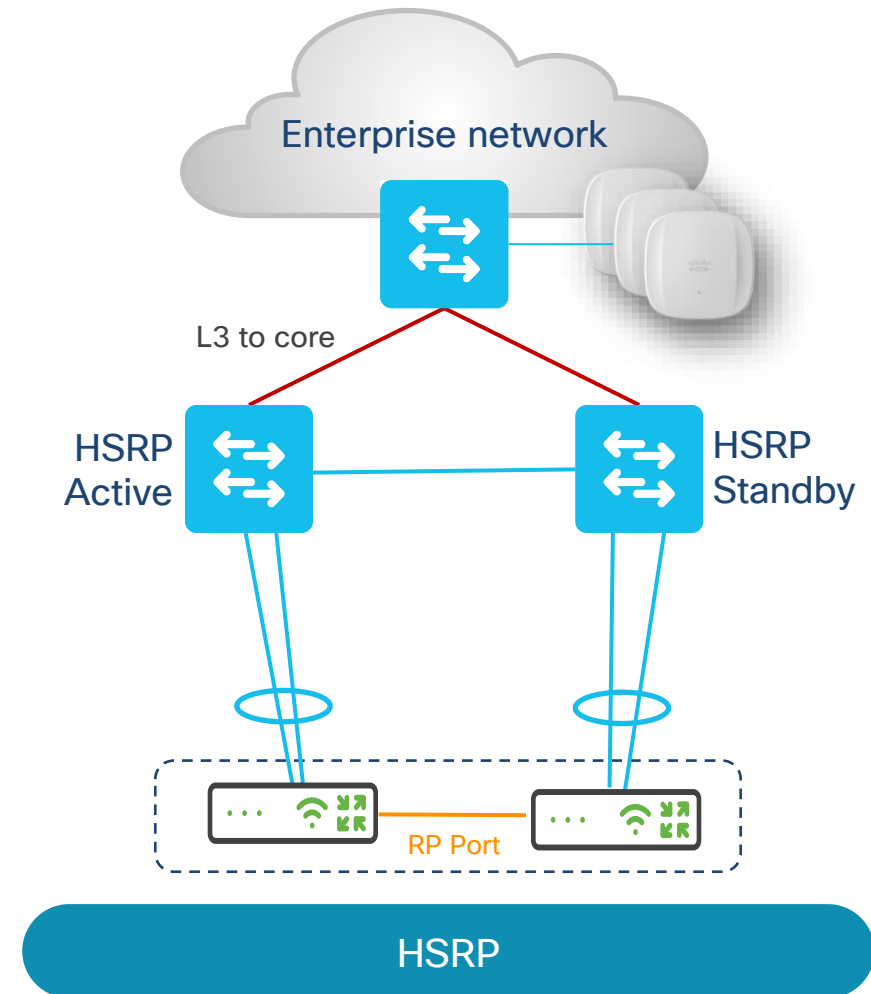
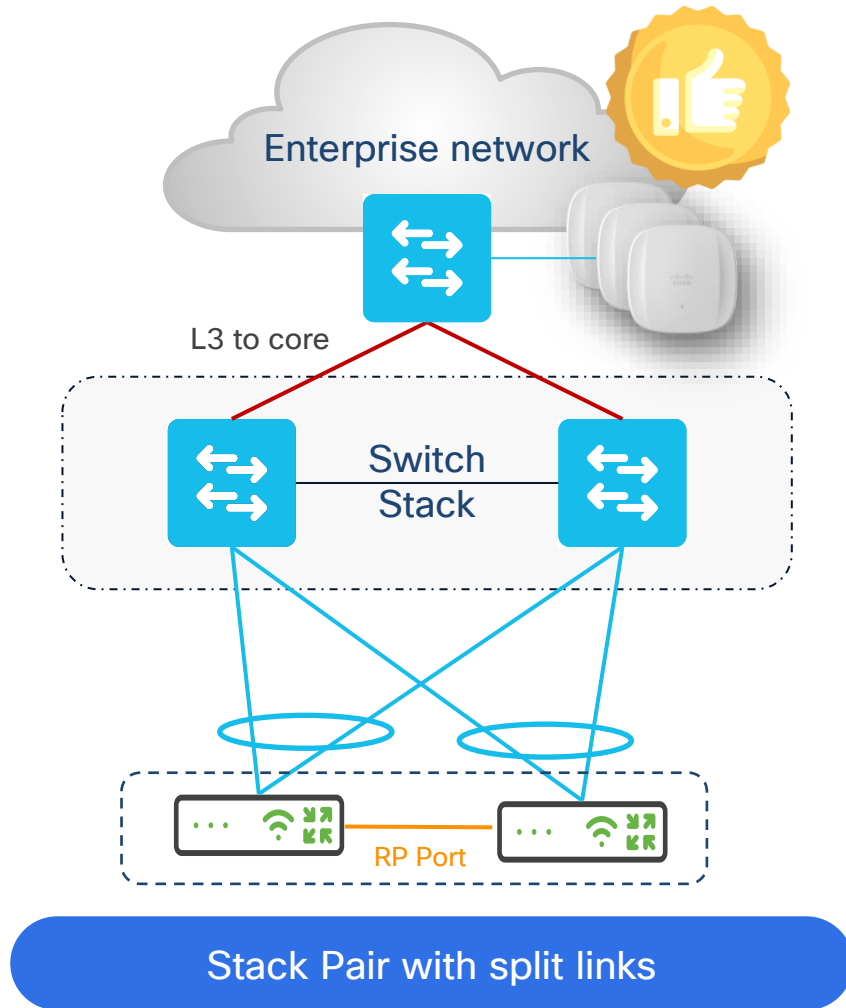
- Client VLANs are only defined centrally; APs are connected to switch ports in access mode
- Single point of Ingress to the wired network, single policy enforcement point for all wireless client traffic
- Single point of failure: how to design for **high availability**?
- Distribution/Core switches need to **scale** for client ARPs and MAC addresses
- Aggregating all clients and traffic in one box: how to optimize **scale**?

Distributed

- Client traffic is bridged at the AP
- SSID/User mapped to a VLAN (at the AP)
- VLAN (and client policy) carried to the switch via uplink 802.1q trunk or VXLAN tunnel
- Multiple switches to configure and maintain
- Multi-point of Ingress to the wired network, policy enforcement applied at every access layer switchport

**Aggregator (= CCG/WLC)
Cluster Redundancy**

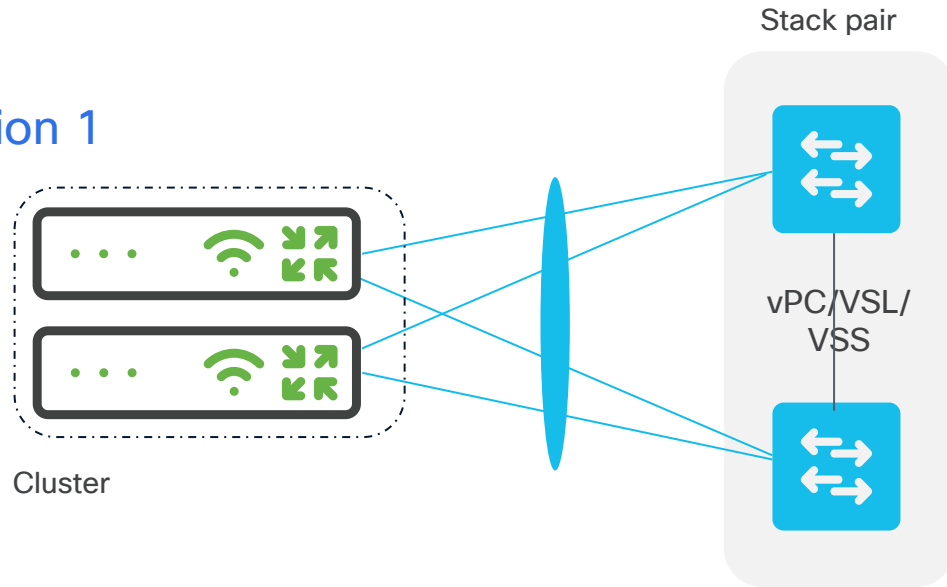
Cluster redundant topologies



Note: RP can be connected back-to-back or via L2 switches

Cluster redundant topologies

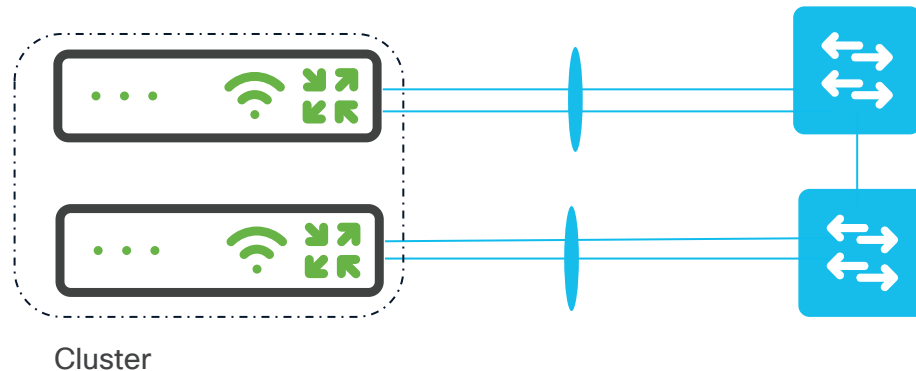
Option 1



Option 1 - Stack Pair with split links:

- Connect each Edge to a single logical switch at the distribution/core using VSS/vPC/VSL stacking.
- Connect links across chassis, one single EtherChannel
- **Maximize redundancy**: if one of the distribution/core switch goes down, the cluster availability is not affected but throughput is lower

Option 2

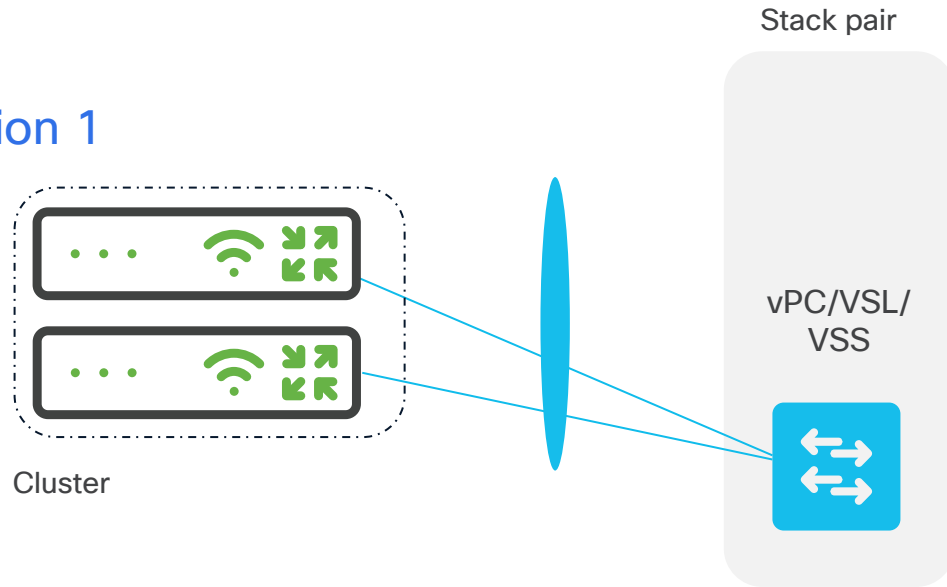


Option 2 - HSRP:

- Connect to two standalone switches, one single EtherChannel to each upstream switch;
- Need a L2 link between switches and run HSRP (or equivalent)
- **Maximize throughput**: if one of the switch goes down, the cluster is affected but throughput remains full

Cluster redundant topologies

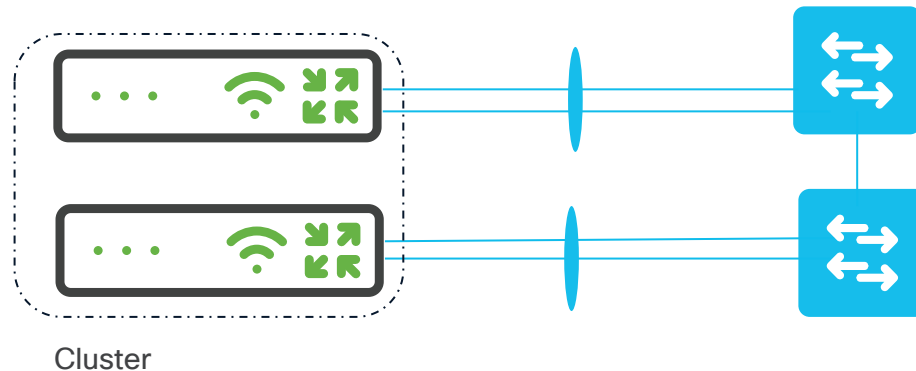
Option 1



Option 1 - Stack Pair with split links:

- Connect each Edge to a single logical switch at the distribution/core using VSS/vPC/VSL stacking.
- Connect links across chassis, one single EtherChannel
- **Maximize redundancy**: if one of the distribution/core switch goes down, the cluster availability is not affected but throughput is lower

Option 2

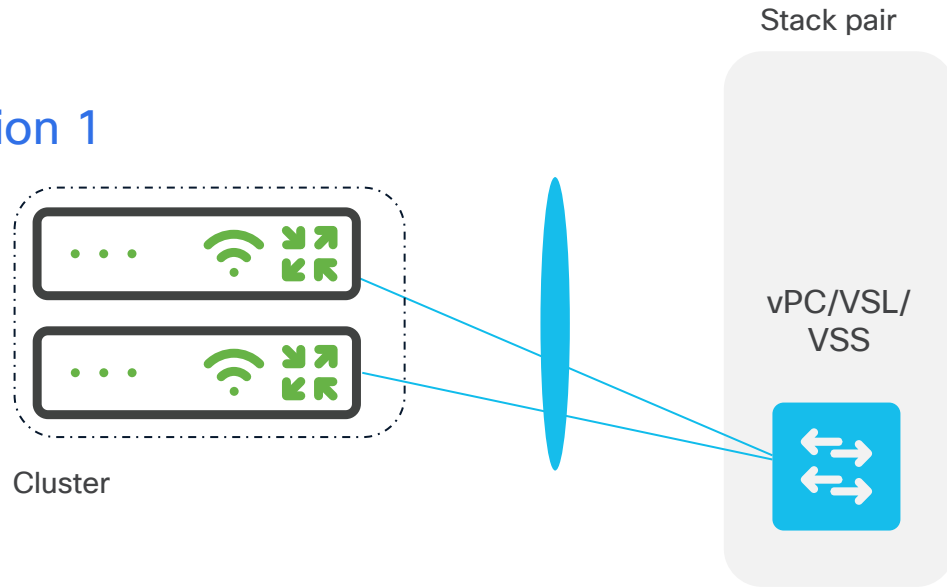


Option 2 - HSRP:

- Connect to two standalone switches, one single EtherChannel to each upstream switch;
- Need a L2 link between switches and run HSRP (or equivalent)
- **Maximize throughput**: if one of the switch goes down, the cluster is affected but throughput remains full

Cluster redundant topologies

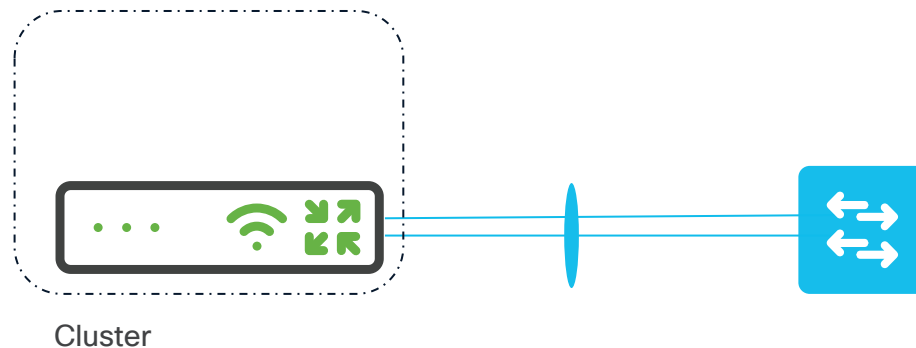
Option 1



Option 1 - Stack Pair with split links:

- Connect each Edge to a single logical switch at the distribution/core using VSS/vPC/VSL stacking.
- Connect links across chassis, one single EtherChannel
- **Maximize redundancy**: if one of the distribution/core switch goes down, the cluster availability is not affected but throughput is lower

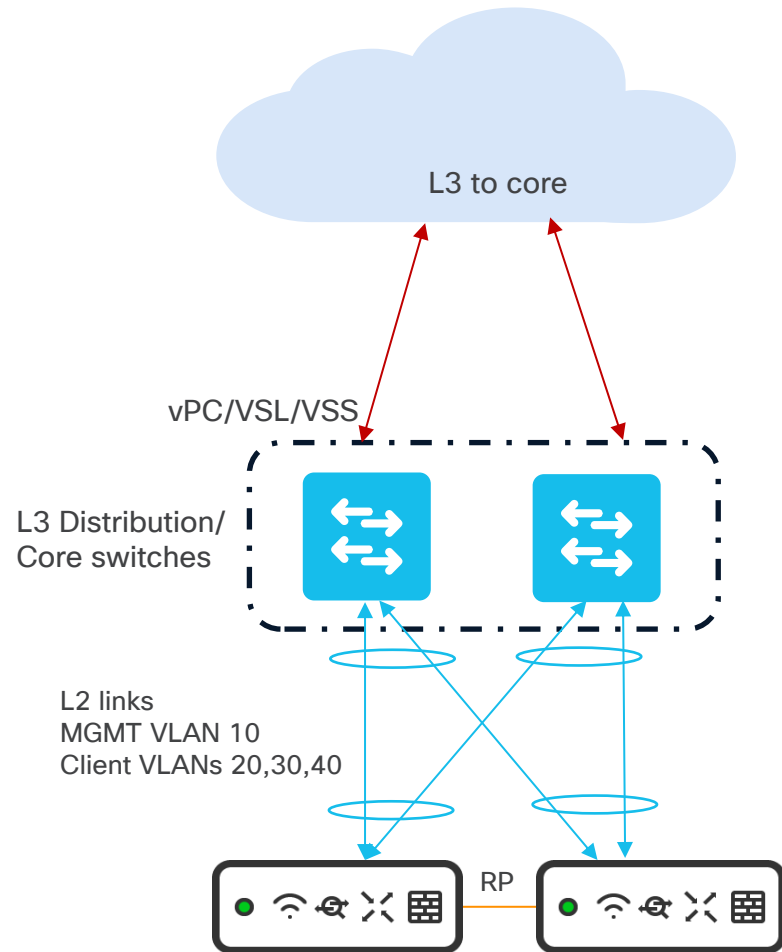
Option 2



Option 2 - HSRP:

- Connect to two standalone switches, one single EtherChannel to each upstream switch;
- Need a L2 link between switches and run HSRP (or equivalent)
- **Maximize throughput**: if one of the switch goes down, the cluster is affected but throughput remains full

Cluster redundant topologies - single Data Center (DC)



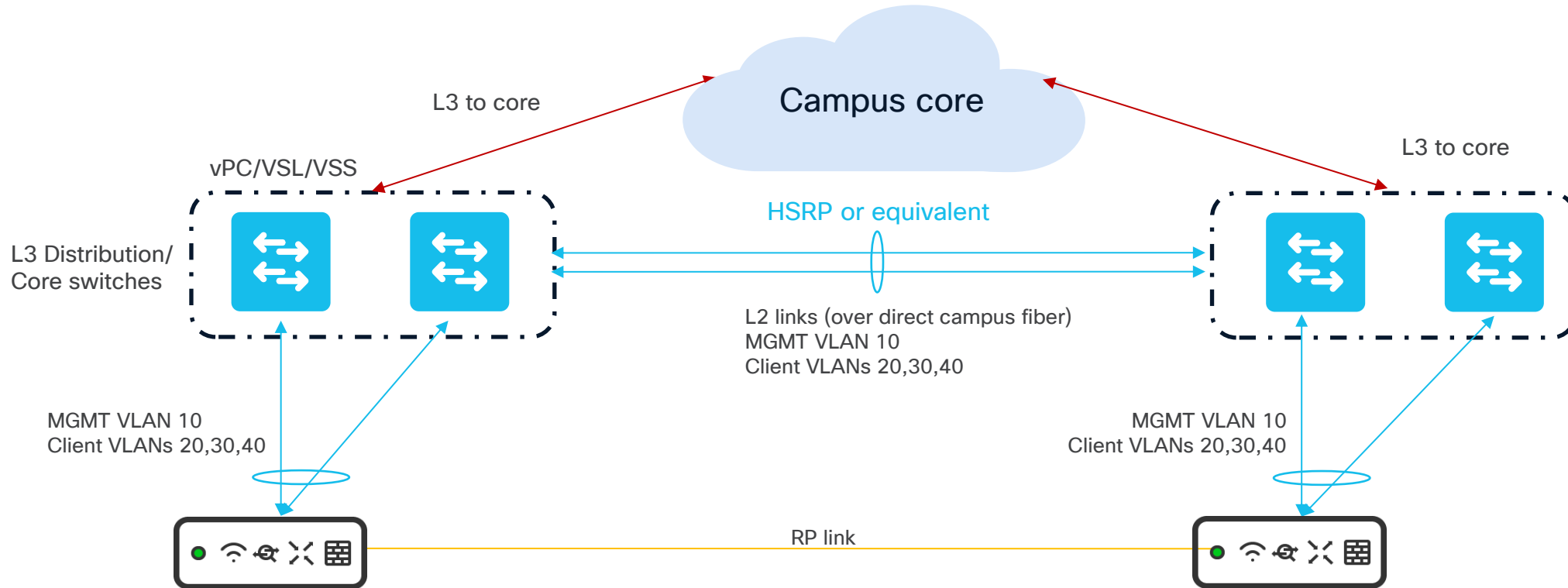
Recommended topology:

- Connect the cluster to a logical switch
- Cross-connect uplinks to a logical distribution switch (vPC, VSL, Switch stack, etc.)
- RP ports directly connected
- Two EtherChannel configured on the switch side: LACP and trunk enabled with the list of allowed VLANs
- Consider the ARP and MAC address table size to select the right switch platform

Distribution/core switch scale

Switch Family		MAC Table Scale	ARP Table Scale
Cisco Catalyst 9600	C9600X-SUP-2 (S1 Q200 based)	256K	128K
	C9600-SUP-1 (UADP 3.0 based)	128K	90K
Cisco Catalyst 9500	UADP 3.0 based	82K	90K
	UADP 2.0 based	64K	80K
Cisco Catalyst 9500X		256K	256K
Cisco Nexus 7000/7700	M3 Line Cards	384K	128K
	F3 Line Cards	64K	128K
Cisco Nexus 9500		90K	90K (60 v4 / 30 v6)
Cisco Nexus 9300	9300 Platforms	90K	65K (40K v4 / 25K v6)
	9300-FX/X/FX2	92K	80K (48K v4 / 32K v6)
Cisco Nexus 9200		96K	64K (32K v4 / 32K v6)

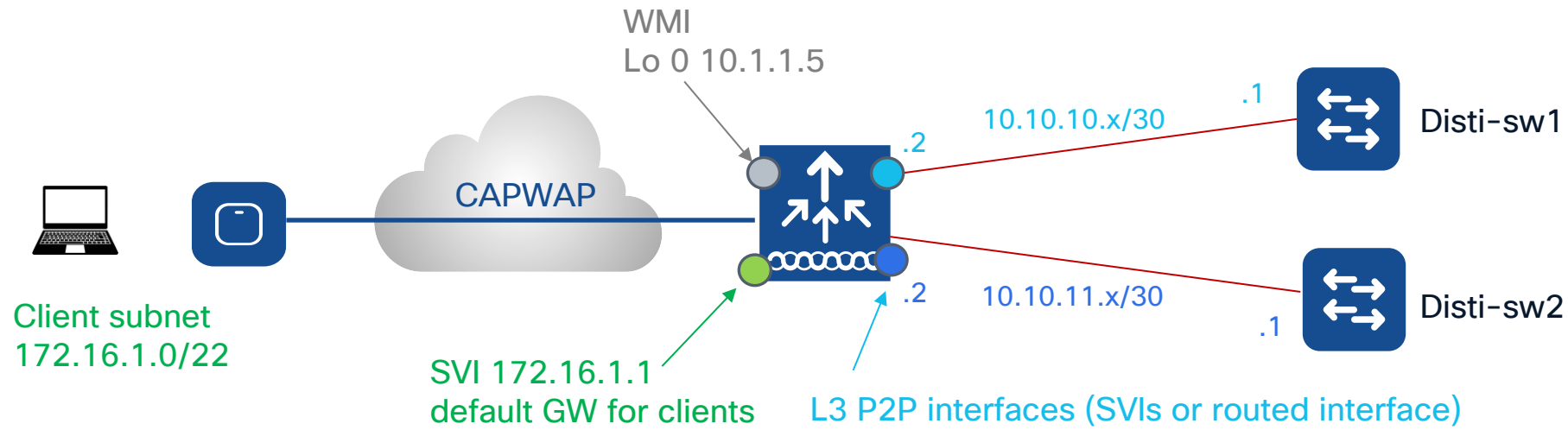
Cluster redundant topologies – Across DCs



- This option is to increase redundancy, splitting the cluster among two different areas of the Data Center with different power distributions or among two different Data Centers
- The requirements: 80 ms latency, 1500B MTU and min 60 Mbps. Inter-DC connection to be L2 (usually fiber), so that Management and client VLANs can be shared
- If distance allow it, it's always recommended to connect the two boxes directly via RP (using fiber port)

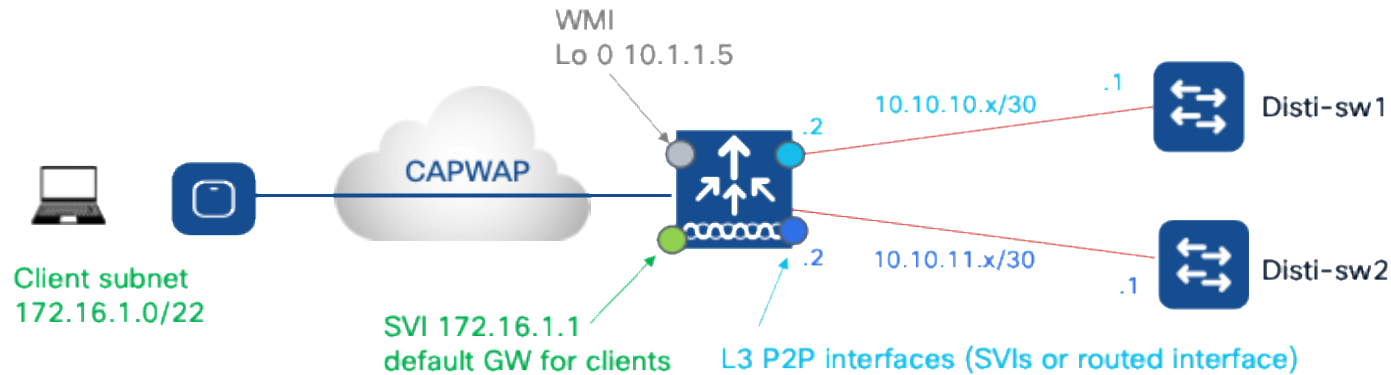
What if your uplink switch does not scale?

L3 forwarding on WLC



- Supported on 17.13 and above
- Configured on an WLAN basis: customer can decide if they want to L2 or L3 forwarding
- If L3 forwarding/routing, a L3 client SVI is configured as the default gateway for clients
- Client traffic is terminated on C9800 and routed out of the uplinks as per routing table
- WMI still terminates CAPWAP and can be defined as a loopback to be reached across multiple uplinks
- Each uplink connection to the switch would be an Ether-channel and use ECMP

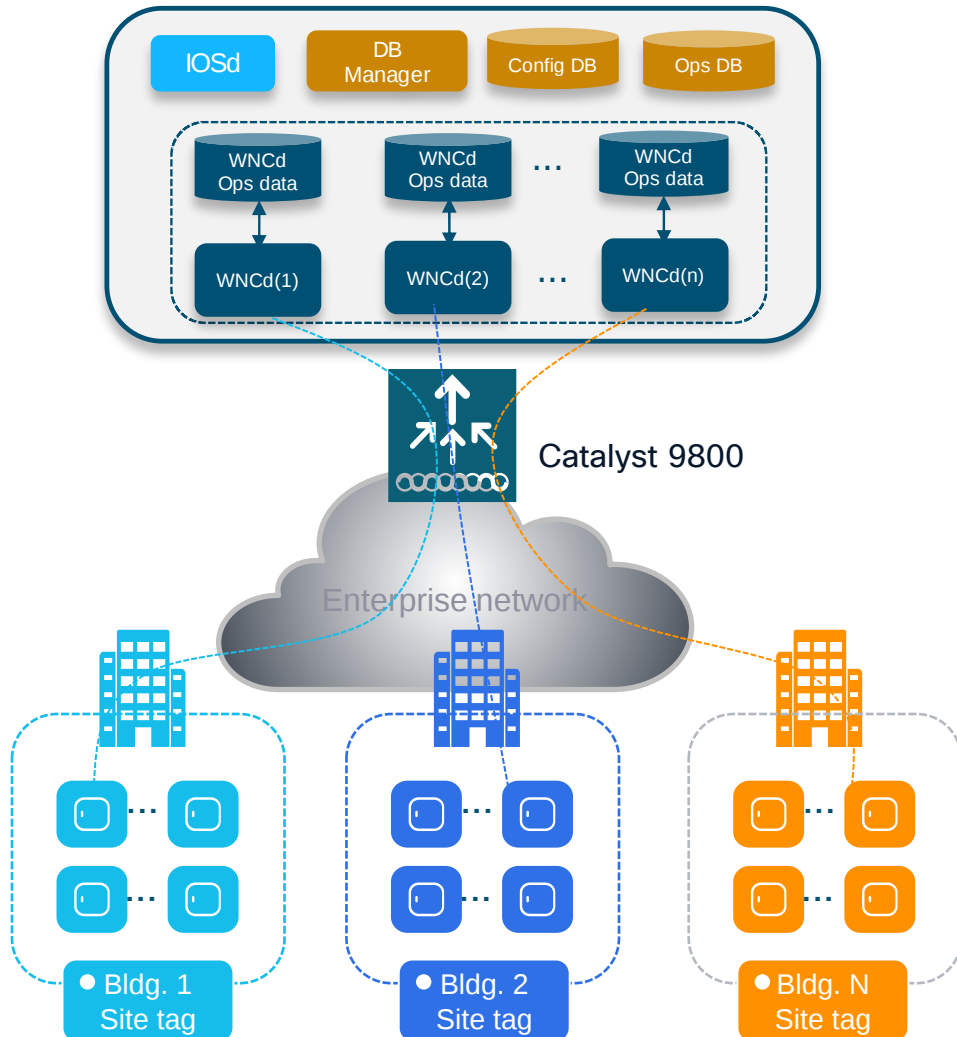
Why L3 forwarding?



