



Cloud-Native Control Plane Design for Modern Network Automation Systems

Manevannan Ramasamy

Software Engineering Senior Technical Leader, Milpitas, California, United States of America

manevannan3@gmail.com

ABSTRACT: The modern network automation is increasingly demanding cloud-native orchestrated networks capable of dynamically delivered workloads, distributed functions, and heterogeneous networks. However, containerization and orchestration are insufficient to create a reliable network control plane because changes in configuration are long-lived, partially observable, and frequently implemented on devices that do not have a transactional updating semantic. This study proposes a cloud-native control plane that incorporates a declarative API, long-term desired-state storage, autonomous reconciliation, workflow coordination, workflow adapters, and an event-driven observation plane. The proposed architecture provides reconciliation as a work model that can maintain the convergence of the desired and observed states of the network in real-time and support idempotent processes and recovery of failure. Mechanisms of consistency and clear ownership boundaries are established to minimize conflicting updates, duplication of actions, and propagation of the stale state. The automation workloads are further kept in check by back-pressure mechanisms with high request volumes and dependency failures. The framework is analyzed conceptually in the workload scaling, dependency malfunctioning, lifeless conditions and in executing scenarios, to investigate its resilience and functioning traits. The analysis shows that network-specific safety controls, the history of workflow permanence, and the reconciliation of workflow state are needed as first-class architectural elements in order to accomplish effective cloud-native network automation.

KEYWORDS: Cloud-native, Control Plane, Network Automation, Reconciliation, Kubernetes, Distributed Systems

I. INTRODUCTION

Cloud-native architecture has stopped the distributed application structure utilizing containerized applications, stateless resource description, automation coordination, and elasticity to achieve scalability and continuous conditioning. Even though such technologies are usually related to the deployment of applications, the architectural principles of these technologies are especially pertinent to the network automation of the modern time. The key question in a network control environment is not an easy question in terms of packaging automation services, but rather one of ensuring a consistent relationship between the desired network structure and the state that is currently being seen amongst a heterogeneous population of devices. Kubernetes controllers was an effective conceptual foundation consisting of identities of reconciliation whereby controllers compare the wished and realized status constantly and engage in corrective measures to reduce the discrepancy amid the two [1]. This solution fits quite well in a network environment where configuration drift, unavailability of devices temporarily, slow responsiveness and asynchronous operation is a way of life.

Classical network automation typically adopts the request response model where a configuration request is sent to an outside program, where a device-specific API or protocol is invoked, and where command execution is successful is simply defined to be network state transition successful. In a distributed environment, such an assumption gets even more problematic. A device may respond to a configuration request asynchronously by applying it or give incomplete information, or fail to respond, before the end state can be discovered. The other similarity is that a controller is able to abort when it provides a command, yet before it proceeds after the status of the respective execution. Network events can be replicated, delayed or lost as well. Moreover, the advent of microservices or containers is not assured to provide reliability, consistency and recoverability that is required within the control of the network.

There are also short-term stateless application requests which are very different to network configuration operations. A network change may involve a network of several devices, dependencies, validation, retries, approval limit and validation of changes. Some of these workflow may require several minutes or even hours when running since in case of failures of their respective services, the workflow should not be lost. In addition, different network devices



demonstrate dissimilar capabilities, design models, administration windows and also functionality limitations. A cloud-native Control plane therefore must have a tool to stored the history of workflows, learn (what) devices can do, coordinate the dependent actions and keep an eye on the resulting state.

The other area of significance is the difference between the desired state, the execution state as well as the observed state. Desired state is what the operator or the automation system would like to see (i.e., what they want to accomplish) whereas execution state corresponds to the amount of actions taken in regards to this desired state. Observed state is data which has been scanned and scanned by network devices and monitoring systems. The fact that the states can be considered the same might contribute to erroneous expectations concerning a configuration that will arrive at convergence. A healthy control plane should then ensure that there are clear consistency viewpoints.

This paper introduces a cloud-native control plane which has been conceived in accordance with the network properties. The architecture separates the resource declaration, persistent desired-state control, reconciliation, workflow coordination, device adaptation, and observation into separate but coordinating fragments. Declarative APIs bring a standardized interface by which it is possible to express network intent and durable state storage those configurations and other information related to workflows in the case of service failures. The independent reconcilers never stop reporting discrepancies between the desired and observed conditions taking remedial action. Workflow coordination is independent of the concrete interactions of particular devices, but orchestrates operations of many stages and dependencies.

