

Evaluating Redundancy and Failure Detection

A Study of FHRP and BFD-Based Network Availability

Mai Elbaabaa¹, Ahmed Ben Hassan², Mahmud Mansour³, and Najia Ben Saud⁴

^{1,2,3,4} Faculty of Information Technology, University of Tripoli-Libya
n.ben_saud@uot.edu.ly

Abstract. The exponential growth of the Internet and its integration into daily life underscore the critical importance of resilient networks. Service outages can cause significant financial losses and damage reputation. First-hop redundancy protocols (FHRPs) are commonly used to enhance virtual gateway resilience and reduce downtime, but they can suffer from slow failure detection, leading to packet loss. Bidirectional routing detection (BFD) provides a rapid mechanism for link failure detection and connectivity monitoring. This paper explores the intricate landscape of network reliability, investigates the benefits of combining BFD with three prominent FHRPs (HSRP, VRRP, and GLBP) to improve network performance, increase availability, and reduce downtime. The evaluation is based on metrics of convergence time, packet loss, CPU utilization, and bandwidth consumption. Results from PNETLAB simulations indicate that using BFD greatly speeds up the detection of failures and reduces packet loss for all three protocols. GLBP achieved the fastest convergence, while VRRP exhibited the lowest CPU utilization. The findings indicate that the integration of Bidirectional Forwarding Detection with First Hop Redundancy Protocol gateways significantly enhance network convergence times, thereby improving overall network reliability and stability.

Keywords: Network Reliability, High Availability, FHRP, HSRP, VRRP, GLBP, BFD.

1 Introduction.

The exponential growth of the Internet and the ubiquitous integration of computing systems into various facets of our lives have propelled us into an era where dependence on reliable networks is more critical than ever. In the early 1990s, the concept of electronic mail connectivity marked a novel leap in communication. Fast forward a decade, we witnessed a seismic shift where reputable companies transformed their websites into platforms for direct consumer transactions. This transformative journey reflects the staggering progress of the Internet and networking technologies, becoming integral to our daily lives and professional landscapes [1].

Our contemporary reliance on the Internet is profound, impacting not just personal communication but also on the functioning of businesses, organizations, and critical infrastructure. With this increased dependence comes an essential requirement – the

need for these networks to be not only operational but highly reliable. We now lean on networks for communication, collaboration, data exchange, and even the fundamental functioning of critical systems. In this context, network availability emerges as a linchpin for sustaining the seamless operation of our interconnected world [1].

As networks become an integral part of our lives, the demand for their reliability has surged. The consequence of a network failure extends beyond inconvenience; it reverberates financial losses, disruptions to productivity, and damage to the reputation of organizations. A key strategy in fortifying network reliability is redundancy. Redundancy, as part of a broader spectrum of services including resiliency, load balancing, and security, plays a pivotal role in achieving higher availability. The efficacy of redundancy, however, is contingent on meticulous network design to support its implementation [1].

1.1 Methods of Achieving High Availability

Several strategies can be employed to enhance the availability of a network, ensuring that it remains operational and reliable under various conditions [2]:

1. **Hardware Availability:** High availability can be bolstered through reliable hardware components that offer failover mechanisms. This includes using redundant power supplies, hot-swappable components, and uninterruptible power supplies (UPS) to minimize downtime caused by hardware failures. Manufacturers often provide Mean Time Between Failure (MTBF) data to help assess the expected reliability of hardware components.
2. **Software Availability:** Reliable software configurations and robust operating systems are crucial for achieving high network availability. Software updates, patches, and efficient error-handling mechanisms play an important role in preventing failures. Additionally, virtualization and containerization technologies can offer seamless transitions in case of software failures, as they allow applications to move between different environments without interruption.
3. **Network Based on Fault-Tolerant Devices:** Fault-tolerant devices are engineered to handle failures gracefully, allowing the network to continue functioning even when a component malfunctions. Such devices often have built-in redundancy, such as multiple network interface cards (NICs) or backup processors, which can take over seamlessly when the primary unit fails. This approach significantly minimizes the risk of network outages due to equipment failures.
4. **Network with Redundant Topologies:** Implementing redundant network topologies, such as ring, star, or mesh configurations, is a widely used strategy to enhance availability. These topologies ensure that if one connection fails, data can be rerouted through alternate paths, maintaining uninterrupted service. Techniques like load balancing and automatic failover ensure that the network remains resilient against disruptions, providing continuous connectivity even during unexpected outages.
5. **Traffic Engineering and Quality of Service (QoS) Mechanisms:** QoS mechanisms represent a crucial approach to achieving high availability in network infrastructures. These methodologies are designed to optimize network traffic flow,

thereby guaranteeing predictable performance and mitigating the impact of packet loss, jitter, and delay, even amidst network congestion or partial degradation. By strategically prioritizing critical data and efficiently utilizing available bandwidth, QoS mechanisms are instrumental in maintaining requisite service levels and enhancing overall network resilience. Furthermore, properly implemented traffic engineering solutions facilitate the rerouting or reshaping of traffic to circumvent bottlenecks, thus sustaining availability during fault conditions or overload events. For example, Multi-Protocol Label Switching (MPLS) networks frequently employ advanced QoS queuing strategies for the classification and prioritization of diverse traffic flows. As detailed in [3], these mechanisms are pivotal in preserving service quality through the judicious allocation of resources to high-priority applications.

1.2 The Cost of Downtime

As we delve deeper into the realm of network reliability, it becomes increasingly crucial to shed light on a critical yet often underestimated aspect – the financial repercussions of network downtime. According to Gartner Research, the average cost of IT downtime stands at a staggering \$5,600 per minute, with potential hourly costs ranging from \$140,000 to \$540,000. Alarming, 98% of organizations assert that a single hour of downtime exceeds \$100,000 in costs [4].

However, these financial impacts are far from mere statistical abstractions; they translate into tangible losses that encompass various facets of business operations, including productivity, revenue, reputation, and overall financial performance. The repercussions of network downtime extend beyond the immediate financial toll, influencing the overall health and competitiveness of an organization.