- No scale constraints on the upstream switches (ARP, MAC tables). No client MAC/IP is visible to the upstream switch
- Better uplink network design:
 - When deploying with standalone switches (no stacking technology), provides a more efficient uplink load balancing with Equal Cost Multi-path (ECMP) vs. spanning tree (not supported by WLC/CG)
 - Better network convergence using L3 protocols like OSPF
- Can be combined with VRF support and Client Overlapping IPs
- **Keep in mind:**
 - Only supported on C9800 WLC, not supported on Campus Gateway (yet)
 - High Availability: N+1 officially supported, SSO not officially supported

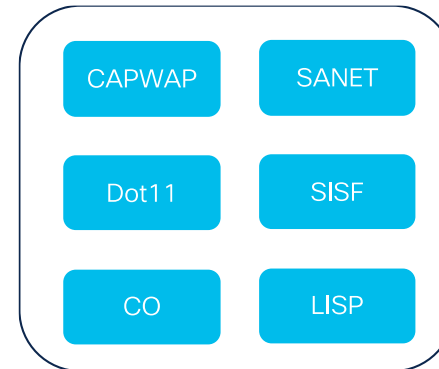
Centralized: How to optimize it for scale?

Wireless Network Control demon - WNCd



Important facts:

- C9800 has a multi-process software architecture
- AP and client sessions are handled by the Wireless Network Controller processes (WNCd) process within a C9800

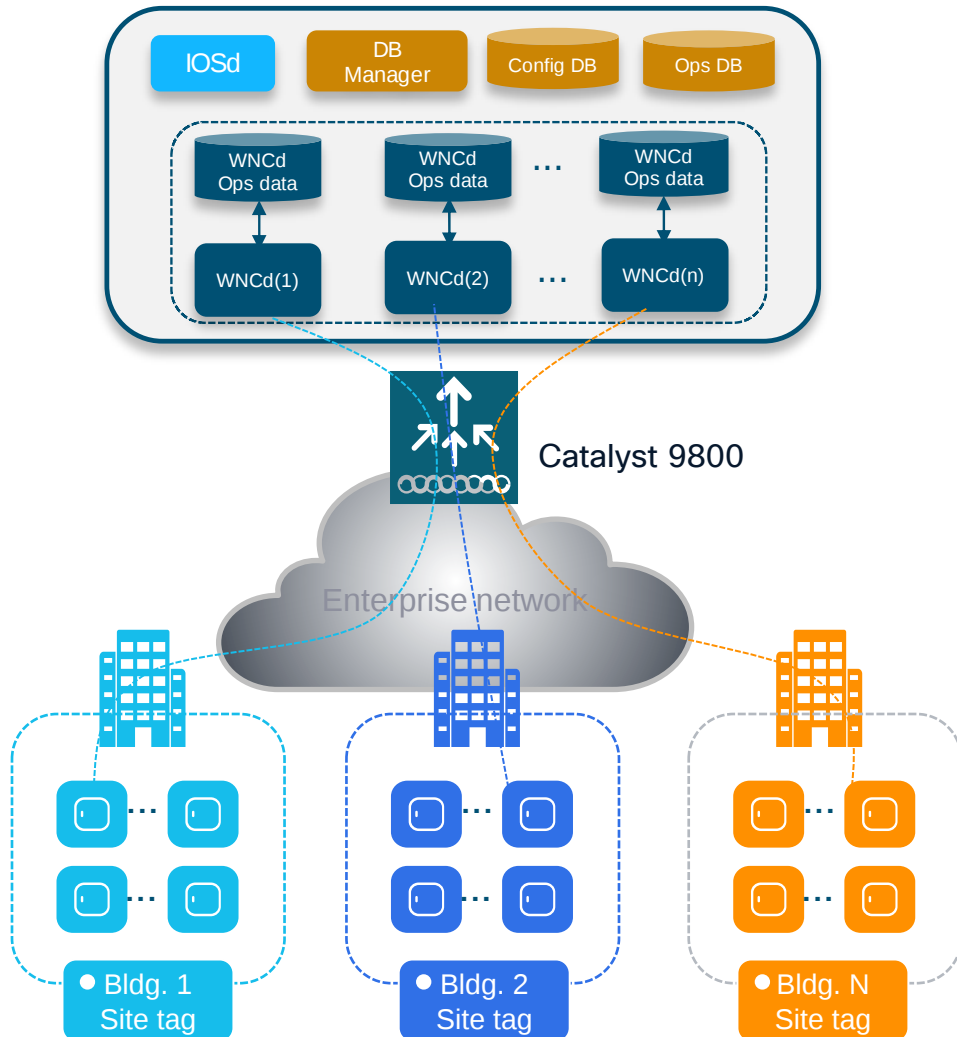


WNCd

Wireless Network Controller demon

- CAPWAP : AP discovery
- Dot11 : Client dot11
- Session Aware Networking (SANet)/AAA: Client authentication
- Switch Integrated Security Features (SISF)
- Client Orchestrator : Client state transitions
- LISP-agent for Fabric deployment

AP to WNCd distribution – Site tags based



How AP distribution across WNCds works:

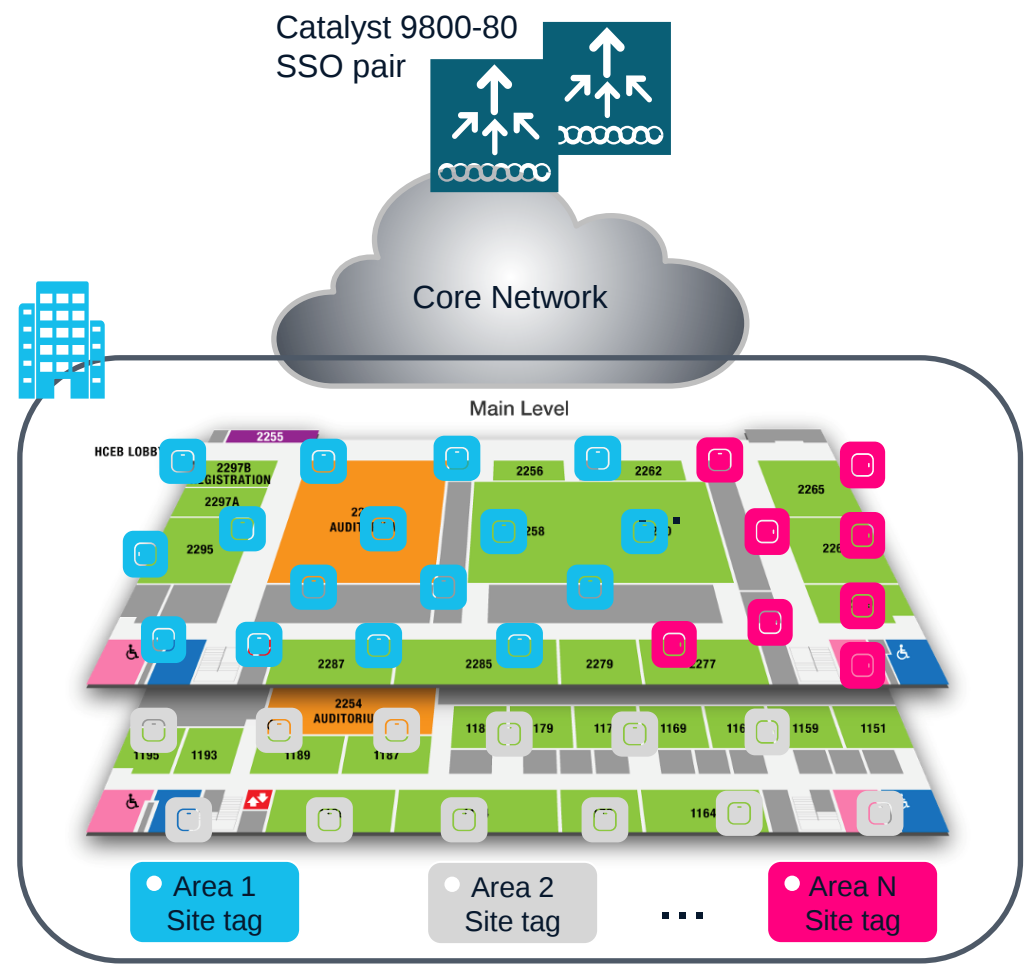
- **APs to WNCd** mapping happens **at AP joining time**. Mapping is considered only for the first AP joining with the new site tag
- **AP distribution to WNCd** processes is based on **Site Tag**: APs with the same site-tag are managed by the same WNCd
- Except when you using **default-site-tag**: in this case APs are distributed among WNCds in round robin fashion
- Site tags are distributed among WNCds using the **least loaded** criteria based **on the number of site tags** (not the # of APs)
- **IMPORTANT**: the site tag doesn't have to coincide with a geographical physical site but should be limited by a geo site. The **site tag is a logical group of access points**

Need to worry about WNCds?

- If you don't see problems with WNCd CPU
- Problem: CPU is $> 70\%$ for at least 5 mins
- Otherwise, relax....
- Always good to follow the recommendations



AP to WNCd distribution – Site tags based



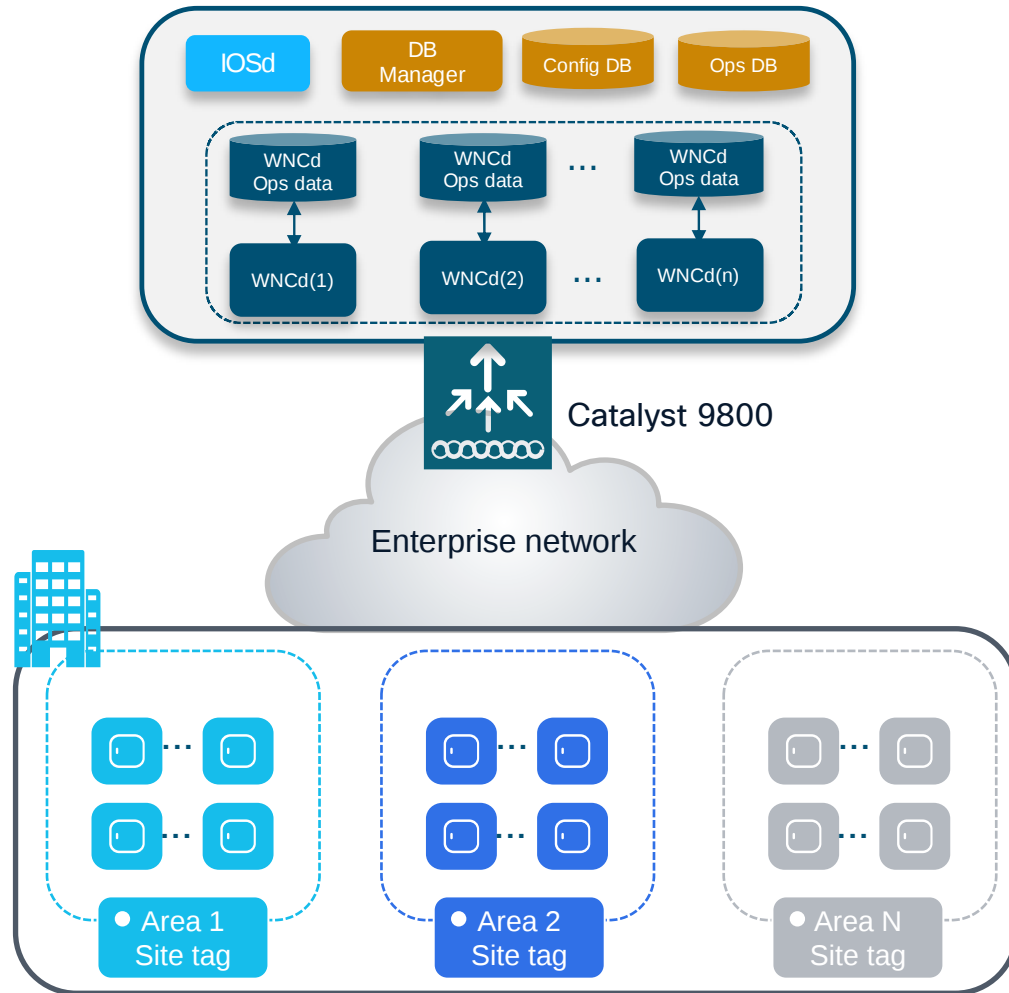
Recommendations

- Use custom site tags
- If you can design site tags and AP assignment: **use the number of site tags equal to the number of WNCd processes per WLC platform, and evenly assign APs to the site tags**

Platform	Recommended # of site tags
C9800-80 / CW9800H1/H2	8
C9800-CL (large)	7
C9800-40 / CW9800M	5
C9800-CL (Medium)	3