Capability-sensitive adapters provide an interface between a generic control-plane implementation, and heterogeneous network devices. The adapters reveal device capabilities, and the control-plane actions are converted to the appropriate device-specific operations, instead of making the assumption that all devices do exactly the same operations. An event-driven observation plane supplements this mechanism with generating state changes, operational events, and verification information. A combination of all these components will enable uninterrupted feedback between the network conditions and control decisions.

The ability to perform idempotent operations, explicit ownership, back-pressure mechanisms and controlled retry behavior are even more reliable. Idem potency will minimize the impact of an overlapping request, whereas ownership limits can guarantee that overlapping controllers do not end up modifying the identical resources without consensus. Uncontrolled propagation of the workload exists in the back pressure in the form of unavailable services or devices that are used. Such mechanisms are especially necessary since network automation has to be safe even when operated partially as opposed to assuming that any given workflow is successful or unsuccessful.

The architecture proposed is therefore focused on reconciliation as opposed to the one time command execution. It has its appraisal based on the scaling of workload, dependency failure, stale state, asynchronous behavior and partially executed state to check the resilience of architectures and consistency of their behavior. The work identifies the following properties of cloud-native network automation as important control-plane properties: network-specific safety mechanisms, ability to store workflow history, continuity, ability to store workflow history, and ability to continue state reconciliation. A combination of these principles can make the proposed design provide a system of network automation that is scalable, resilient, and observable.

II. CLOUD-NATIVE RECONCILIATION CONTROL PLANE ARCHITECTURE

The proposed control plane is a collection of decoupled functional layers which when put together implementable control and observable transitions between high-level network intent. The architecture tries to separate the expression of the resources, long-term state, the reconciliation, the orchestration of workflows, adapting devices and the observation. This separation serves to not only package operator-facing services and heterogeneous network infrastructure, but also to provide each individual part to scale, or recover independently.

2.1 Declarative API and Admission Layer

The API layer provides declarative and versioned resource models of network outputs such as site configuration, policy implementation, software compliance, connectivity services, and other assets managed. The API itself actually is a reflection of the desired state at the level of the form of structured resources instead of the exposure of low-level device commands. All the accepted resources are specified, generated, have metadata and lifecycle information.

There is a schema validation, an authorization check, policy check, an invariant check and a capability check performed in an admission process followed by persistence. This immediate verification removes structurally invalid or outright



unfounded requests which could sometime get into the control system. Admission should be deterministic and computational lightweight in the sense that the core of admission is to make a decision on whether a request is admissible. Components and complex planning as well as dependency analysis are deferred to asynchronous control components.