Additionally, the excerpt provides valuable insights into the impact of server downtime on organizations, underscoring the increasing demand for reliability in server hardware, server operating systems (OS), and mission-critical applications. It emphasizes the critical importance of maintaining high levels of uptime, with over 90% of corporations now requiring a minimum of "four nines" (99.99%) reliability, and nearly 40% striving for "five nines" (99.999%) uptime or higher. The distinction between these uptime levels is significant, with each additional nine representing a substantial reduction in annual per server downtime.

Table 1. Reliability/Uptime [4].

Reliability (%)	Downtime per year	Downtime per month	Downtime per week
90% (one nine)	36.5 days	72 hours	16.8 hours
99% (two nines)	3.65 days	7.20 hours	1.68 hours
99.9% (three nines)	8.76 hours	43.8 minutes	10.1 minutes
99.99% (four nines)	52.56 minutes	4.32 minutes	1.01 minutes
99.999% (five nines)	5.26 minutes	25.9 seconds	6.05 seconds
99.9999% (six nines)	31.5 seconds	2.59 seconds	0.605 seconds

To compute the theoretical availability of a network, it is segmented into individual components such as hardware, software, physical connections, and power supplies. Typically, manufacturers provide availability expectations for most equipment, often in the form of Mean Time between Failure (MTBF). However, components lacking this data, such as power sources, statistical information and estimations are employed. The anticipated time required to repair each component, known as Mean Time, to Repair (MTTR), must be estimated. The availability of each component is then determined using the formula [5]:

$$\text{Availability} = \text{MTBF}/(\text{MTBF}+\text{MTTR}) \quad (1)$$

Moreover, the excerpt highlights the significant financial implications of server downtime, with a single hour of downtime potentially resulting in losses exceeding \$300,000 for a majority of mid-sized and large enterprises. Notably, almost half of these organizations report hourly outage costs surpassing one million (\$1M) to over five million (\$5M), underscoring the considerable financial risks associated with server downtime.

These findings from ITIC's 2022 Hourly Cost of Downtime survey underscore the imperative for organizations to invest in reliable server infrastructure to ensure continuous, uninterrupted data access, maintain regulatory compliance, and mitigate risk. In today's interconnected digital landscape, where downtime can disrupt operations across data centers, clouds, remote work environments, and the network edge, the need for robust, resilient server hardware, and applications has never been greater. This financial impact is not merely a statistical abstraction; it translates into tangible losses encompassing productivity, revenue, reputation damage, and impaired financial performance. The repercussions of network downtime extend beyond the immediate financial toll, influencing the overall health and competitiveness of an organization. To provide a visual representation of the escalating financial implications of network disruptions, is presented below [4].

Figure 1 and Table 2 illustrate the escalating financial impact of network downtime, emphasizing the urgent need for effective redundancy and high availability measures.

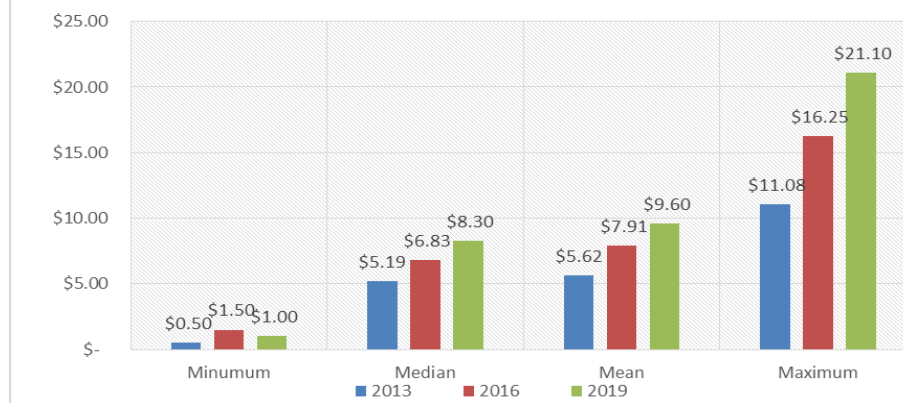


Fig. 1. Estimated Total Cost per Minute of Unplanned Downtime.

Table 2. Monetary Cost of Hourly Server Downtime: Per Minute/Per Server (s) [4].

Hourly Cost of Downtime	Per Minute/Per Server	Per Minute/Per 10 Servers	Per Minute/Per 100 Servers	Per Minute/Per 1,000 Servers
\$10,000	\$167	\$1,670	\$16,700	\$167,000
\$100,000	\$1,667	\$16,670	\$166,667	\$1,666,670
\$300,000	\$4,998	\$49,980	\$499,800	\$4,999,800
\$400,000	\$6,666	\$66,660	\$666,600	\$6,666,670
\$500,000	\$8,333	\$83,330	\$833,300	\$8,333,300
\$1,000,000	\$16,667	\$166,670	\$1,666,700	\$16,667,000
\$2,00,000	\$33,333	\$333,330	\$3,333,300	\$33,333,000
\$3,000,000	\$49,998	\$499,980	\$4,999,800	\$49,998,000
\$5,000,000	\$83,333	\$833,330	\$8,333,300	\$83,333,000
\$10,000,000	\$166,667	\$1,666,670	\$16,666,700	\$166,667,000

In summary, as we navigate the intricate landscape of network reliability, redundancy, and the efficacy of FHRPs, this paper aims to shed light on the substantial investments made by large companies to enhance availability. It underscores the financial implications of network downtime and emphasizes the critical role of redundancy in mitigating these risks. As the paper unfolds, readers will gain insights into the operational dynamics of FHRPs and their pivotal role in fortifying network resilience.

2 Related Work.

Mehdi Berrish et al. [6] examined the performance of Hot Standby Routing Protocol (HSRP) with and without BFD in terms of packet loss, convergence time, CPU utilization, and bandwidth consumption in their study, "Performance Analysis of Bidirectional Forwarding Detection (BFD) over the HSRP.

In the study "Performance Evaluation of Bidirectional Forwarding Detection (BFD) over Virtual Router Redundancy Protocol (VRRP)" by Ben Hassan [7], the authors examined how well VRRP performed in terms of convergence time, CPU usage, bandwidth consumption, and packet loss while working with and without BFD.

In the research study "Design and Implementation of VRRP and BFD Linkage Technology in Campus Information Service Platform Network" by Niu and Li [8], they looked at how to incorporate BFD and VRRP technologies into networks that already existed. It hasn't been assessed, nevertheless, how BFD impacts CPU utilization, bandwidth consumption, packet loss, and convergence time.