- Group APs at a roaming domain level > **Site Tag = Roaming Domain**

AP to WNCd distribution – Site tags based

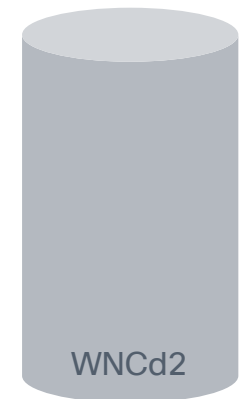
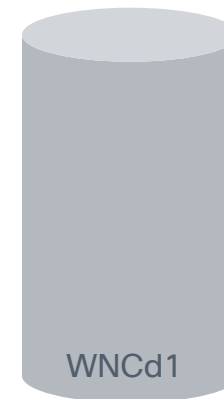
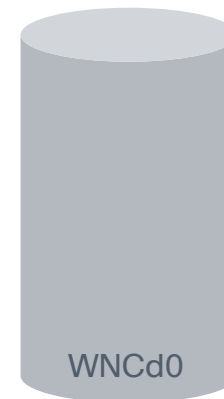
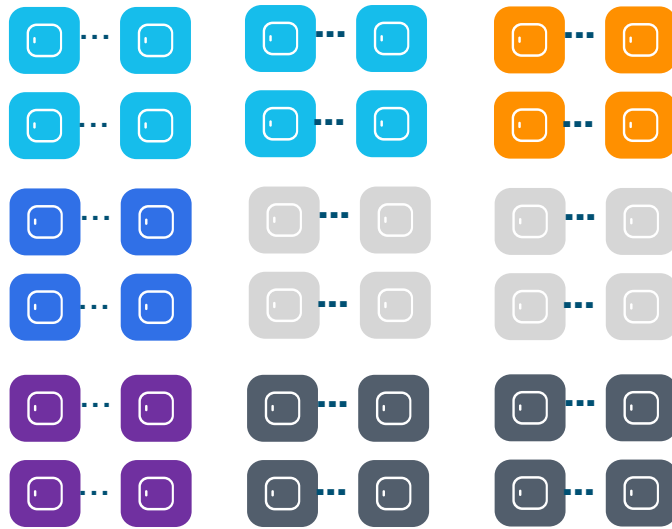


Recommendations

- What if site tags have already been created and customer does not want to change the design?
- What if the number of site tags is different than recommended and/or distribution of APs is uneven?
- Use the **Load** command:
 - Load is an estimate of the WNCd capacity reserved for that site tag and it's used to assign site tags (APs) to the WNCd
 - What contributes to the load of the WNCd: all control plane traffic > client joining, authentication, roaming, client probes, but also features like mDNS that require CPU time
 - To start, set the “load” equal the # of APs on each site tag and reboot the WLC. Then monitor and tweak
- What if you cannot design with site tags? You don't have AP names, you don't have APs on maps, or simply you don't want to. Load command will not help...
- Use RF based Load Balancing. Need 17.12 and above

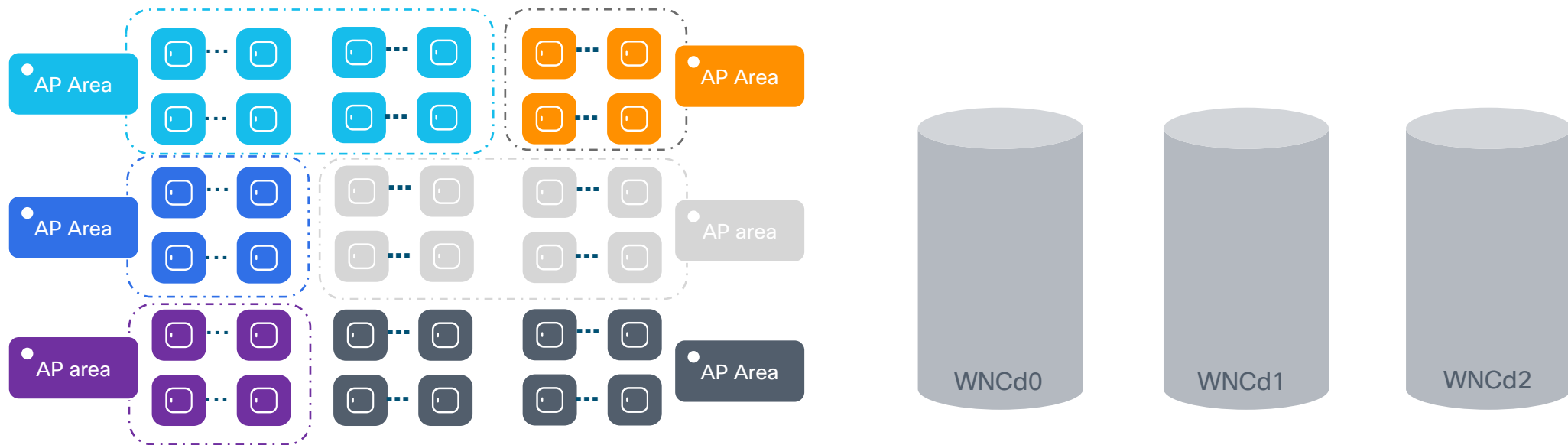
AP to WNCd distribution – RF based

- RRM-based, automatic way of clustering APs and evenly distribute them across WNCds.
- RF based clusters (**AP Areas**) are formed using RSSI info received from RRM AP neighbour reports
- The algorithm can be run on demand or scheduled. The resulting AP load balancing is applied upon WLC reboot or admin trigger which causes AP CAPWAP restart
- When RF based is applied, **load balancing based on site tag is not considered**. So yes, you can have all the APs in default site tag
- It's off by default and **it requires the APs to be deployed and a stable RF** (APs have their neighbours discovered).



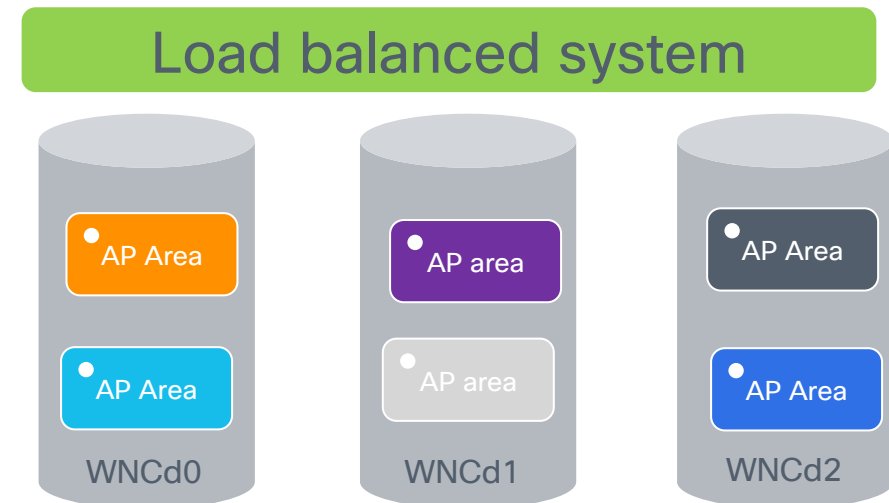
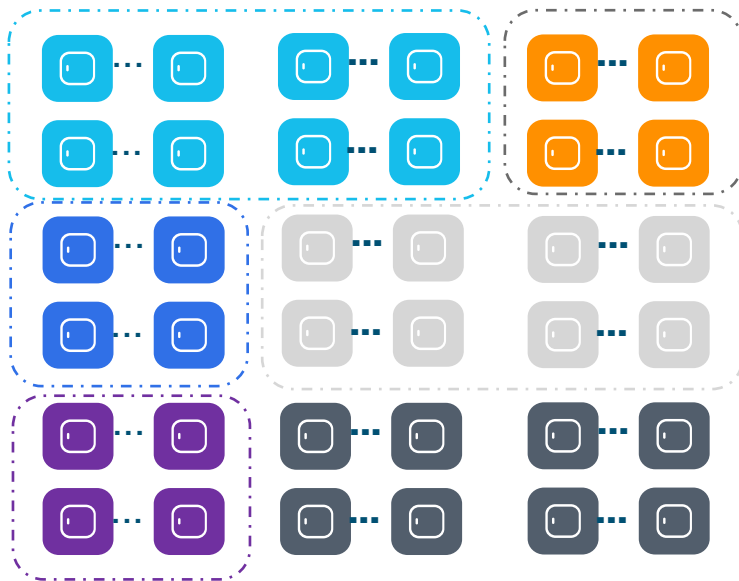
AP to WNCd distribution – RF based

- RRM-based, automatic way of clustering APs and evenly distribute them across WNCds.
- RF based clusters (**AP Areas**) are formed using RSSI info received from RRM AP neighbour reports
- The algorithm can be run on demand or scheduled. The resulting AP load balancing is applied upon WLC reboot or admin trigger which causes AP CAPWAP restart
- When RF based is applied, **load balancing based on site tag is not considered**. So yes, you can have all the APs in default site tag
- It's off by default and **it requires the APs to be deployed and a stable RF** (APs have their neighbours discovered).

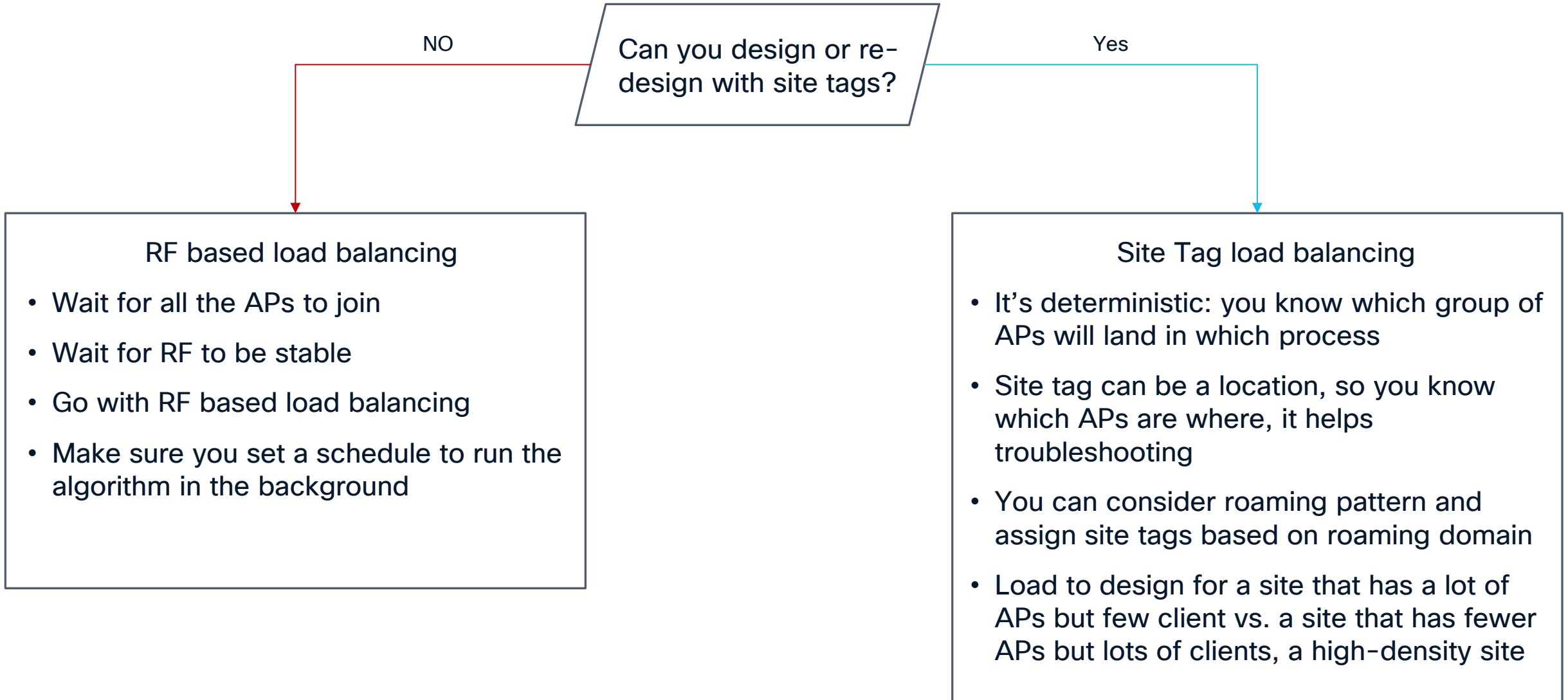


AP to WNCd distribution – RF based

- RRM-based, automatic way of clustering APs and evenly distribute them across WNCds.
- RF based clusters (**AP Areas**) are formed using RSSI info received from RRM AP neighbour reports
- The algorithm can be run on demand or scheduled. The resulting AP load balancing is applied upon WLC reboot or admin trigger which causes AP CAPWAP restart
- When RF based is applied, **load balancing based on site tag is not considered**. So yes, you can have all the APs in default site tag
- It's off by default and **it requires the APs to be deployed and a stable RF** (APs have their neighbours discovered).



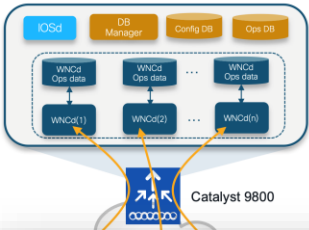
Site tag vs. RF based load balancing



Site Tag Design: What about Campus Gateway?

C9800

Site Tags – AP to WNCd distribution



Recommendations:

- Use custom site tags
- Whenever possible, have less than 500 APs per site tag
- Do not overwhelm a site-tag and WNCd. Do not exceed the following max number of APs per site tag:

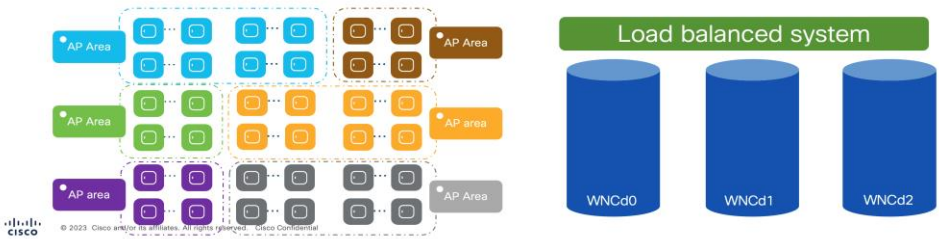
Platform	Max APs per site tag
9800-80, 9800-CL (Medium and Large)	1600
9800-40	800

Catalyst 9800 IOS-XE 17.12

RRM based Auto WNCd load balancing

What is it?

- RRM-based, automatic way of clustering APs and evenly distribute them across WNCds.
- RF based clusters (**AP Areas**) are formed using RSSI info received from RRM AP neighbour reports
- The algorithm can be run on demand or scheduled. It's off by default and it requires the APs deployed and a stable RF (APs have their neighbours discovered). Works with any site tag configuration.
- The resulting AP load balancing is applied upon WLC reboot or admin trigger which causes AP CAPWAP restart
- When applied, it overwrites any other load balancing based on site tag and load



Site Tag Design: What about Campus Gateway?

C9800



**No need to worry anymore about
AP load balancing to internal
processes!**

AP to WNCd distribution

Recommendations:

- Use custom site tags
- Whenever possible, have less than 500 APs per site tag

Do not exceed the maximum number of APs per site tag:

Platform	Max APs per site tag
9800-80, 9800-80-40 (Medium and Large)	1600
9800-40	800

Catalyst 9800 IOS-XE 17.12

RRM based Auto WNCd load balancing

What is it?

RRM-based, automatic way of clustering APs and evenly distribute them across WNCds.

- The algorithm can be run on demand or scheduled. It is off by default and it requires the APs deployed and a stable RF (APs have their neighbours discovered). Works with any site tag configuration.
- The algorithm is applied to all APs in the network.
- When applied, it overwrites any other load balancing based on site tag and load.

balanced system

WNCd0 WNCd1 WNCd2

Site Tag Design: What about Campus Gateway?