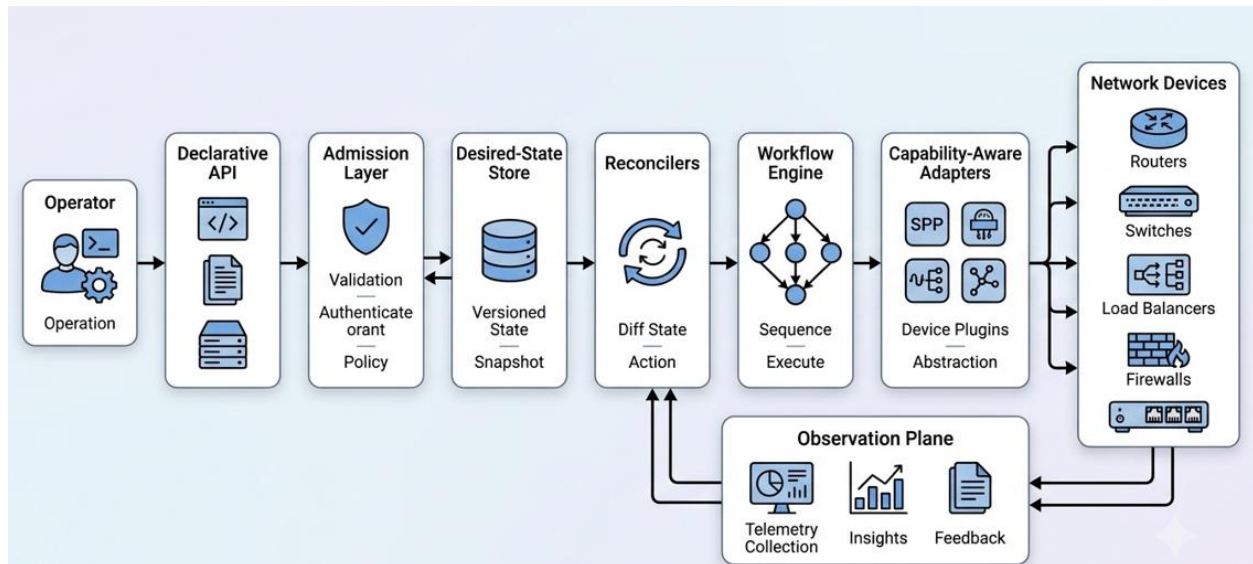


Figure 1: Cloud-Native Network Control Plane Architecture

It also has an evolution mechanism in the control plane with versioned resource definitions. New resources can be created with new features, without breaking the current consumers. Generation identifiers help identify (to the controllers) a new requested configuration and an old one and apply a sound platform to which they follow-up the progress of the reconciliation.

2.2 Durable Desired-State Store

The desired-state store is the desired representation of the desired state of the control plane. It needs to be long-lasting because the controller operations could be resumed, distinct services might not be online and network operations could be occurring in the event of the service failure. The requested configuration must then be resilient to failure, but not necessarily to be handoff by the lifecycle of the service on which it was initially handed down.

Resources updates should be based on optimistic concurrency control. Version or generation checks are used to guarantee that two actors acting independently do not happen to overwrite each other. The control plane can on competitive updates observe the conflict and explicitly require an explicit resolution, rather than take over the last writer.

Status is what is happening at present and specification is what is desired to be the outcome. Status can be other conditions like Accepted, Planned, Progressing, Degraded or Satisfied. This isolation only allows the controllers to report about progress in the operations to the operator without distorting his original plan.

2.3 Ownership-Based Reconciliation

The control plane proposed, which is the focus of the operating model, is the reconciliation one. Reconcilers can be separated based on resources ownership, which enables the separation of various resource types to be managed separately. Reconciler reads the desired state which is persistent, recovers suitable observation by calculating difference between desired and actual condition and derives a finite next action.

Where possible a reconciliation cycle is supposed to be deterministic and must be idempotent since retries in a distributed system are anticipated. When the same request of reconciliation is done more than once then it should not inadvertently make changes in the network twice. The controllers should maintain a reasonable state that will determine whether a certain action is initiated or not.



This type of categorization of errors is also essential. An explicit failure state where the request needs to be fixed should be caused by permanent errors, e.g. invalid configuration, unsupported capabilities. Controlled retries with back off should be made in situations of transient errors, including temporary connectivity loss or unavailability of service. This discrepancy will make sure that the failures of the stable retries will not generate the redundant traffic, but rather it will progress the successful recoverable errors to converge by themselves.