They compared the performance of FHRPv4 and FHRPv6 in terms of packet loss, convergence time, and CPU utilization without assessing bandwidth consumption in by Ben Saud [9], "Performance Evaluation of First Hop Redundancy Protocols in IPv4 and

IPv6 Networks." They also used IP SLA as a technique for identifying ISP failures, which has a failure detection time of at least one second.

The study "Performance analysis and functionality comparison of first hop redundancy protocol IPV6" by M. Mansour [10] used IP SLA as a technique to detect ISP failures and concentrated on the FHRPv6 performance in terms of packet loss and convergence time.

Building upon this groundwork, Julia extended the comparative assessment, emphasizing throughput, jitter, packet loss, and downtime. Julia's findings positioned GLBP as a standout performer, showcasing superior results when compared to both VRRP and HSRP [11].

Anwar [12] contributed to this body of research by conducting a comparative study on HSRP, VRRP, and GLBP. However, the specific performance metrics employed in Anwar's analysis were not explicitly detailed. Complementing these endeavors, Zemtsov [13] shifted the focus to the recovery time of FHRPs within an industrial enterprise network context. These collective studies not only provide a nuanced understanding of the operational intricacies of HSRP, VRRP, and GLBP but also underscore the importance of evaluating their performance under diverse scenarios.

The use of fast detection BFD with VRRP to enhance failure detection and a failover was examined in the paper "FDVRRP: Router implementation for fast detection and high availability in network failure cases" by Kim et al. [14]. However, it was limited to an on-premises failure scenario involving the failure of a master router.

3 Redundancy Protocols (FHRP).

3.1 HSRP

Hot Standby Routing Protocol (HSRP) was invented by Cisco to provide dynamic failover between routers within the HSRP group in case of failure. Cisco enhanced the second version of HSRP (HSRPv2) to support IPv6. The active-standby model supports end-user traffic with one device at a time and one on standby to take over if the active device fails. It enables a set of router interfaces to work together to present the appearance of a single virtual router or default gateway to the hosts on a LAN [9].

HSRP Operation. When you use HSRP, a set of routers works in concert to present the illusion of a single virtual router to the hosts on the LAN. This set is known as an HSRP group or a standby group. A single router elected from the group is responsible for the distribution of the packets that hosts send to the virtual router. This router is known as the active router. Another router is elected as the standby router. In the event that the active router fails, the standby assumes the packet-forwarding duties of the active router. Although an arbitrary number of routers can run HSRP, only the active router forwards the packets sent to the virtual router [15].

Once the protocol has finished the election process, the active and standby routers are the only ones sending periodic Hello messages to maintain the group and manage the transition between states. The active router is selected based on the device priority. By default, the priority value is 100, which can be configured. If multiple routers share

the same priority, the router with the highest IP address will be elected as the active router [15].

A LAN can include multiple hot standby groups simultaneously. Each HSRP group emulates a single virtual router. The individual routers can participate in multiple groups. In this case, the router maintains a separate state and timer for each group. Each standby group has a single, well-known MAC address and IP address [15].

As part of the HSRP operation, when a router becomes the active router, it sends ARP responses with virtual the IP and virtual MAC addresses. These gratuitous ARP responses are crucial because they help switches and learning bridges update their port-to-MAC mappings, ensuring that the correct MAC address is associated with the virtual IP address. Unlike the typical ARP responses sent when an interface first becomes active, HSRP-specific ARP responses carry the virtual MAC address in the packet header. These responses ensure that the network devices correctly map the virtual IP address to the virtual MAC address.

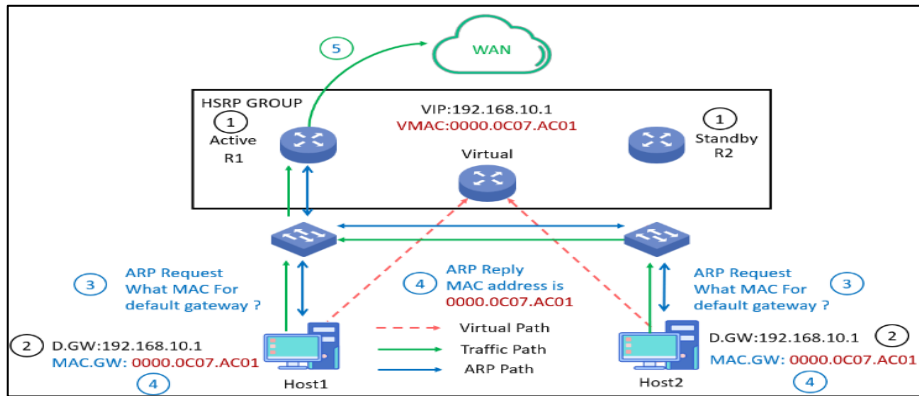


Fig. 2. HSRP Operation.

Figure 2 illustrates the HSRP operation with the following steps:

1. Routers in the HSRP group are configured with the same virtual IP address, and then a virtual MAC address is automatically generated. One router is chosen as the active router to forward traffic, while the other is the standby router. Based on the highest priority, if equal, it is based on the highest IP address.
2. The host configures the virtual IP address as its default gateway.
3. Since the host does not know the virtual MAC address, it sends an ARP request (broadcast).
4. The active router responds to the ARP request with the virtual MAC address so that the host can reach the default gateway.
5. Once the ARP is complete, the traffic is sent to the virtual MAC address, and the active router forwards it to the appropriate destination. If the active router fails, the standby router automatically takes over its role.

HSRP States. HSRP has 6 states: Initial, learn, listen, speak, standby and active, as shown in Figure 3.