C9800



AP to WNCd distribution

Recommendations:

- Use custom site tags
- Whenever possible, have less than 500 APs per site tag

Do not exceed the maximum number of APs per site tag:

Platform	Max APs per site tag
Catalyst 9800 (9800-80, 9800-100, 9800-110, 9800-120, 9800-130, 9800-140, 9800-150, 9800-160, 9800-170, 9800-180, 9800-190, 9800-200, 9800-210, 9800-220, 9800-230, 9800-240, 9800-250, 9800-260, 9800-270, 9800-280, 9800-290, 9800-300, 9800-310, 9800-320, 9800-330, 9800-340, 9800-350, 9800-360, 9800-370, 9800-380, 9800-390, 9800-400)	1600
Catalyst 9800 (9800-40)	800

Catalyst 9800 IOS-XE 17.12

RRM based Auto WNCd load balancing

What is it?

RRM-based, automatic way of clustering APs and distribute them across WNCds.

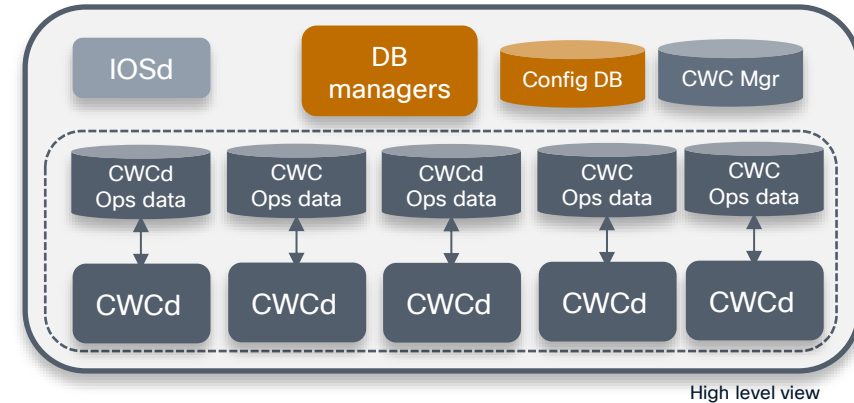
- The algorithm can be run on demand or scheduled. It is off by default and it requires the APs deployed and a stable RF (APs have their neighbours discovered). Works with any site tag configuration.
- The algorithm is applied to all APs in the system.
- When applied, it overwrites any other load balancing based on site tag and load.

balanced system

WNCd0 WNCd1 WNCd2

No need to worry anymore about AP load balancing to internal processes!

Campus Gateway

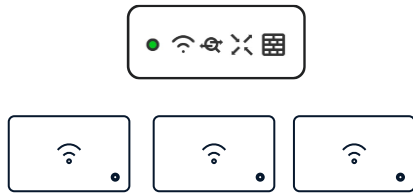


- Campus Gateway's software architecture is different. Main wireless process is [Cisco Wireless Concentrator Demon](#) (CWCd)
- APs and Clients are load balanced across CWCd based on least loaded process, independent from each other
- In C9800 there is the need to have AP and clients in same WNCd because of BSSID and radio info
- In Campus Gateway, the client state machine is handled at the AP, so we don't need APs and clients in same CWCd process > load balance of APs and clients is independent.

Let's revisit the Scale factor...

Large, Medium Campus (no Fabric)

> Centralized architecture



km^2/mi^2

E.g., University Campus



Distribute Enterprise: Branch, Small Campus

Large, Medium Campus with Fabric

> Distributed architecture



m^2/ft^2

E.g., Campus, Retail



Key Design Factors



Features



Services



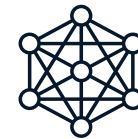
Wired Network Design



Scale

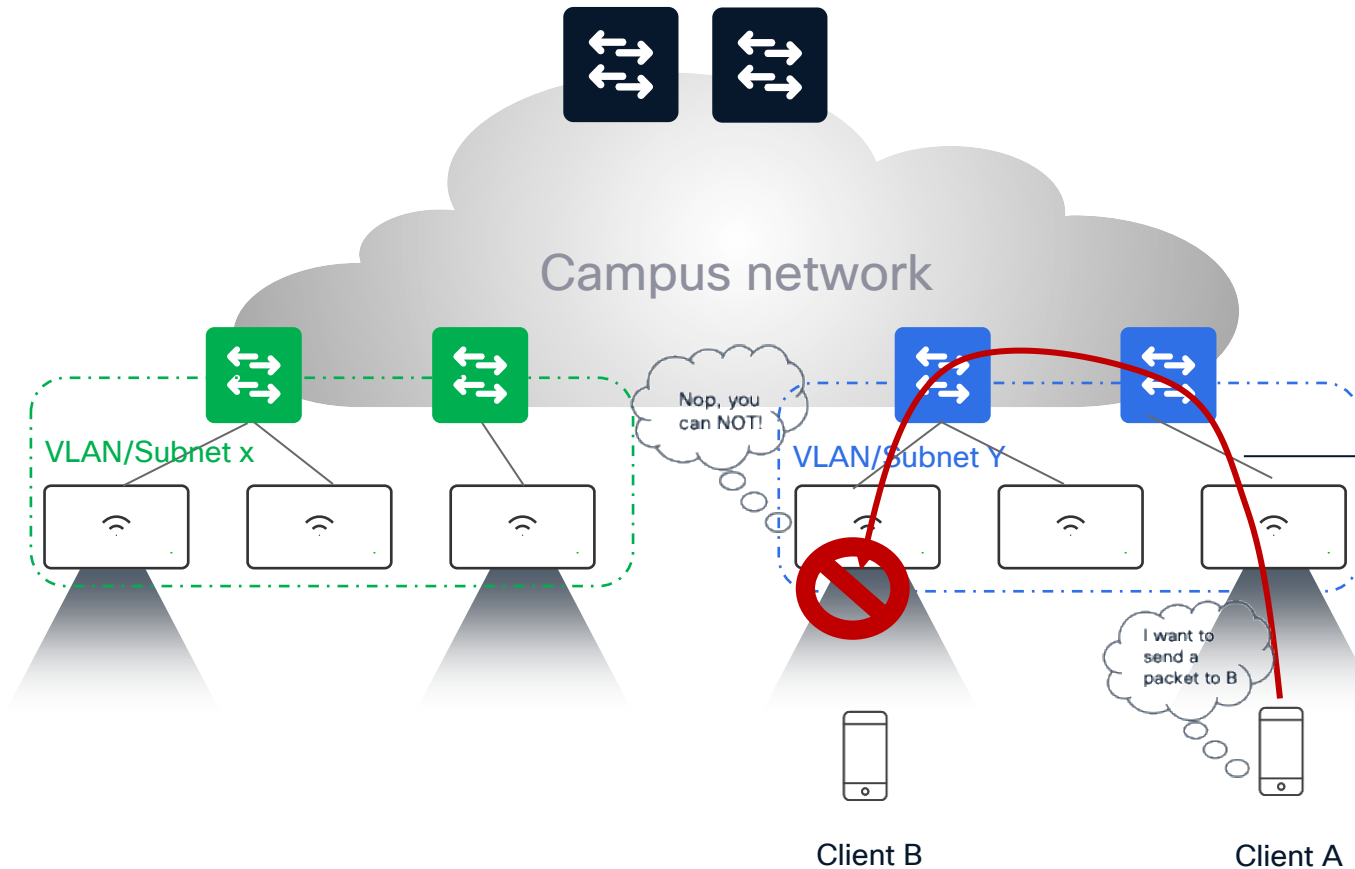


Policy



Use Cases

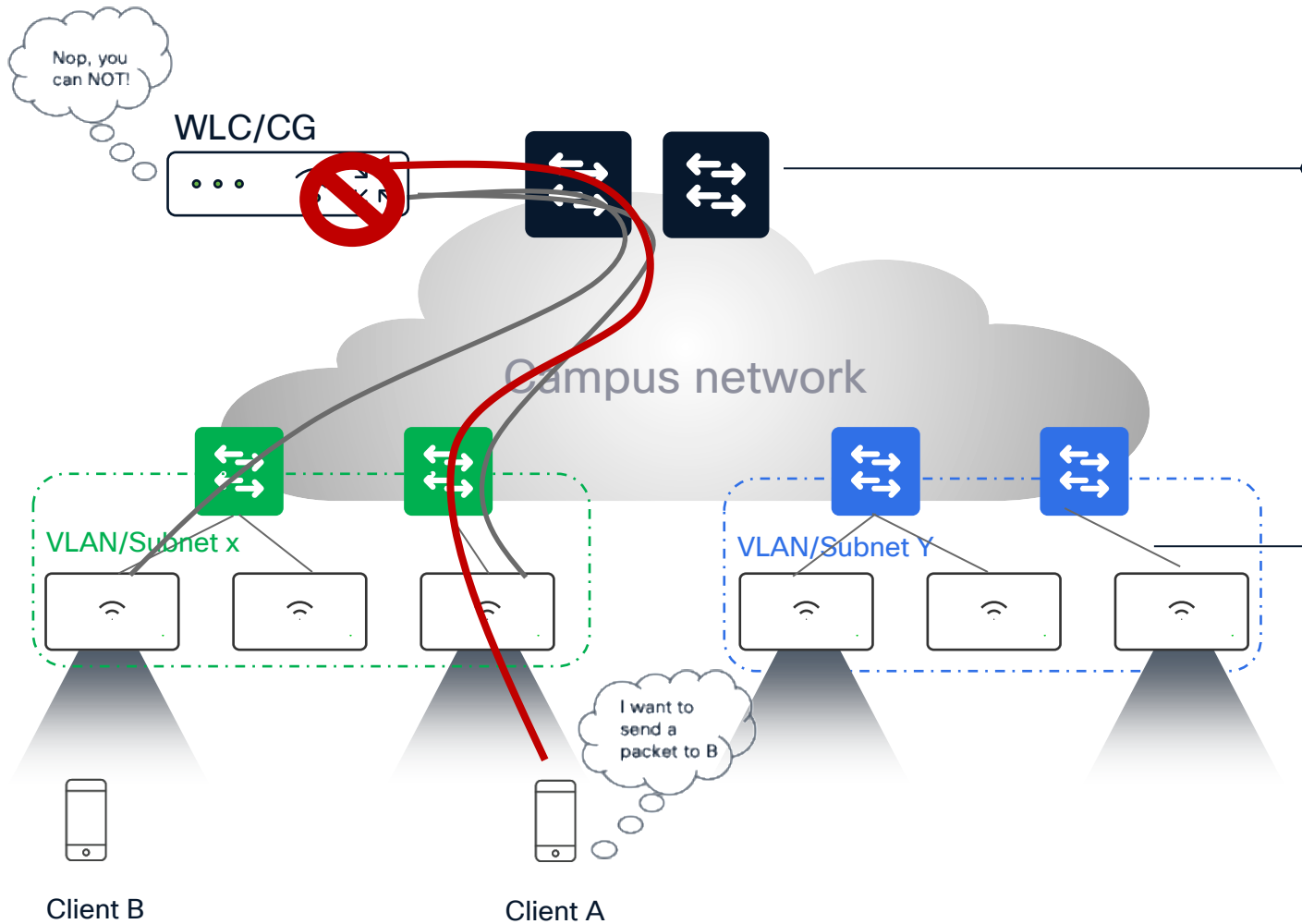
Policy Enforcement



Distributed

- Policy is applied at edge of the network: the AP (Flex/Bridge mode) or the switch with SDA Fabric mode
- For Flex and Meraki bridge mode: ACL/SGACL policies, these are applied at the AP in the Egress direction
- For Fabric mode: ACL/SGACL policies, these are applied at the switch in the Egress direction

Policy Enforcement



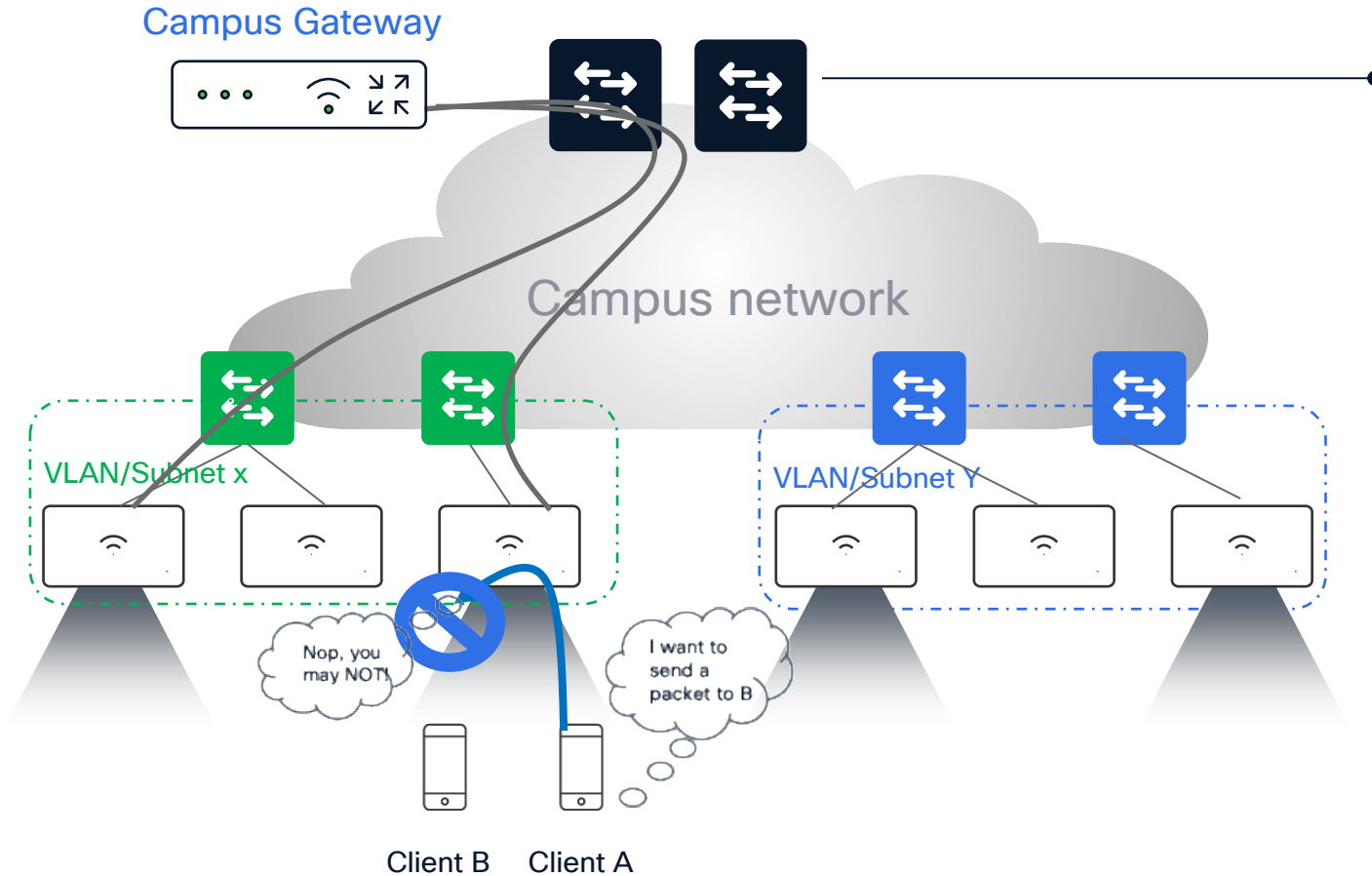
Centralized

- Policy is applied centrally at the aggregator
- For Catalyst Local or Flex Central switching mode: Policy is always applied at the WLC, even if clients are connected to the same AP