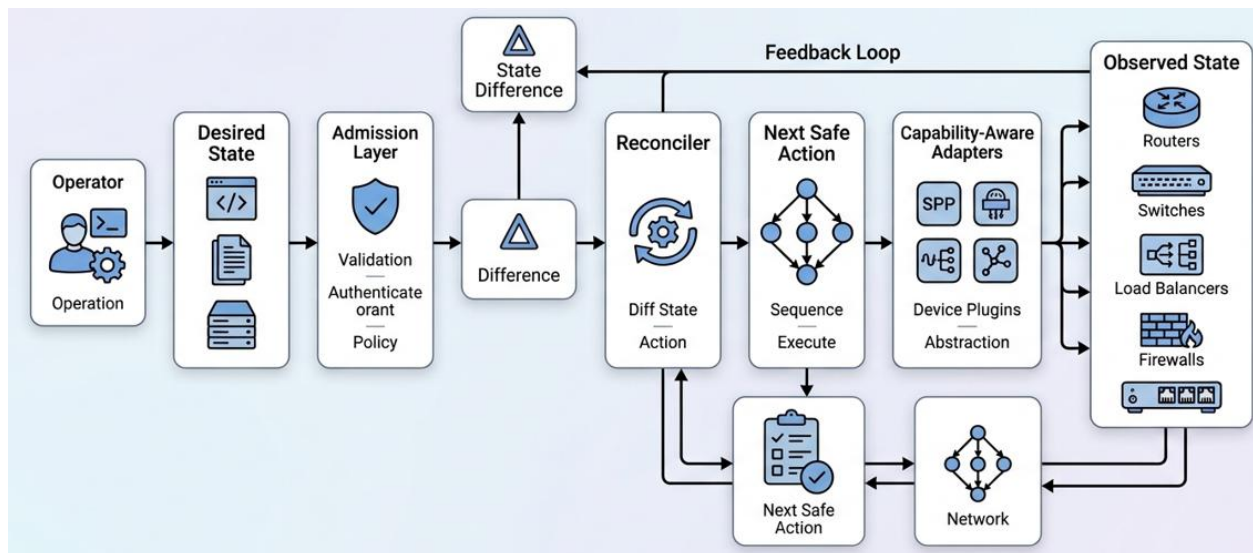


Figure 2: Desired-State Reconciliation Control Loop

2.4 Durable Workflow Coordination

Multifaceted network mutations can be very many sequenced acts and dependent. Here, one may add a durable workflow engine to the control plane and records the execution history irrespective of different controller processes. A workflow step has the option to retain its precondition, external operation identifier, execution result and compensating action which can be used.

In extremely cases where uncertainty exists as to an external operation, durability is a huge issue. A request can be given by a controller to a device and an abort can be made prior to the response being received. There is a possibility that it could lead to an undesirable re-action after carrying out an operation again, which is blindly repeated upon completion of the recovery. To ascertain the next option, the system can first identify once more the previous option will be saved in the history of the previous workflows and external operation identifiers since it will hold the previous operation that has been already taken by the system before proceeding to the next action.

Reconciliation is not done away with but replaced by workflow coordination. The choice of whether to work is made by the reconciler and it is the workflow that establishes a way in which a selected multi-stage process should be conducted. This difference makes workflow state to be an alternative to the constant assessment of the real network conditions.

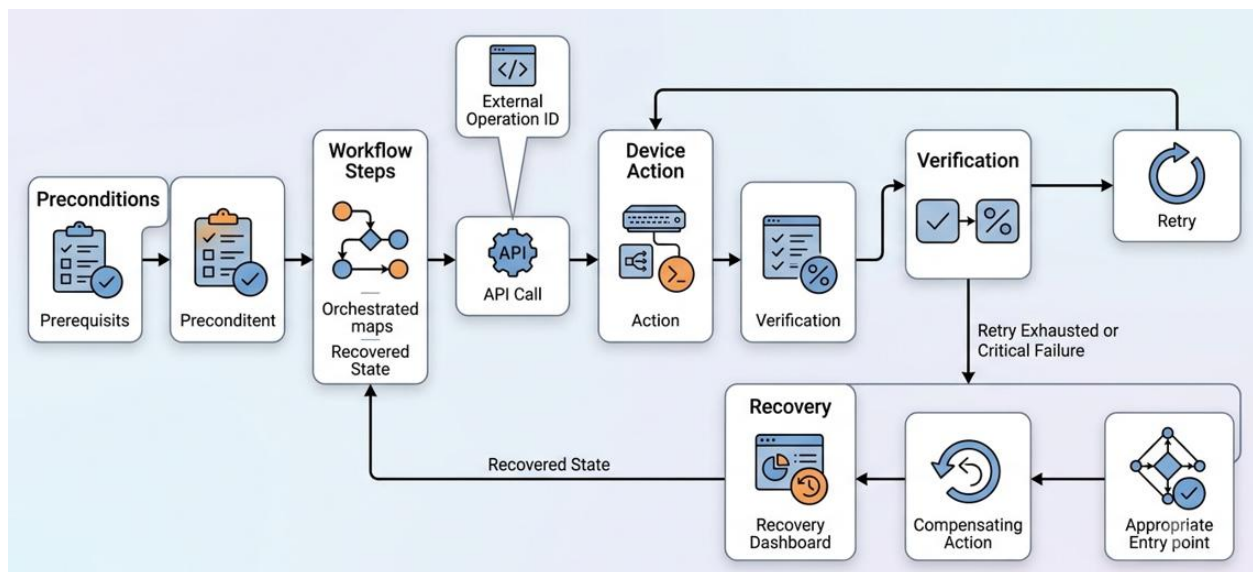


Figure 3: Durable Workflow and Recovery Model

2.5 Capability-Aware Adapter Layer

The adapter layer provides a separation of high level control logic and device-specific management mechanisms. Network infrastructure may expose NETCONF, RESTCONF, command-line interfaces, vendor APIs or domain-specific controllers and each can have different transactional semantics, error formats and configuration capabilities [2][3].