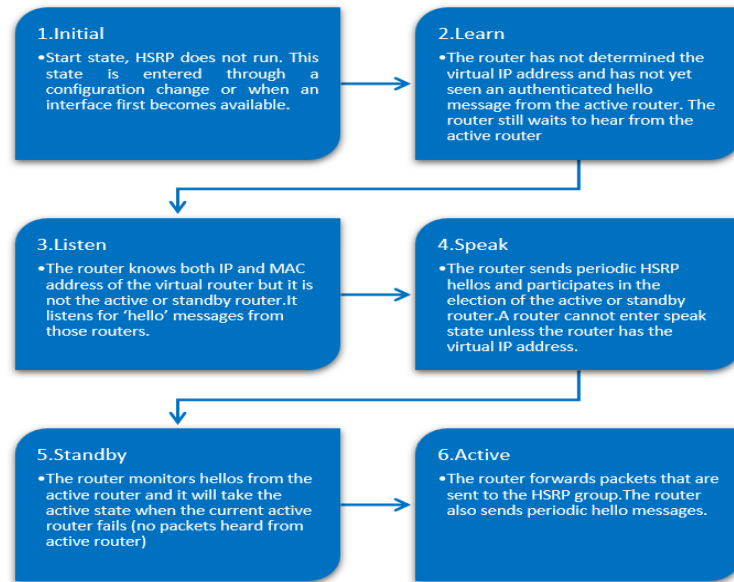


Fig. 3. HSRP Status.

1. **Initial:** HSRP is not operational in the initial state. This state is entered by changing the configuration or when an interface first becomes available.
2. **Learn:** The router has yet to determine the virtual IP address and has not encountered an authenticated hello message from the active router. The router is still waiting for a response from the active router.
3. **Listen:** The router is aware of the virtual router's IP and MAC addresses, but it does not serve as the active or standby router. It listens to hello messages from these routers.
4. **Speak:** The router takes part in sending regular HSRP hello messages and engages in the process of electing the active or standby router. Unless a router possesses the virtual IP address, it cannot enter the speak state.
5. **Standby:** The standby router monitors hello messages from the active router. If the active router fails and hello messages are no longer received, the standby router will assume the active role to maintain service continuity.
6. **Active:** The active router is responsible for forwarding traffic addressed to the HSRP group. It also periodically transmits hello messages to inform other HSRP routers of its active status.

HSRP Timers. Hello time is the estimated time that routers send in a hello message to signal that the peer router is active, with a default value of 3 seconds.

Hold time is the estimated period during which the standby router will announce that the peer is down and become active, with a default value of 10 seconds. These timers

are tuneable and tweaked to achieve the lowest convergence, making a network highly accessible.

3.2 VRRP

Virtual Router Redundancy Protocol (VRRP) is an open standard redundancy protocol for constructing a fault-tolerant default gateway.

VRRP is a redundancy protocol to LAN routers. It provides an alternate route path for hosts without altering the IP address or MAC that the host knows. VRRP follows the same principle as Cisco's HSRP, with certain changes.

VRRP Operation. VRRP provides a set of routers that work in concert to present the illusion of a single virtual router to the hosts on the LAN. This set is known as a VRRP group. The master router is the only router that is responsible for packet forwarding, it is chosen during electing process; another router is chosen as the backup router. If the master router fails, the backup will assume the master router's packet forwarding duties. This process is transparent to users. Although multiple routers can run VRRP, only the master router sends the packets to the virtual router. According to device priorities, routers in a VRRP group elect the master [16].

VRRP advertisement packets are sent to all backups in the VRRP group to communicate the operating status and configuration periodically to reduce network traffic [16]. For load splitting, HSRP and VRRP support multiple groups, with separate states and timers for each group.

ARP is also used in VRRP; the correct virtual MAC address is associated with the virtual IP address at network devices and helps switches update their port-to-MAC mappings.

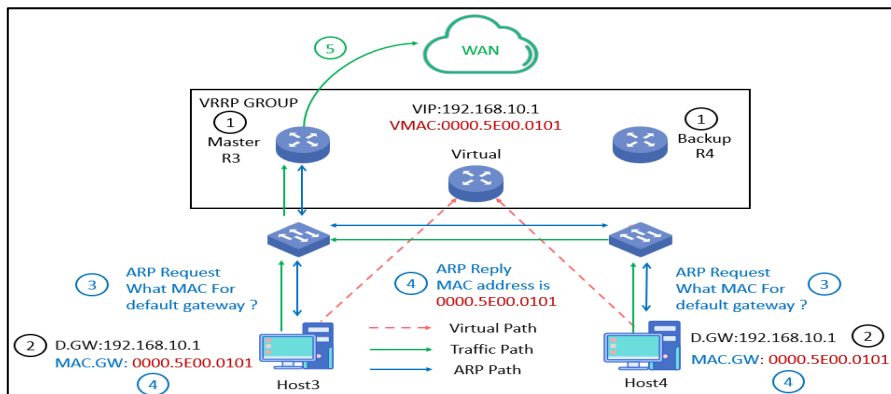


Fig. 4. VRRP Operation.

Figure 4 illustrates the VRRP operation with the following steps:

1. Routers in the VRRP group are configured with identical virtual IP address, and then a virtual MAC address is automatically created. One router is chosen as the master

- router to forward traffic, while the other is the backup router. Based on the highest priority, in case equal, then the highest IP address is preferred.
2. The host configures the virtual IP address as its default gateway.
 3. Since the host does not know the virtual MAC address, it sends an ARP request (broadcast).
 4. The master router responds to the ARP request with the virtual MAC address so the host can reach the default gateway.
 5. Once the ARP is complete, the traffic is sent to the virtual MAC address, and the master router forwards it to the appropriate destination. If the master router fails, the backup router automatically takes over its role.

VRRP States. HSRP has 3 states: initialize, backup and active as shown in Figure 5.

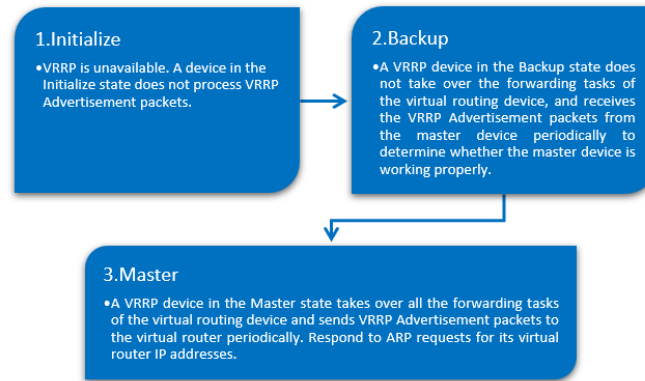


Fig. 5. VRRP Status.

VRRP Timers. The time period that routers send out a hello message to indicate that a peer router is operational is known as the "hello time.", with a default value of 1 second.