Distributed

- Policy is applied at edge of the network: the AP (Flex/Bridge mode) or the switch with SDA Fabric mode
- For Flex and Meraki bridge mode: ACL/SGACL policies, these are applied at the AP in the Egress direction
- For Fabric mode: ACL/SGACL policies, these are applied at the switch in the Egress direction

Policy Enforcement



Centralized

- Policy is applied centrally at the aggregator
- For Catalyst Local or Flex Central switching mode: Policy is always applied at the WLC, even if clients are connected to the same AP
- For **Campus Gateway**, the exception is clients connected to the same AP; in that case the policy is applied locally at the AP and not centrally

Key Design Factors



Features



Services



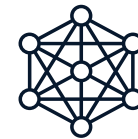
Wired Network Design



Scale/Performances



Policy



Use Cases

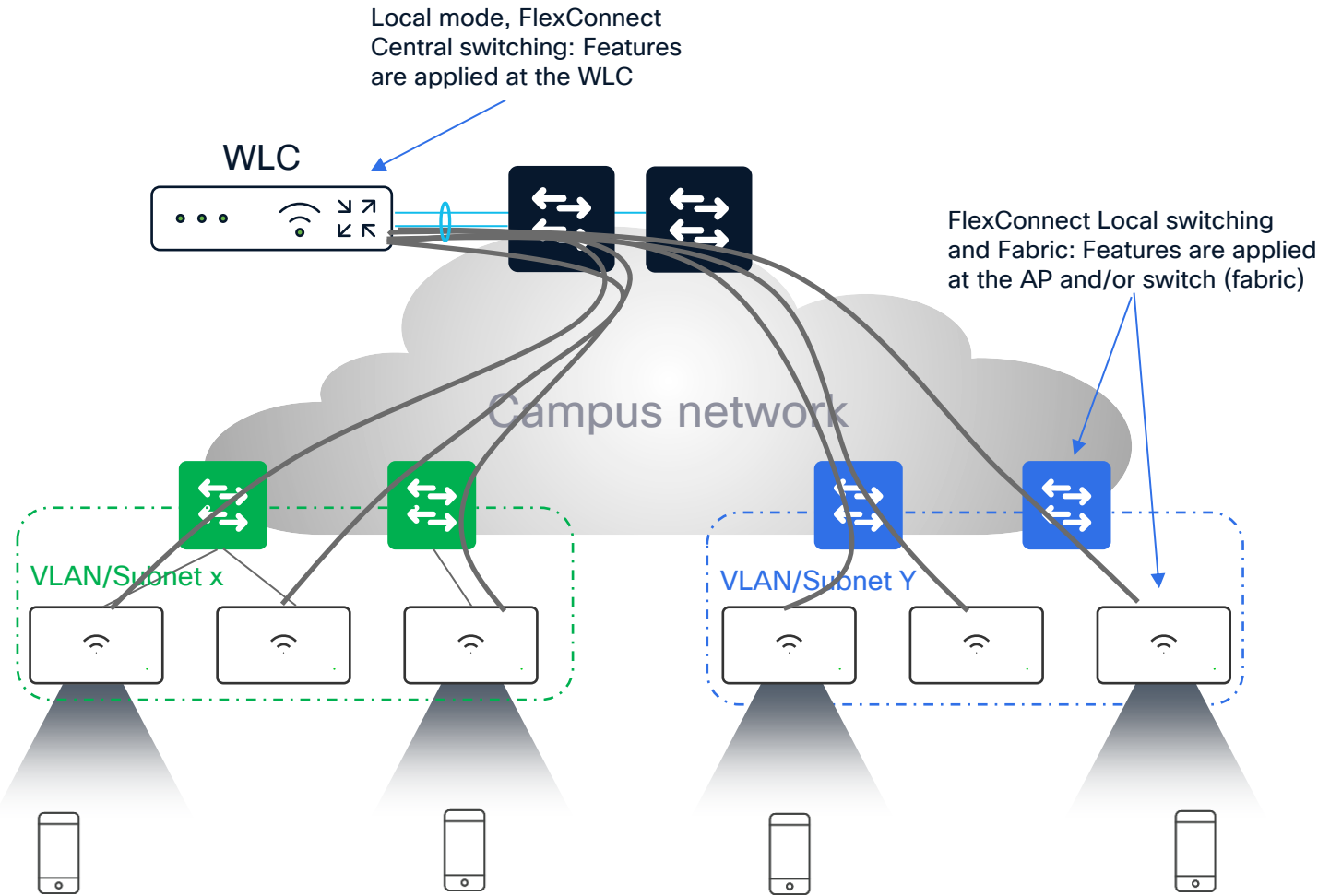
Features

What are the critical features for the customer?

Are these supported on both network architecture models?

Is the same true for On-prem vs. Cloud?

Features – On-prem

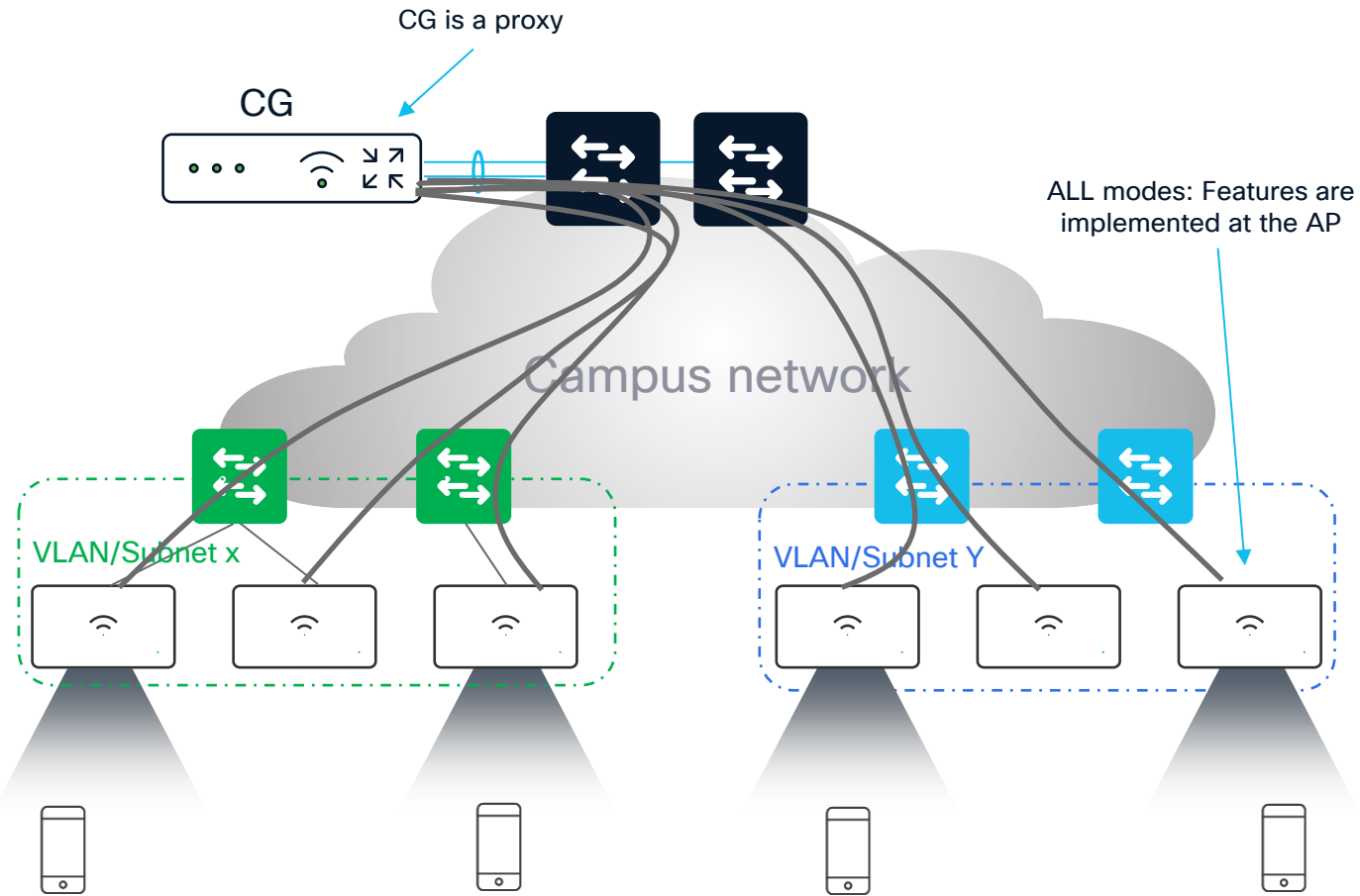


On-Prem



- Centralized (Local mode): Data plane features are applied at the WLC
- Distributed (Flex Local switching and Fabric): Data plane features are applied at the AP
- Feature difference is minimal between Centralized and Distributed but is there and Centralized is more complete
- Check the Feature Matrix
https://www.cisco.com/c/en/us/td/docs/wireless/access_point/feature-matrix/ap-feature-matrix.html

Features – Cloud



Cloud



- There is mainly no difference in distributed vs. centralized because features are implemented at the AP
- Go to Centralized with Campus Gateway if you want AAA proxy or mDNS Gateway
- As of today, Open Roaming and Access Manager are not officially supported with Campus Gateway (roadmap items)

Key Design Factors



Features



Services



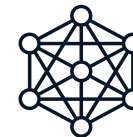
Wired Network Design



Scale/Performances



Policy



Use Cases

Use Cases

APP requirement: High BW and super low latency?

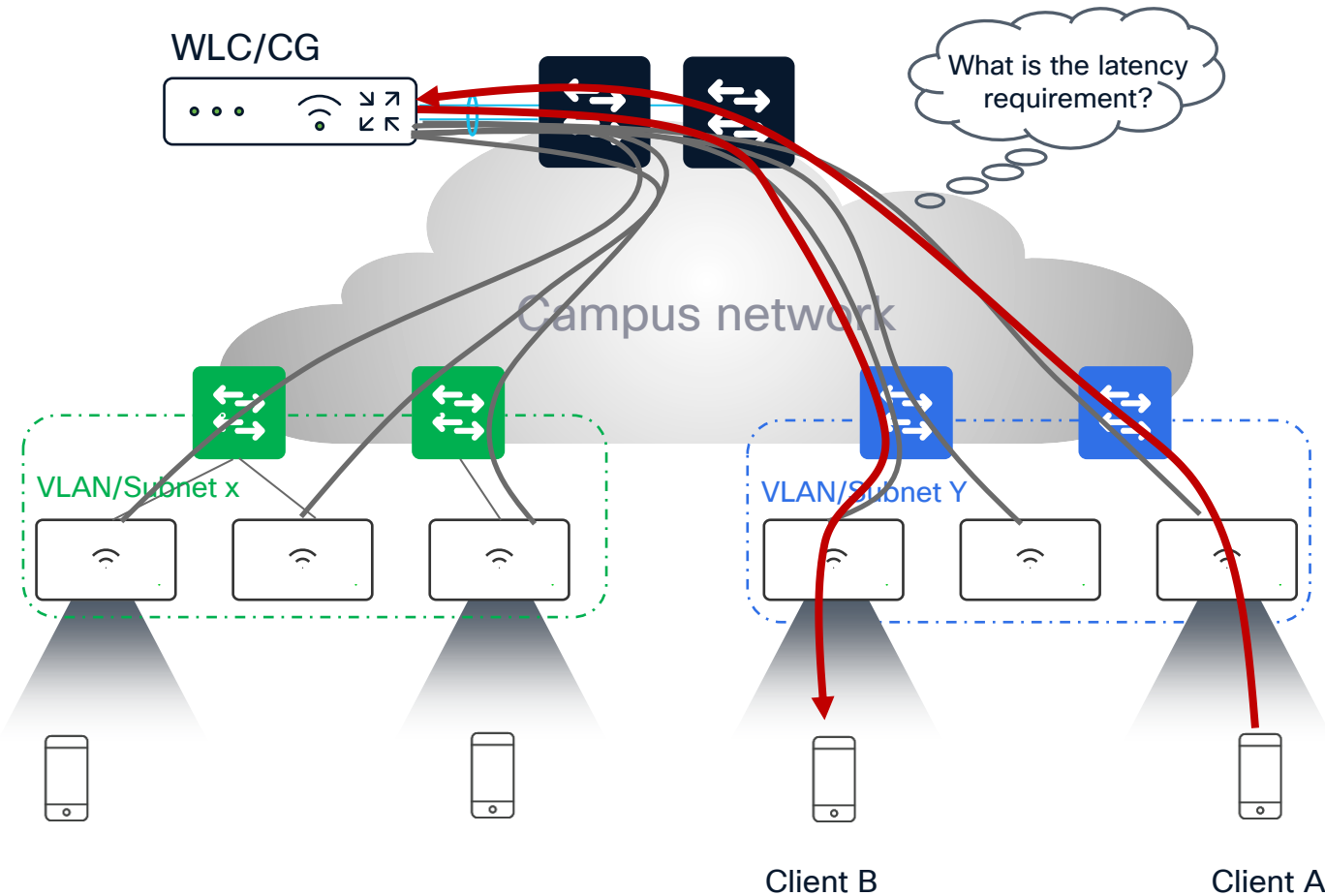
Is Wi-Fi 7 a factor for architecture choice?

Is Guest Access centralized or DIA?