Adapters translate shrinkable control-plane actions into appropriate target-specific actions. They also publish capability descriptors which describe features that they support, like candidate configuration, atomic commit, rollback, transactional behavior or asynchronous task execution. The elements of planning can thus select procedures with regards to the actual device capabilities rather than making assumptions regarding the uniformity of behaviour across the infrastructure.

Sensitive security data, like device credentials, certificates and transport-specific data, does not cross the adapter boundary or the infrastructure-management boundary. The higher level services then load with the abstract resources/capability information without detailing the irrelevant connection information.

2.6 Event-Driven Observation Plane

Observation plane is the source of the feedback needed to do continuous reconciliation. It accumulates configuration state, topology data, telemetry, task results, reachability data and operational events of managed infrastructure. These inputs are standardized into common events representations and they are used to maintain materialized views that can be accessed by reconcilers.

At-least-once event delivery can be taken when the consumers are configured to safely handle a duplication of events. The consumers can isolate the recurrent and real new observations values in the transitions in the state with the help of event identifiers, version number of resources, and timestamps, and ordering information. This is the phenomenon of idempotent event processing thereby eliminating architectural valuable delivery cost.

Also there should be some difference between not being detected and detected failure in the plane at which the observation occurs. A lost telemetry update is not always an indicator a device or service is not healthy since the mechanism of collecting it may be offline. Normalized state ought therefore to be supplied with collection status, freshness indications and observation timestamps. This assists the control plane to avoid making incorrect choices as a result of silence.



2.7 Integrated Control Loop

The key thing is that the whole architecture is based on an endless control loop: the declarative intent will be processed into the API layer, admissibility will control the possibility to alter the desired outcome, durable storage will ensure that the desired result will be maintained, and ownership-based reconcilers will outline what should be done. Similarly, the synchronization of complex processes is achieved with workflows, the transformation of processes into operations to the target is achieved with adapters and the provision of state back into reconciliation is achieved with an observation plane.

This hierarchical structure creates definite boundaries of consistency and lets the failures confined to local proportions. Better yet, it sparks network automation of running converge rather than once execution of commands,

III. DISTRIBUTION AND CONSISTENCY

The distributed character of a cloud-native network control plane demands explicit solutions to deal with resource ownership, consistency, assignment of workload, as well as isolation of failure. Such mechanisms are employed to coordinate horizontally scaled services in a way that it can avoid transient infrastructure conditions that can lead to conflicting or unsafe network operations.

3.1 Resource Ownership and Consistency Boundaries

One logical writer can always only write to a managed resource at any given time. exclusive ownership may be provided using the techniques of leader election, partition leases, or other coordination methods between controllers replicated across multiple computers. The information on ownership has to be durable though it has to be capable of withstanding rescheduling of the processes and recovery of service. The records of desired-state store and critical invariance, therefore, require the strong consistency to guarantee that the incompatible decisions taken by overlapping controllers can never take place.

Not all the data categories can have the same level of consistency. Eventual consistency can be applied to high-volume observations (e.g. telemetry) the freshness of which is explicitly maintained. Each observation has to include an age or a time, which will allow the reconcilers of a particular action to determine whether the information can still apply to a particular action or not. Any high impact network change should be done by the reconciler when they are certain that all the observations required are within established safety limits. And it is only stale information that can cause one to pick at a slow action or slow action that is under control rather than make a premise that the network has not altered.

3.2 Scalable Work Distribution

Horizontal scaling assumes predictable split of the reconciliation workloads. The distributable resources can utilize stable ownership keys, including site, device, and service or resource identity. The work queuing in each partition must be done and the update frequencies consolidated so that there can be more than one update of the same resource but one processing cycle will be unnecessary.

Concurrency limits can be sensitive to network-facing operations especially. Limits of individual devices, sites and of groups of resources can be configurable by control plane. This is implemented as follows controls that will guarantee that the occurrence of high volume of simultaneous operations will not cause the unloading of the network equipment or down-stream services.

The propagation to the back-pressure should not be within internal queues but across the control plane. The Admissions and conditions to the operator should reflect the delaying requests on capacity constraint, dependency failures as well as observation overload. This offers operational visibility and avoids abrupt telemetry bursts.