Hold time is the approximate time that the backup router will declare that the peer is dead and become the master, with a default value of 3 seconds. These hello times are adjustable and set to achieve minimum convergence, thus making the network highly available. while hold time is not but is calculated by hello time: **Hold timer = (3 x Hello timer) + Skew time**. VRRP has the best default timer, its default timings enable it to converge more quickly than HSRP or GLBP.

3.3 Gateway Load Balancing Protocol (GLBP)

Gateway Load Balancing Protocol (GLBP) is one of the First Hop Redundancy Protocols (FHRP), which provides redundancy, like other First Hop Redundancy Protocols, as well as load balancing. It is a Cisco proprietary protocol that can perform both functions. It provides load balancing over multiple routers using a single virtual IP address and multiple virtual Mac addresses [17].

GLBP uses a single virtual IP address and several virtual MAC addresses to distribute load balancing over multiple routers (gateways). Every router in the Virtual Router Group participates in packet forwarding, and each host is configured with the same virtual IP address.

GLBP Operation. GLBP works by making use of a single virtual IP address, which is configured as the default gateway on the hosts. When the routers are set to a GLBP group, they first choose one gateway to be the Active Virtual Gateway (AVG) for that group. The election process is based on the highest priority of each gateway, in case of the same priority, then the gateway with the highest actual IP address becomes the AVG [18].

In a Gateway Load Balancing Protocol (GLBP) group, member devices serve as redundant backups for the Active Virtual Gateway (AVG). The AVG is responsible for assigning a unique virtual MAC address to each member of the GLBP group. Consequently, each gateway assumes the responsibility of forwarding packets destined for the virtual MAC address specifically assigned to it by the AVG.

These gateways are designated as Active Virtual Forwarders (AVFs) for their respective virtual MAC addresses [18].

The AVG plays a pivotal role in responding to Address Resolution Protocol (ARP) requests for the virtual IP address. To achieve load sharing, the AVG strategically replies to these ARP requests by distributing different virtual MAC addresses among the requesting clients [18].

Each gateway that is issued a virtual MAC address is termed an Active Virtual Forwarder (AVF). A GLBP group only has a maximum of four AVFs. If there are more than 4 gateways in a GLBP group, then the remainder will become standby virtual forwarders (SVFs), which will take the place of an AVF in the case of failure [18].

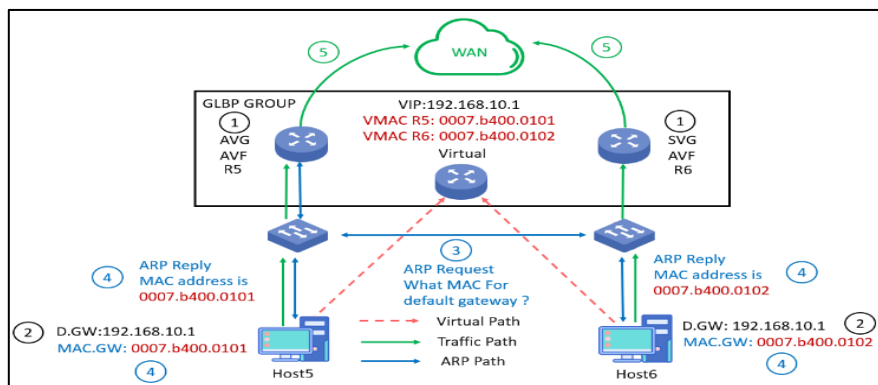


Fig. 6. GLBP Operation.

Figure 6 illustrates the GLBP operation with the following steps:

1. Routers are configured with an identical virtual IP address within a Gateway Load Balancing Protocol (GLBP) group. Subsequently, a unique virtual MAC address is automatically generated for each of the two participating routers. One router is designated as the Active Virtual Gateway (AVG), while the other functions as the

Standby Virtual Gateway (SVG). Notably, the SVG router also operates as an Active Virtual Forwarder (AVF), capable of forwarding network traffic.

2. The host configures the virtual IP address as its default gateway.
3. Since the host does not know the virtual MAC address, it sends an ARP request (broadcast).
4. The AVG receives the ARP requests and replies to Host5 and Host6, but with different virtual MAC addresses. This behavior ensures load balancing by distributing the gateway traffic between the two routers.
5. Once the ARP is complete, Host5 directs traffic through router 5 and forwards it to the appropriate destination, while Host6 uses router 6. In the event of the router's unavailability, GLBP guarantees allowing an alternative router to assume its responsibilities.

GLBP States. AVG has 6 states: disabled, initial, listen, speak, standby and active as shown in Figure 7.

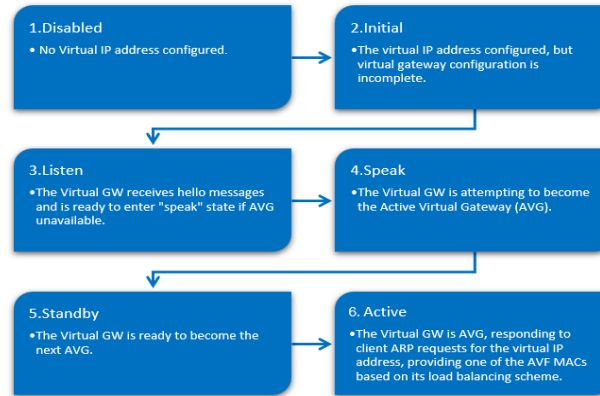


Fig. 7. GLBP States for AVG.

AVF has 4 states: disabled, initial, listen, speak, standby and active as shown in Figure 8.

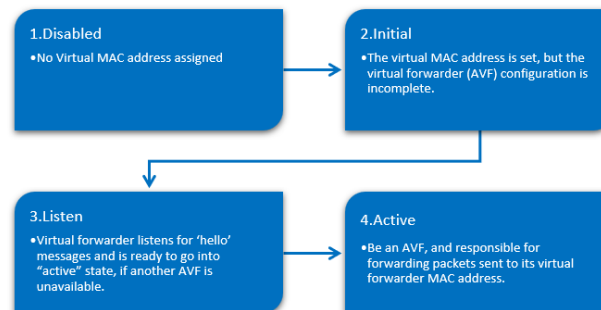


Fig. 8. GLBP States for AVF.