Traffic paths (E-W vs N-S),

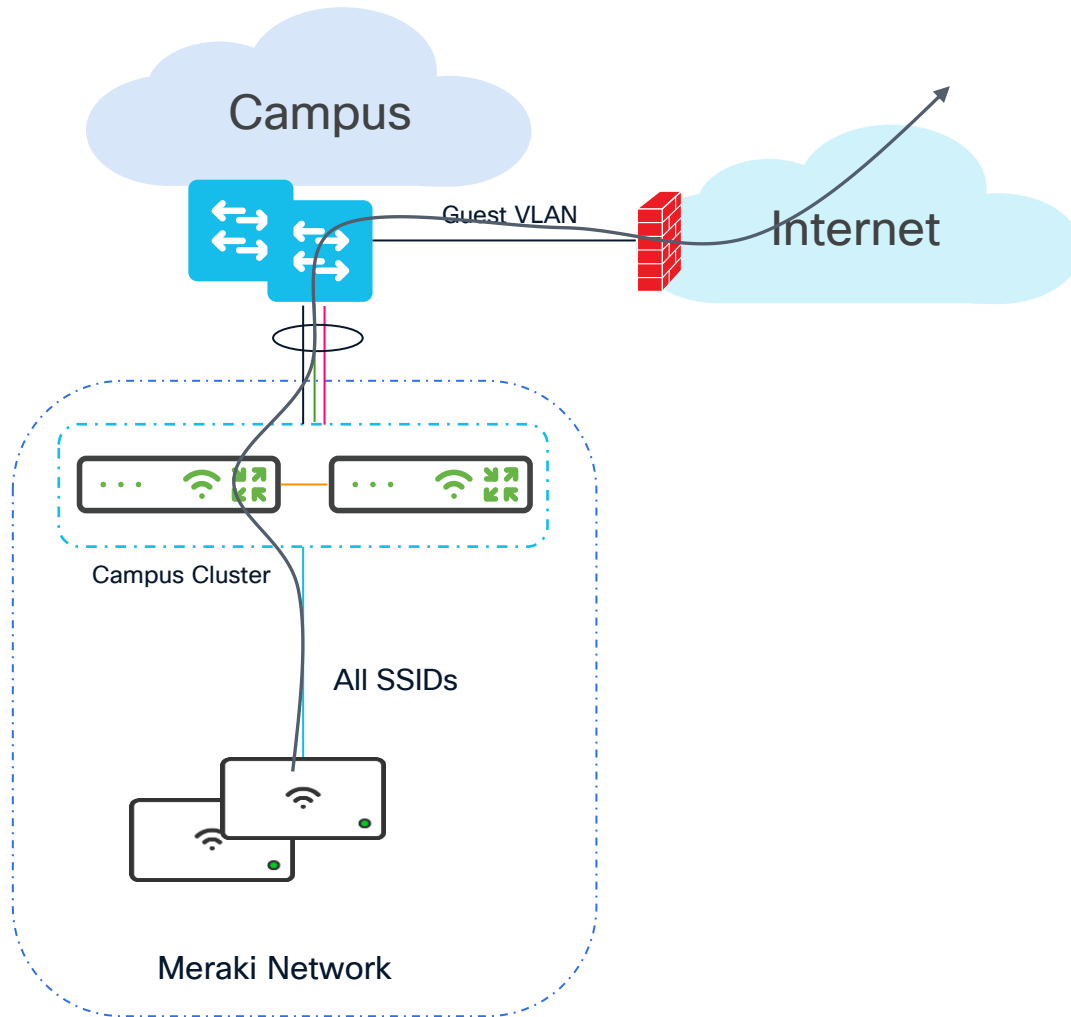
BUM traffic handling

Wi-Fi 7 and APP requirements



- You need to consider the APP requirements and the related traffic path
- Centralized designed is not optimized for East-West traffic > Traffic from client A to B (East-West) travels all the way to the aggregator/edge and back
- What is the latency introduced by the path to the central aggregator? Is this a factor?
- Wi-Fi 7 brings throughput and low latency for applications like XR (eXtended Reality) that may benefit from a distributed data plane approach

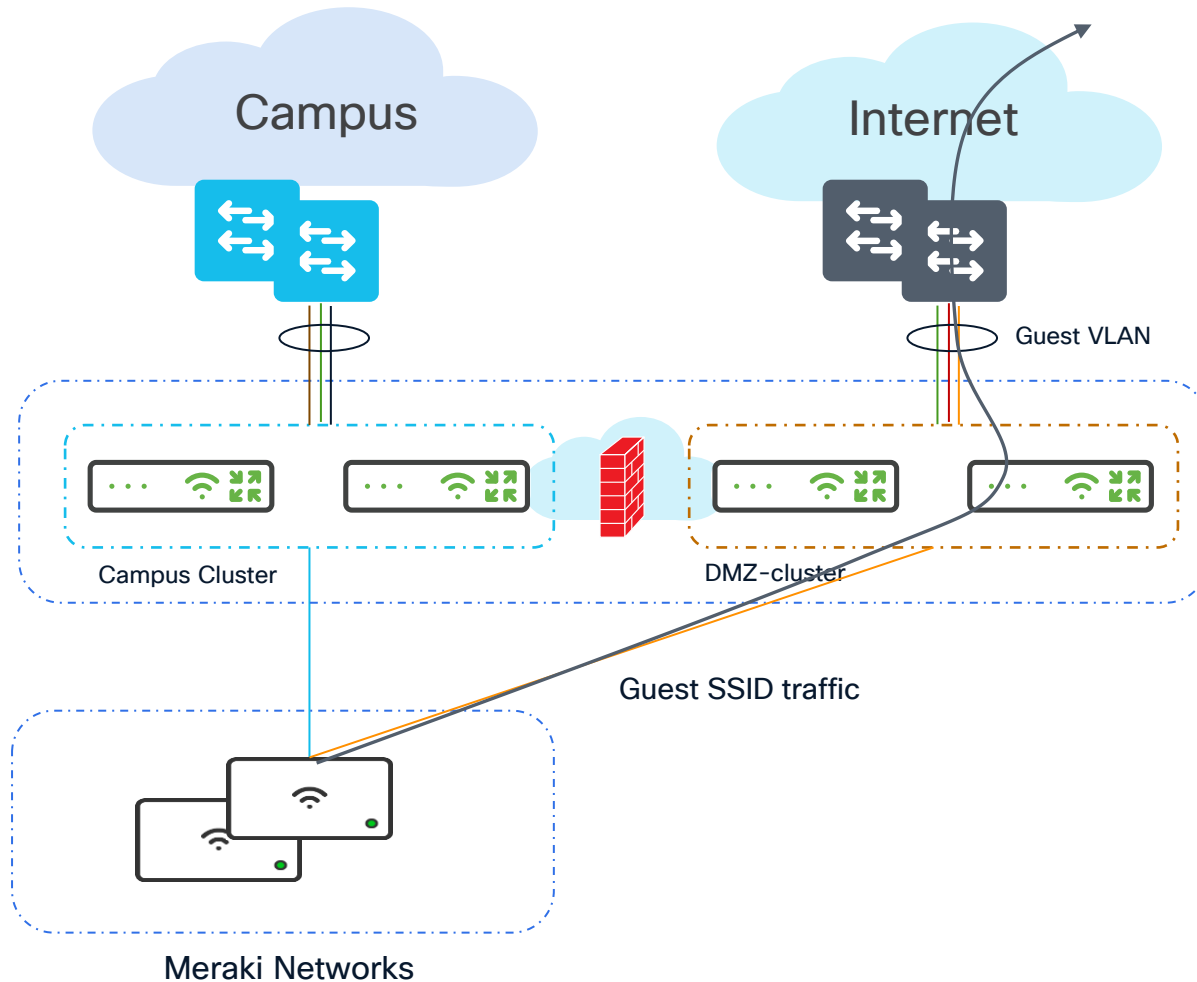
Guest Access Centralized with CG (today)



- Traffic needs to exit from a centralized point in the network (DMZ) > Centralized approach
- AP can only tunnel to one CG cluster
- All SSIDs are terminated at the same CG cluster
- Guest (or other) SSID traffic segmentation is only supported via VLAN segmentation
- Map the Guest SSID to a guest VLAN and have that terminated to a VRF interface on the Firewall, for example
- If the firewall is not L2 adjacent to the CG cluster, then you would need to transport the Guest VLAN using some wired network mechanism (VRF light, BGP EVPN, etc.)
- **Note:** Anchoring (tunneling between CGs) is not supported

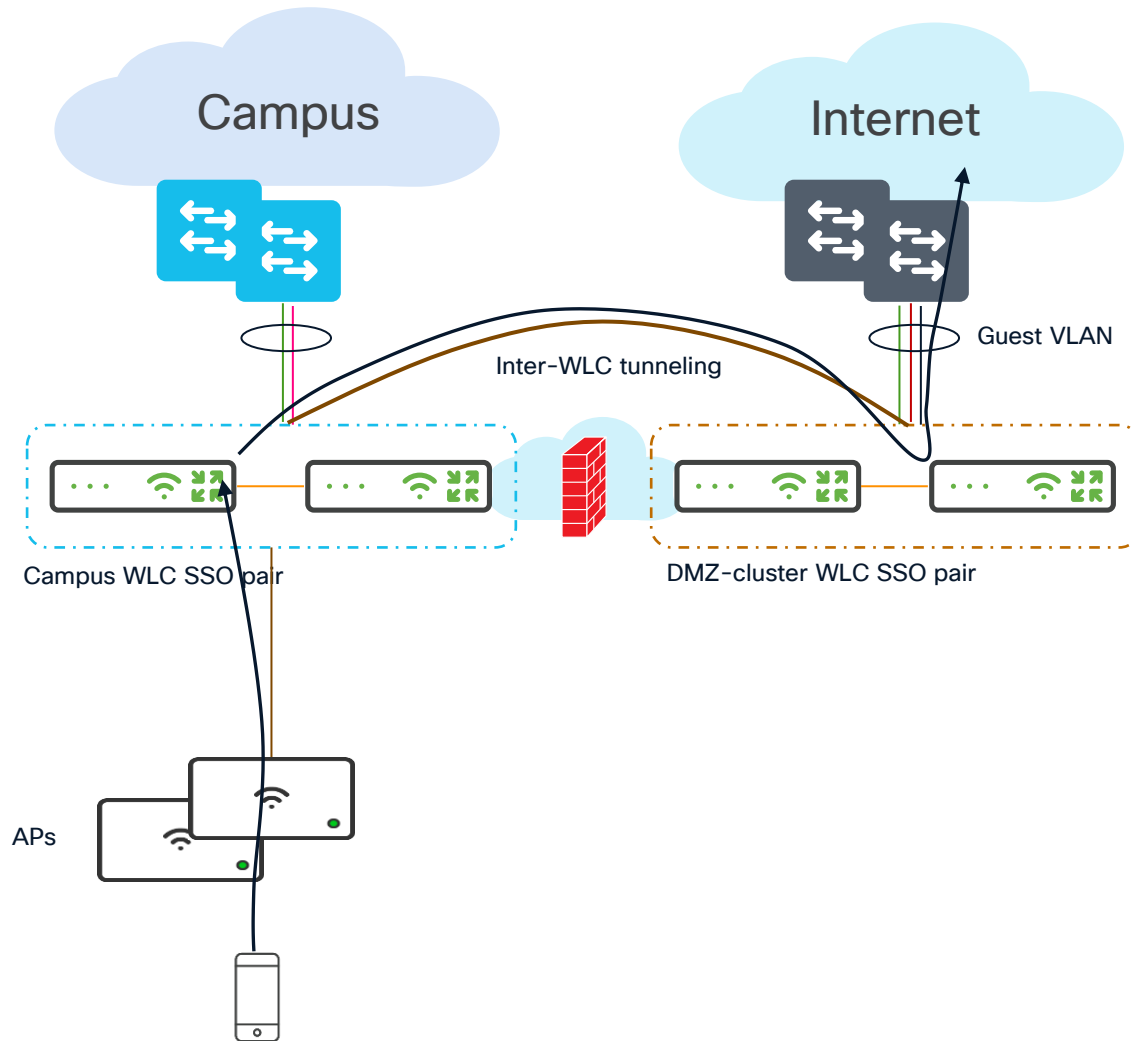
Guest Access Centralized with CG – 32.2

Roadmap
alert



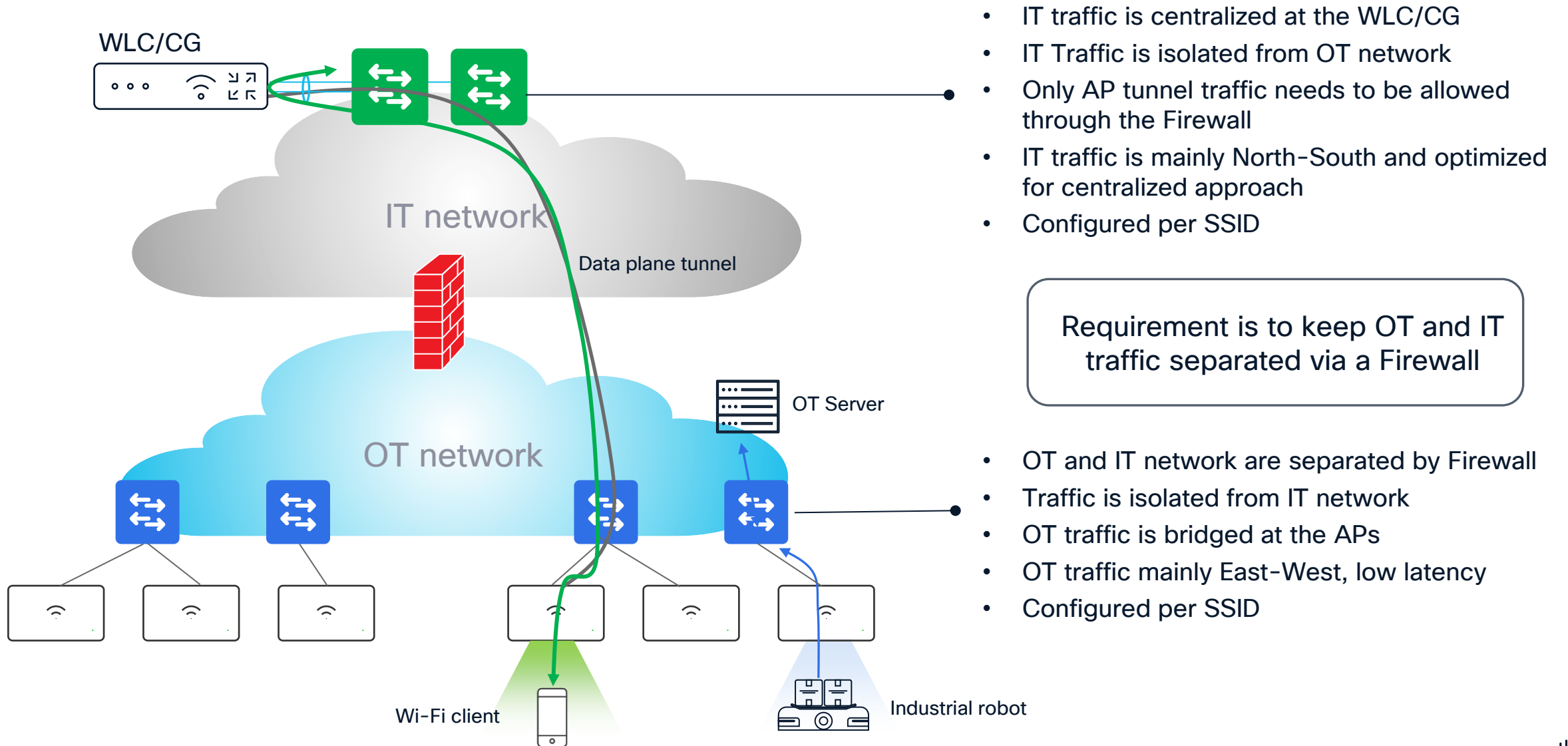
- You can have multiple CG clusters in a Meraki Network
- AP in a Meraki Network can tunnel different SSIDs to different CG clusters
- APs can be in the same Meraki Network as CG clusters or in a separated one
- A cluster in the DMZ of the firewall can be dedicated to Guest traffic
- Configure the Guest SSID to tunnel to the Guest cluster and map it to the guest VLAN
- **Note:** Anchoring is not supported