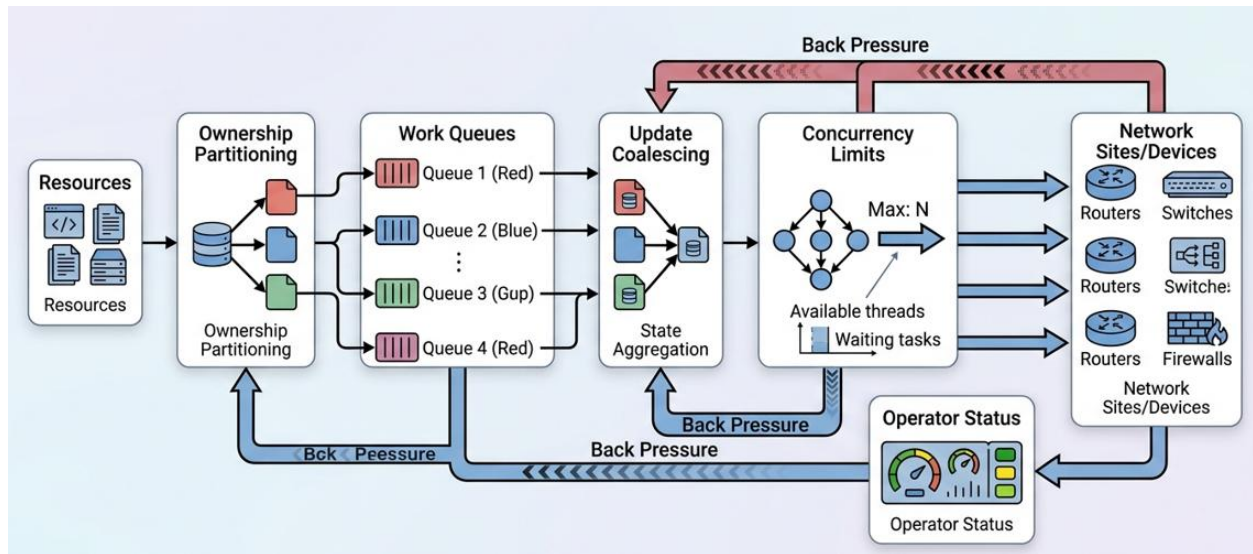


Figure 4: Distribution, Consistency and Back-Pressure Model

IV. RELIABILITY AND OBSERVABILITY

A network control plane built on a cloud will not only be reliable in the sense of service availability, but it will also need to be able to determine to the network whether results requested on the network are converging. That is observability must be able to bridge operator intent, internal processing, actions of devices and resultant network conditions into a state of operation that is understandable.

4.1 End-to-End Operation Traceability

All network operations need to have a stable correlation identity at all times during lifecycle. This identity should be established at the location at which the request arrives at the API layer and subsequently it should be replicated by admission, planning, reconciling, workflow execution, adapter communication, device activities and lastly by state validation. This forms a full circle relationship between an operator request and their measurable external impacts.

Such lifecycle may be considered in various perspectives that may be distributed tracing, metrics and structured logs. OpenTelemetry outlines the uniformity methods to gather traces, metrics and logs between distributed services [4]. Combining these signals can allow operators to determine the locations of delays, differentiate between control-plane failures and issues on the device side, and rebuild the history of long-lasting operations without relying on the local logs of just one service.

4.2 Control-Outcome Indicators

The traditional platform availability measurements are not adequate in measuring network automation behavior. A controller may be maintained accessible as the results of its reconciliation queues grow logarithmically, or where its workflows fail to make any speedy progress. The control plane is thus to monitor the indicators related directly to the control outcomes.

Delay in queue, change success has been verified, length of workflow, workflow stuck, freshness of observations and rollback completion are some of the important service-level measures. An appearance of responsiveness and rightness is accorded by such signs.

V. EVALUATION FRAMEWORK

The suggested control plane must be put through an organized evaluation strategy which assesses scalability, responsiveness, state convergence, consistency, and resiliency under normal network automation conditions. In place of relying on infrastructure level availability, an evaluation of whether or not the control plane can safely transition desired network states that are accepted to verifiable network states is performed. The model takes into account the



growing resource groups, the variety of devices, simultaneous transformations, the devices volume of events, and the managed failures.

5.1 Workload and Scale Dimensions

The test should step-by-step increase the number of resources it manages, devices on the network, reconfigurations within a specific time frame, and events of new telemetry. It is possible to express various working conditions with multiple workload profiles. A tiny autonomous adjustment is customary system management contacts with massive deployments of sites are coordinated interchanges of many devices and reliance. The updates of fan-out policies where a single logical request is hitting a significant number of targets simultaneously are called high fan-out policy updates.