GLBP Timers. The default timings in GLBP are identical to HSRP and tuneable too.

- Hello timer: Hello time default value 3 seconds.
- Hold timer: Hold time default value 10 seconds.
- Redirect time: refers to the duration, set by default to 600 seconds, during which the Active Virtual Gateway (AVG) persists in redirecting client hosts to the previously utilized virtual forwarder MAC address.
- Forwarder timeout: The forwarder timeout is the interval during which the virtual MAC address is valid, with a default value of 14400 seconds.

3.4 IPv6 Equivalents and Considerations

IPv6 routing protocols ensure router-to-router resilience and failover. However, in situations in which the path between a host and the first-hop router fails, or the first-hop router itself fails, first hop redundancy protocols (FHRPv6) ensure host-to-router resilience and failover.

The three main First Hop Redundancy Protocols that support IPv6 are:

- Hot Standby Routing Protocol (HSRPv2)
- Virtual Router Redundancy Protocol (VRRPv3)
- Gateway Load Balancing (GLBP)

The concepts in FHRPV6 are essentially the same as those in FHRPV4. With slight differences described below:

Hosts on the LAN segment use Neighbor Discovery Protocol (NDP) to learn the virtual IPv6 Link-Local Address and associated virtual MAC addresses. The Neighbor Discovery Protocol (NDP) is a protocol of the Internet protocol suite used with Internet Protocol Version 6 (IPv6) and responsible for gathering various information required for network communication, including the configuration of local connections and the domain name servers and gateways [19].

IPv6 hosts learn of available IPv6 routers through IPv6 neighbor discovery RA messages. These are multicast periodically or may be solicited by hosts. HSRP is designed to provide only a virtual first hop for IPv6 hosts [20].

Cisco devices support manually configured millisecond timers for VRRPv3, a feature distinct from VRRPv2, requiring careful implementation due to conditional performance, yet offering compatibility with other VRRPv3 vendors within a 100-40,000ms range [21].

3.5 Comparison between HSRP, VRRP, and GLBP

Table 3 displays a comparison summary between HSRP, VRRP, and GLBP in terms of Terminology, Virtual object, Communication Method, Communication Protocol, Load Balancing, Authentication, Active Selector, Hello and Hold Time, and Preemption.

4 Bidirectional Forwarding Detection (BFD).

Bidirectional Forwarding Detection (BFD) is a network protocol that provides fast failure detection times between two forwarding engines, with minimal traffic overhead. Used across various network types, including MPLS, Ethernet, and IP routed networks, BFD operates independently of media, data protocols, and routing protocols. The primary function of BFD is to detect faults in the path between two endpoints at a speed that traditional protocols cannot achieve, thereby helping in swift network convergence and resilience.

4.1 BFD Detection Modes

BFD has two operation modes: asynchronous mode and demand mode. In asynchronous mode, two end nodes send control packets to each other periodically. If they do not get some of the control packets, they decide that there is a failure. It is the default mode for BFD. While in demand mode, two end nodes send control packets only for a short time to detect whether there are any failures or not. If the connectivity is found to be up, no more BFD control packets are sent until the next command is issued [7]. In demand and asynchronous mode, control packets flow in each direction.

4.2 BFD Echo Mode

The BFD Echo function is an optional feature of the BFD protocol that offloads part of the monitoring process to the data plane. When enabled, it allows devices to send echo packets that are looped back to the sender to check for link failures without necessitating processing by the remote device's control plane. BFD echo mode works with asynchronous BFD; a device sends a BFD echo packet towards its peer, which then routes the packet back to the originating device without involving the control plane of the peer, thereby substantially reducing control plane processing.

Protocol	HSRP CISCO-PROPRIETARY	VRRP Multi-Vendor	GLBP CISCO-PROPRIETARY
Terminology	One Active Router, one Standby Router, other Routers in Standby group	One Master, one or more Backup Virtual Routers	Active Virtual Gateway (AVG), Standby Virtual GW (SVG), (AVFs)
Virtual object	0000.0C07.ACXX (v1, XX is Group ID) 0000.0C9F.FXXX (v2, XXX is Group ID) 0005.73A0.0000 - 0005.73A0.0FFF (IPv6)	0000.5E00.01XX (v1,v2,v3,XX is VRID) 0000.5E00.0200 - 0000.5E00.02FF (IPv6)	0007.b400.XXYY (XX is Group ID, YY is the Gateway number)
Communication Method	IP Multicast 224.0.0.2 (v1) 224.0.0.102 (v2) FF02::66 (IPv6)	IP Multicast 224.0.0.18 (IPv4) FF02::0:0:0:0:0:12 (IPv6)	IP Multicast 224.0.0.102 FF02::66 (IPv6)
Communication Protocol	Pv4, UDP port 1985 IPv6, UDP port 2029	IPv4 and IPv6, protocol 112 (IANA)	IPv4 and IPv6, UDP port 3222
Load Balancing	NO	NO	YES
Authentication	Default: No authentication Plain text authentication MD5 authentication (newly added)	Default: No authentication Plain text authentication MD5 authentication	Default: No authentication Plain text authentication MD5 authentication
Active Selector	Priority – with the highest value electing the Active router and another becoming Standby. The remaining routers listen. The default priority is 100.	Priority – Highest value wins. Default: 100, 254 for router with the same IP as the virtual IP	Priority - One gateway is elected as AVG; another is elected as (SVG). The remaining routers are in a listen state. Highest value wins. Default: 100
Hello and Hold Time	HELLO - Interval between successive HSRP Hello messages from a given router. Default: 3 sec HOLD - Interval between the receipt of a Hello, and the presumption that the sending router failed. Default: 10 sec	Unlike HSRP and GLBP, VRRP does not learn timers from the master router. VRRP requires the hello timer of all routers in the group match. HELLO – Default: 1 sec, HOLD - Default: 3 sec	HELLO - Interval between successive GLBP Hello messages from a given router. Default: 3 sec HOLD - Interval between the receipt of a Hello, and the presumption that the sending router failed. Default: 10 sec
Preemption	Use of preemption allows a HSRP device whose priority has become higher to take over the role as the active router in HSRP. Default: preempt off	With preemption enabled, VRRP switches to a backup if that backup comes online with a priority higher than the new master. Default: preempt on. Exception: The router that owns the IP address (es) associated with the virtual router always preempts.	Preemption allows a backup virtual gateway to become AVG, if it has a higher priority than the current AVG. Default: preempt off AVF (Forwarder) Preemption is similar, except that the weighting is used instead of priority, and it is enabled by default with delay of 30 seconds.