Guest Access Centralized with WLC



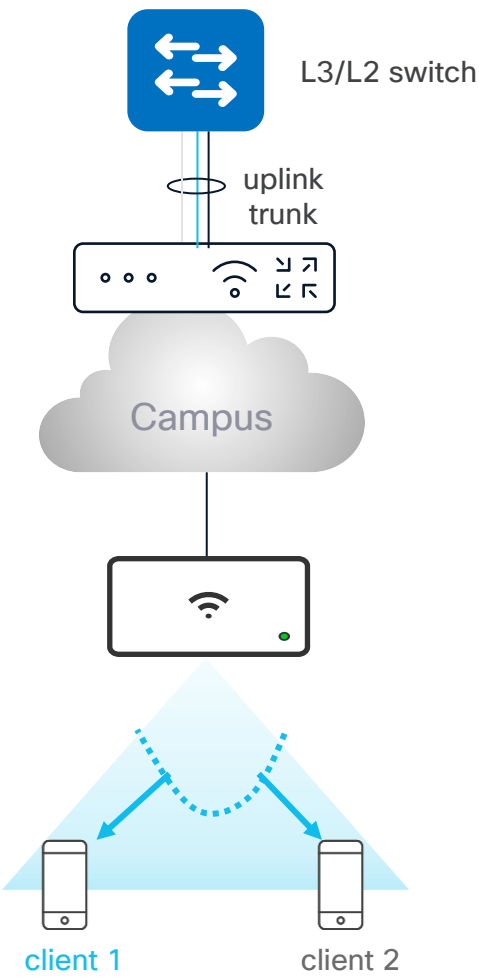
- Traffic needs to exit from a centralized point in the network (DMZ) > Centralized approach
- Guest auto-anchoring on the WLC provides a simple solution to segment guest traffic directly to the DMZ without touching the underlay network
- The WLC (Foreign) terminates CAPWAP traffic from APs, decapsulates the client traffic and re-encapsulate it into the inter-WLC CAPWAP tunnel to the DMZ-WLC (Anchor)
- The Anchor decapsulates the client traffic and bridges into the wired network

IT/OT Separation: mixing Centralized and Distributed



BUM Traffic handling - Centralized

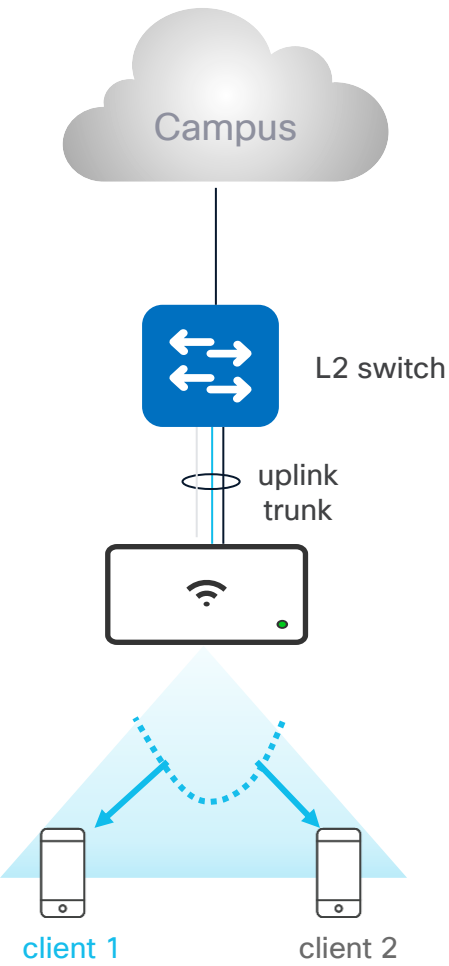
Broadcast Unknown Unicast Multicast (BUM)



	ARP for known client from wireless	ARP for known client from wired	ARP for unknown client from wireless	ARP for unknown client from wired
WLC (local mode)	Broadcast to unicast and ARP sent to client to reply. Proxy ARP, replies on behalf of client	Broadcast to unicast and ARP sent to client to reply. Proxy ARP, replies on behalf of client	Bridged onto the VLAN	Drop
Campus Gateway	Broadcast to unicast and ARP sent to client to reply	Broadcast to unicast and ARP sent to client to reply	Bridged onto the VLAN	Drop

	mDNS	Broadcast	IP multicast
WLC (local mode)	Can be configured to drop/bridge/gateway	By default Drop, exception is DHCP, mDNS active queries, IPv6 NS and NA. Broadcast can be enabled per VLAN, and WLC floods it.	By default Drop, it can be enable via IGMP snooping and MoM forwarding
Campus Gateway	Can be configured to drop/gateway	Drop only, exception DHCP, mDNS active queries, IPv6 NS and NA.	Drop only, exception is IPv6 NS and NA.

BUM Traffic handling - Distributed

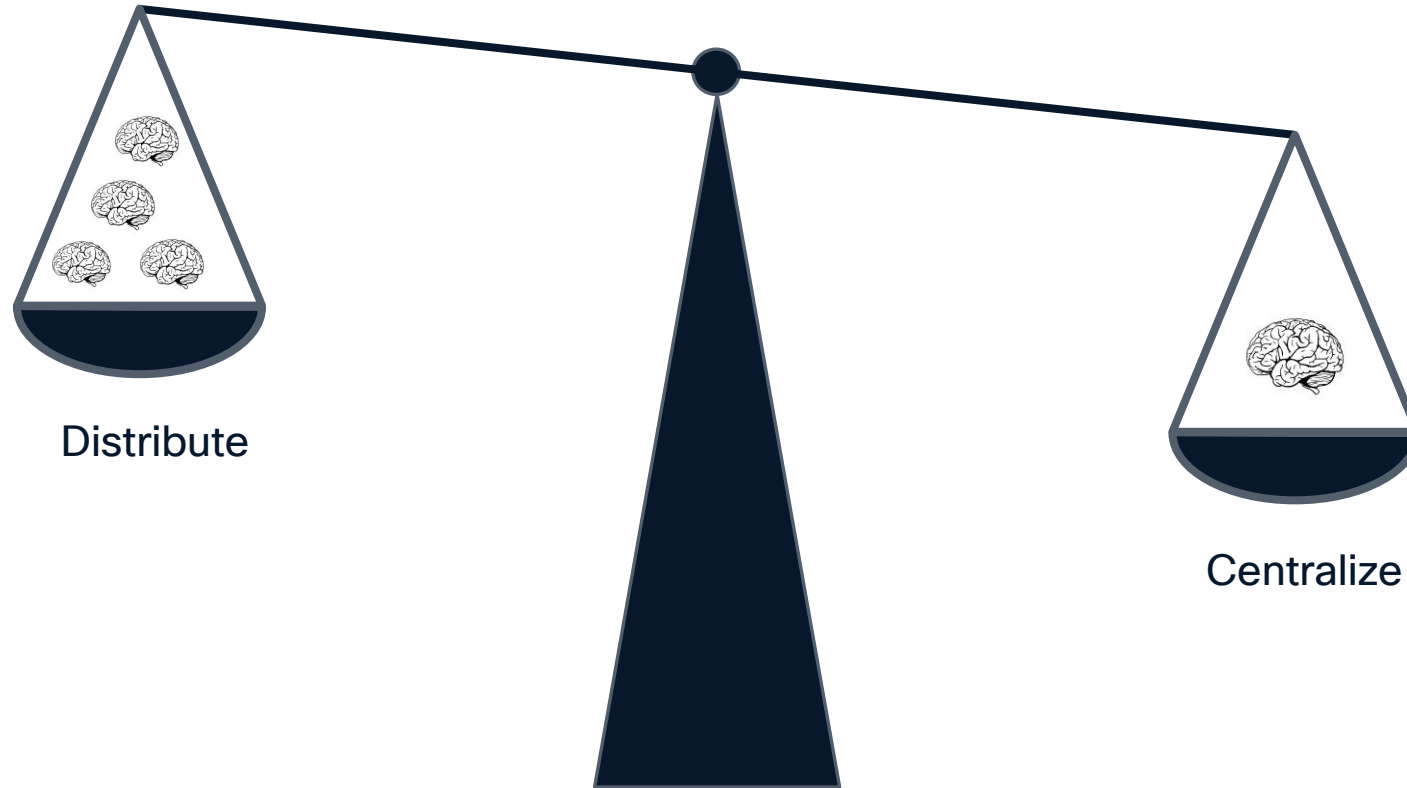


	ARP for known client from wireless	ARP for known client from wired	ARP for unknown client from wireless	ARP for unknown client from wired
AP (Flex Local switching or Fabric)	Proxy ARP, replies on behalf of client	Proxy ARP, replies on behalf of client	Bridged onto the VLAN	Drop
AP (Meraki mode)	Proxy ARP, replies on behalf of client	Proxy ARP, replies on behalf of client	Bridged onto the VLAN	Drop

	mDNS	Broadcast	IP multicast
AP (Flex Local switching or Fabric)	Can be configured to drop/bridge/gateway	By default floods, can be configured to drop	By default floods, can be configured to drop
AP (Meraki mode)	drop/bridge/forward (across VLANs). mDNS Gateway feature not available	By default floods, can be configured to unicast	Configurable to perform a multicast-to-unicast packet conversion using the IGMP protocol

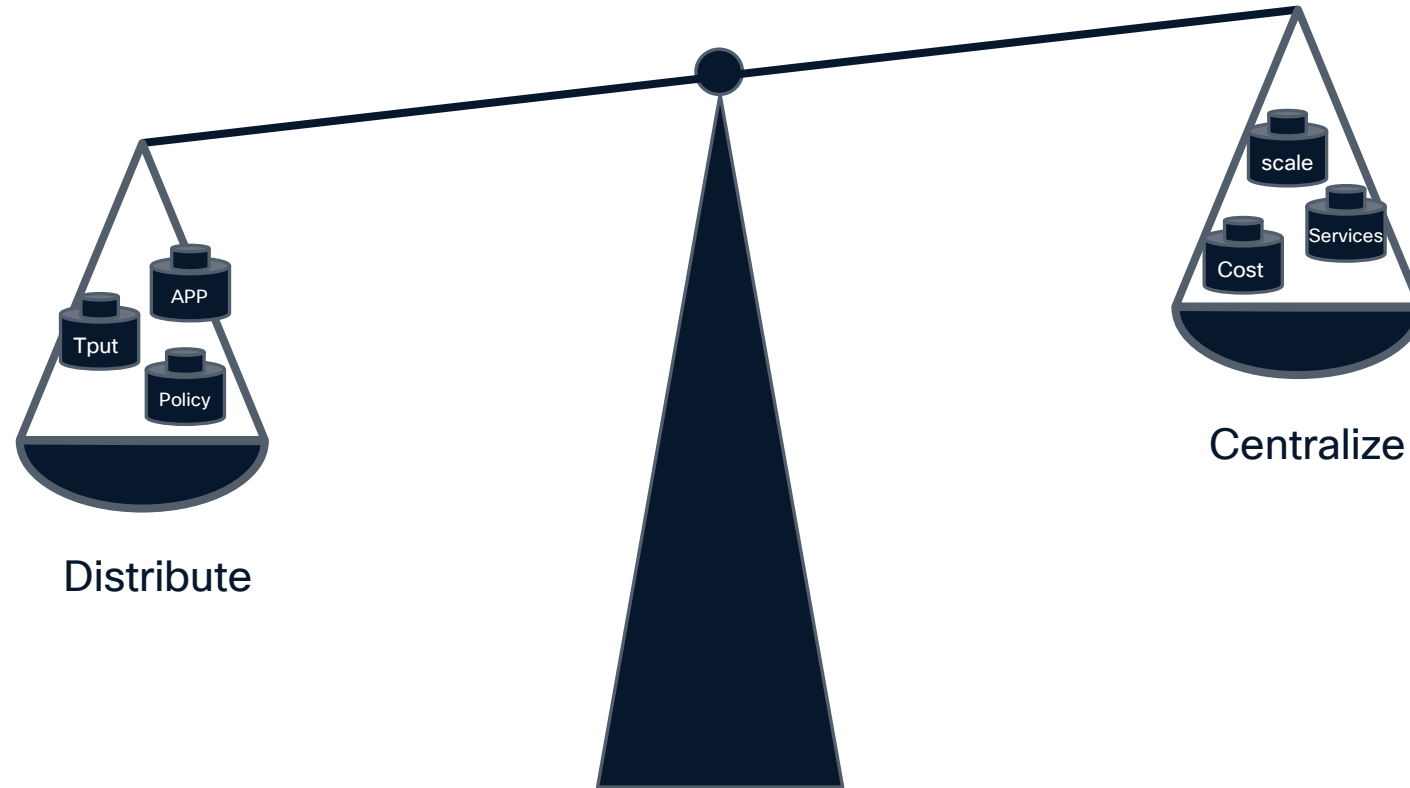
Conclusions

It's a balancing act...



Where do you distribute the intelligence?

It's a balancing act...



You need to consider all the design factors

Summary of Design Factors

Distributed	Centralized
Guest (Direct Internet Access)	Centralized Guest services
Roaming at lower scale (unless Fabric is used)	Roaming @ scale
Scale limited by L2 access switches	Scale limited by distribution/core switches
Need to consider VLAN design (unless Fabric is used)	Simplified VLAN/DHCP design
Multiple entry points into the wired network (multiple touch points for config and policy)	Single entry point to the wired network, single policy enforcement point
Services need to be distributed (AAA, DHCP, etc.)	Services need to be Centralized (AAA, mDNS, etc.)
Suitable for IT teams that manage wired and wireless	Suitable for siloed IT teams (less changes needed on wired, easier to get changes done)
\$ (less equipment to deploy and manage)	\$\$ (more CAPEX and more devices to manage)
Optimized for East-West traffic	Optimized for North-South traffic

Complete your session surveys



Complete your surveys in the Cisco Events App.



Complete a minimum of 4 session surveys and the overall event survey to receive a unique Cisco Live t-shirt.
(from 11:30 on Thursday, while supplies last)

Continue your education



Visit the Cisco Showcase for related demos



Book your one-on-one Meet the Engineer meeting

Visit the Technical Solutions Clinics to discuss your technical questions



Attend the interactive education with DevNet, Capture the Flag, and Walk-in Labs



Visit the On-Demand Library for more sessions at CiscoLive.com/On-Demand

Contact me at: siarena@cisco.com

Thank you

CISCO Live !