Possessing such types of workloads allows testing the control plane in different configurations with different complexities of requests and processing volumes. Resource allocation, sharing ownership, queuing and concurrency may be measured to assess whether proportional or linear growth in processing delay can be attained as the demand increases or whether they experience unwanted contention. The potential of device level concurrency constraints should also be considered because network infrastructure may have constraints different than any normal cloud services.

5.2 Performance and Responsiveness Indicators

A variety of quantitative indicators provide a base to measure the control-plane behavior with. Latency of API is also utilized to measure the responsiveness of the admission layer but queue delay is utilized to measure the time a request has to wait before it can be processed. Reconciliation throughput is the rate of measurement of the resources through which the resources are to be transformed to the desired level.

Network-specific behavior should be provided to more indicators. The number of redundant external operations might be indicative of weaknesses in idempotency and recovery logic. The difference between the time to achieve the desired state acceptance and the network result achieved should be measured by convergence time. Status freshness is used to reflect the up-to-date information provided to the operators and the controllers. Resource usage (processor usage, memory usage, number of queues and storage activity) can also shed more light on scalability.

These indicators need to be viewed as a package and not individually. To illustrate, higher throughput with significantly older observations might not be a significant change towards better due to increasingly stale information being used to make decisions.

5.3 Controlled Failure Scenarios

In the resilience assessment disruption controlled has to be established in distinct locations of the control loop. Reconcilers are terminated during admission processing, planning, execution of workflow, interaction with devices and verification. They can purposely delay in the delivery of events, in order to test how the consumers will behave when they get the observables later than they are expected. Device interfaces can provide ambiguous results, a scenario where an operation can be accepted, it cannot be determined whether it will be completed or not.

State-store partition scenarios can be studied to understand the behavior of controllers which deal with a scenario in which authoritative state temporarily becomes unavailable. Repeated event delivery should be re-played as well to find out whether repeated observations are safely processed by the consumers. All these situations are representative of typical distributed system states, where partial progress, slow information and adversarial external results coexists.

5.4 Correctness and Safety Criteria

It should validate the correctness with respect to express architectural invariants. Once recovery and reconciliation are made, the resources should be in its most accepted form rather than in an obsolete request. Unauthorized operations should not be directed by the control plane, and there should be no conflicting changes of the resources by independent controllers.

Uncertain external operations should be particularly focused on. The usage of recovery software should also verify the success of external operating before the next command is sent in the event that a controller aborts on transmitting a device demand but any receiver does not receive the command. This practice does away with duplication or conflicting network modifications.



5.5 Recovery and Operational Visibility

Recovery evaluation requires considering not just system restoration but also expertise amongst the operators. Time spent to restart the disrupted workflow, time spent to stabilize again to reach convergence, and the time spent until the resources will remain in the degraded or uncertain form is all relevant measures.

Status such as delayed processing, stale observations, unresolved dependencies and recovery in progress need to be visible operator-friendly status. A device that recovers automatically, without notifying operators that there are unresolved network conditions, partially guarantees its operation.

5.6 Assessment Matrix

The framework can put together situations on four dimensions such as the magnitude of workloads, the health of the control-plane, the health of the network and the expected outcome. The initial sought generation will be characterized by a set of scenarios, the device capabilities will be pertinent, the injected condition, the expected controller behavior and the ultimate convergence requirement. It is a technique which enables a repeat comparison of the architecture behavior at all normal, high load, degraded and recovery conditions.

Overall, the assessment system is based on the measurable control outcomes, and not on the infrastructural availability itself. It provides a unified basis of knowledge, through these concerted efforts on scalability, responsiveness, correctness, state freshness and recovery behavior as well as operator visibility, that the proposed cloud-native control plane is capable of ensuring safe and predictable network convergence into the face of a challenging distributed constraint.

VI. DISCUSSION

The provided architecture reiterates how it does not necessitate taking declarative control as the requisite that all the network operations can achieve their desired state immediately. The network configuration used unlike most other cloud application resources can be ordered, have maintenance windows, device constraints and manage via management protocols that have limited transactional capabilities. Hence, in the process of reconciliation, intermediate states should not be ignored but be continued to achieve the expected result initially.