While BFD control packets maintain the BFD session as shown in Figure 9. Leveraging BFD Echo has a minimal impact on control plane resources, which is especially advantageous in high-speed networks where processor time is limited [22].

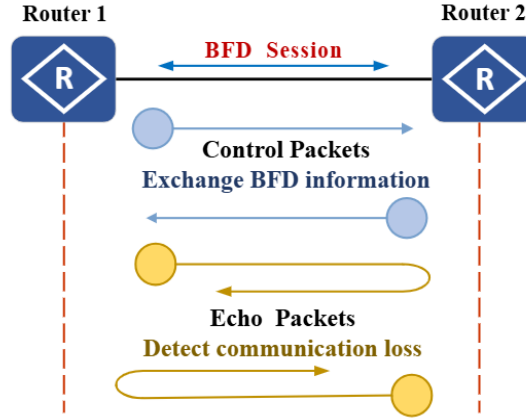


Fig. 9. Asynchronous mode with Echo mode.

5 Experimental Setup.

This paper evaluates the implementation of different first hop redundancy protocols utilizing bidirectional forwarding detection across three locations and assesses their performance relative to FHRP without BFD. Each location is linked to two distinct ISPs to guarantee maximum availability. In the event of a connection failure between an ISP and a gateway, or if the ISP encounters a period of unavailability, the gateway will swiftly detect the disruption. This will enable the backup gateway, which is connected to the other ISP, to take over and assume control. This strategy significantly avoids network downtime, a vital objective for organizations working in contemporary network environments.

In this work, PNELAB network emulator software was utilized to implement network scenarios. Gathered data such as convergence times, packet loss, and bandwidth usage during failover scenarios using the Wireshark program, failover is simulated as a link failure between the active router and ISP1.

The network is structured hierarchically with two default gateway routers, each linked to a distinct ISP. On the LAN side, there are two access switches connecting to end devices. Access switches are connected to gateway routers in a partial mesh network configuration in order to eliminate single points of failure in the company network. The identical design is used for three unique organizations. The topology is illustrated in Figure 10.

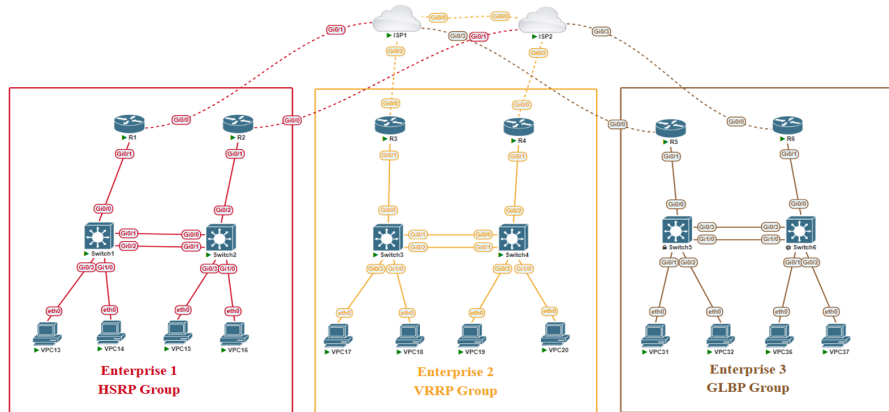


Fig. 10. Network Topology.

5.1 Configuration

HSRP is configured on R1 and R2, VRRP on R3 and R4, and GLBP on R5 and R6, with R1, R3, and R5 initially active due to higher priority. Both routers in each protocol group are configured identically to enable the exchange of hello packets and seamless role transition, simulation parameters are shown in Table 4. Results were measured with FHRP default timers and with timer's optimization. Initially, FHRP is implemented without BFD, using IP SLA to monitor ISP reachability. IP SLA detects connectivity loss (e.g., a link failure between an active FHRP router and ISP1) and triggers a track object. This track object reduces the router's priority, prompting the higher-priority standby router to become active. When implementing FHRP with BFD, the failure detection time is significantly reduced to 50 milliseconds. BFD sessions are configured between the active routers (R1, R3, and R5) and ISP1. Upon detecting a link failure, BFD rapidly notifies the routing process and the associated track object, lowering the active router's priority. Crucially, traffic is immediately redirected to ISP2 via a triggered route, even before the FHRP protocol fully converges and the standby router becomes the new active. This ensures minimal disruption and smoother traffic continuity during the transition.

Table 3. Simulation Parameters.

Parameter	Value (seconds)
HSRP- Hello –Hold time	Default 3 10-With Optimization 1 3
VRRP- Hello –Hold time	Default 1 3
GLBP- Hello –Hold time	Default 3 10-With Optimization 1 3
Forwarder Preemptive Delay	Default 30- With Optimization 0
BFD Detection Time	50 milliseconds
IP SLA Detection Time	1

6 Results.

To evaluate the impact of integrating BFD with FHRP (HSRP, VRRP, and GLBP) protocols, a series of comparative measurements were conducted, the results of which are presented and discussed in this section. The testing methodology included transmitting 3600 ICMP packets continuously for one hour per scenario, allowing us to collect accurate data on data collection on CPU utilization and bandwidth overhead. Following this observation period, an intentional ISP link disruption was induced to monitor and record network behavior upon failure. Multiple repetitions of each experiment were performed to guarantee the consistency and stability of the findings.

6.1 FHRP without BFD Results

1. Convergence Time

- In HSRP, using the default hello and hold timers, we measured a convergence time of 7.3 seconds between routers R1 and R2, and the ISP1-ISP2 convergence took 9 seconds. During this process, we lost 4 ICMP packets. By optimizing the timers, we reduced the R1 and R2 convergence time to 3.78 seconds and the ISP1-ISP2 convergence to 5 seconds. We lost only 2 ICMP packets during this convergence process.
- For VRRP, we recorded a convergence time of 3.48 seconds between routers R3 and R4, while the ISP1-ISP2 convergence required 5 seconds. We observed the loss of 2 ICMP packets during this convergence process.
- In GLBP, with default timers, we found a lengthy convergence time of 36.07 seconds between R5 and R6, and the ISP1-ISP2 convergence took 37.02 seconds. We attribute this to the forwarder preemption delay, which resulted in us losing 18 ICMP packets. However, after we optimized both the timers and the forwarder preemption delay, we achieved substantially faster convergence: 2.86 seconds between R5 and R6 and 3 seconds for ISP1-ISP2 convergence. Consequently, we lost only 1 ICMP packet during this convergence process.