6.1 Declarative Intent and Durable Execution

A declarative resource expresses what the control plane has yet to accomplish, when a single execution processes are blocked or broken. Durable workflows are also therefore complementary to declarative control. They preserve historical record of the performance of multi-stage operations, which include finished operations, pending operations, external identifiers and conditions, which are pre-requisite to an operation.

The existence of workflow coordination under declarative resources provides a good architectural partitioning. Any interruption or failure of a procedure, can nullify no state which we desire. Alternatively, the reconciliation process may re-examine the current state of the network and determine what to do regarding the current workflow: leave it unchanged, restore it, or replace it with some other safe process. The method also eliminates the risks of carrying out some operations on the equipment in an erratic manner, when they are reclaimed after collapsing of the controller.

6.2 Lessons from Cloud-Native Orchestration

The practice utility of declarative specifications, admission control, scheduling, continual reconciliation and automated recovery are testified to by massive cluster management systems [5]. These standards develop a beneficial framework to network automation as they would provide a loosely coupled controller and consistent management of the distributed resources.

But network environments give rise to new requirements, which it is not possible to deal with using generic orchestration mechanisms by themselves. The network controllers should consider the new devices as they make decisions regarding capability of the devices, connections in the topology, dependencies of configuration, and limits of maintenance and protocol sensitive transactions. Most of the resources that are valid at the API level will continue to require a definite implementation process depending on the target devices and their capabilities.

6.3 Network-Specific Extensions

The control plane can use Capability modeling to know the procedures supported by specific targets. Topology awareness thrives the controllers to consider the interdependencies among the interconnected resources so as to end up



in initiation of changes. Change verification gives an indication that the desired result has truly been delivered as opposed to presupposing success on command acceptance. Where the effect of failure in the multi-stage procedures can be directly reversed, measures can be counterbalanced.

The extensions suggest that the network automation developed in the cloud can be described by the intersection of the general principles of reconciliation with the network-specific safe-guards. The resulting architecture is sustained at the expense of the scalability and recovery benefits of declarative control and recognizes the fact that transitions between states in the network can often need to be implemented through coordinated, observable and ordered execution.

VII. CONCLUSION

Cloud-native network automation entails beyond containerized services and orchestration infrastructure. A sound network control plane should be organized on the desired state, sustained execution, and verified results, as can be observed through the analyses done in this study. This model has knowledge of the nature of network environments, where they may be heterogeneous, their operations are asynchronous, they may be interconnected only intermittently, and changes in configuration may be a shared resource.

The proposed architecture distinguishes declarative resource management, permanent desired-state storage, reconciliation, workflow coordination, and adapters that are capable of performing their duties considering abilities and event-driven observability. Such division creates a sense of responsibility and boundary of consistency, and enables separate components to grow and recover on their own. Idempotent reconciliation minimizes the effects of duplicate requests, and ownership controls and use of optimistic concurrency minimize conflicting modifications. Durable workflows omit the execution history and provide reliable information to recover on interrupted or unpredictable processes.

Capability-mindful adapters intertwine the cloud-native model further with the awareness of differences between the network management protocols and the capabilities of the devices. Meanwhile, outcome-oriented observability connects operator requests to network-verified network conditions and internal processing. Measures such as reconciliation age, convergence time, and freshness in observation, queue delay, and approved modifications provide a better perspective of the health of the control plane than just service availability.

The perceived test provides insights into how scalability, consistency, resilience, and operational visibility can be exercised using representative workloads and controlled failure states. Overall, the architecture positions reconciliation as the primary working model and takes into account lasting workflow history, network-sensitive safety mechanisms, capability consciousness and state validation significant control-plane characteristics.

REFERENCES

1. Kubernetes Documentation, Controllers. <https://kubernetes.io/docs/concepts/architecture/controller/>
2. R. Enns et al., Network Configuration Protocol, RFC 6241, June 2011. <https://doi.org/10.17487/RFC6241>
3. A. Bierman et al., RESTCONF Protocol, RFC 8040, January 2017. <https://doi.org/10.17487/RFC8040>
4. OpenTelemetry, OpenTelemetry Specification. <https://opentelemetry.io/docs/specs/otel/>
5. A. Verma et al., Large Scale Cluster Management at Google with Borg, EuroSys 2015. <https://doi.org/10.1145/2741948.2741964>
6. Q. Wu et al., A Framework for Automating Service and Network Management with YANG, RFC 8969, January 2021. <https://doi.org/10.17487/RFC8969>