2. CPU Utilization

- In HSRP, without timer optimization, the HSRP process consumed approximately 0.05% CPU on average on routers R1 and R2. Concurrently, total CPU utilization was measured at 2% on R1 and 1% on R2. Timer optimization resulted in a slight increase in the HSRP load to 0.09% on R1 and 0.07% on R2. However, the total measured CPU utilization remained unchanged at 2% on R1 and 1% on R2.
- For VRRP, CPU consumption for the VRRP process was measured at 0.06% on R3 and 0.02% on R4. Total CPU utilization was recorded at 2% on R3 and 1% on R4.
- For GLBP, using default timers, the GLBP process accounted for 0.05% CPU on R5 and 0.04% on R6. Total CPU utilization was measured at 3% on both routers during this period. Following optimization, the GLBP process load increased to

0.12% on R5 and 0.09% on R6. Nevertheless, the observed total CPU utilization remained constant at 3% for both routers.

3. Bandwidth Consumption

- In HSRP, during the testing period with default timers, HSRP consumed approximately 172 KB of bandwidth. This consumption increased to about 499 KB following timer optimization. The HSRP hello packet size is 62 bytes.
- In VRRP, during the testing period, the total measured bandwidth consumption for VRRP was approximately 240 KB. The VRRP hello packet size is 60 bytes.
- In GLBP, during the testing period with default timers, GLBP consumed roughly 270 KB of bandwidth; this value rose to approximately 800 KB following timer optimization. The GLBP hello packet size is 102 bytes.
- In IP SLA, during the testing period, IP SLA consumption was recorded at approximately 562 KB. IP SLA packet size is 78 bytes.

6.2 FHRP with BFD Results

1. Convergence Time

- In the HSRP protocol, using the default hello and hold timers, we measured a convergence time of only one second between ISP1 and ISP2, and the convergence time between R1 and R2 took 6.38 seconds. During the fast convergence process, no ICMP packets were lost. By optimizing the timers, we reduced the convergence time between R1 and R2 to 3.33 seconds, and the convergence time between ISP1 and ISP2 remained the same at one second. During the fast convergence process, no ICMP packets were lost.
- For VRRP, we recorded a convergence time of only a second between ISP1 and ISP2, while the R3 and R4 convergence required 1.57 seconds. No ICMP packets were lost during this fast convergence process.
- In the GLBP protocol, using the default hello and hold timers, we measured a convergence time of only one second between ISP1 and ISP2, and the convergence time between R5 and R6 took 35.2 seconds. During the fast convergence process, no ICMP packets were lost. By optimizing the timers, we reduced the convergence time between R5 and R6 to 0.59 seconds, and the convergence time between ISP1 and ISP2 remained the same at one second. During the fast convergence process, no ICMP packets were lost.

2. CPU Utilization

- In HSRP, during the testing period, BFD consumed an average of 2.57% of the CPU usage on router R1, while the CPU usage was 5% on R1 and 1% on R2.
- In HSRP, the BFD process consumed an average of 2.57% of CPU on R1. Total CPU utilization on R1 reached 5%, while R2 remained at 1%.
- For VRRP, the BFD process required 2.05% of CPU on R3. Total CPU utilization on R3 reached 4%, while R4 remained at 1%.

- In GLBP, the BFD is enabled on two routers, so the BFD process consumes 2.55% of CPU on R5 and 2.62% of CPU on R6. Consequently, total CPU utilization rose to 7% on both routers.

3. Bandwidth Consumption

- In BFD, it exhibited significantly higher bandwidth consumption, totaling approximately 16.54 MB. This elevated consumption is attributed to the frequent packet transmissions necessary for rapid detection, during which BFD echo packets are sent every 50 milliseconds and BFD control packets are sent every second. BFD control packet size is 66 bytes, while BFD echo packet size is 54 bytes.

6.3 Results Summary

This section summarizes the results from sections 6.1 and 6.2. Table 5 displays the results for all three FHRPs, both before and after implementing BFD. The results were measured with default protocol timers and optimized timer settings. In the case of VRRP, no optimization was needed, as its default timer settings are appropriate and much lower than the default timers of other protocols (HSRP and GLBP). Convergence time in seconds was calculated for enterprise routers and also for ISP routers. Packet loss is calculated per ICMP packet loss. For CPU percentage utilization, two measurements were included: the total CPU consumption for both enterprise routers, which together form the FHRP group, is measured.

Table 4. Results.

Protocol	BFD	Timers	Convergence Time		Packet Loss	CPU Utilization
			Enterprise Routers	ISP Routers		
HSRP	Without	Default	7.5	9	4	2%,1%
	BFD	Optimized	3.78	5	2	
	With	Default	6.38	1	0	5%,1%
	BFD	Optimized	3.33	1	0	
VRRP	Without	Default	3.48	5	2	2%,1%
	BFD					
	With	Default	1.57	1	0	4%,1%
GLPB	BFD					
	Without	Default	36.07	37.02	18	3%,3%
	BFD	Optimized	2.86	3	1	
	With	Default	35.2	1	0	2.62%,7%
	BFD	Optimized	0.59	1	0	

Due to the immediate redirection triggered by BFD failure detection, which leverages the shorter ISP convergence time, there is no ICMP packet loss, even though FHRP full convergence may take several seconds.

7 Performance Analysis.

This section compares and evaluates the performance of FHRP before and after implementing BFD, identifying the most effective protocol when BFD is used.

7.1 Convergence Time

Figure 11 shows the convergence time results for the three FHRP protocols before and after integrating BFD and optimizing timers. The results indicate that integrating BFD significantly reduces convergence time compared to IP SLA to one second in all cases. This advantage is due to the BFD failure detection time of 50 milliseconds, compared to FHRP without BFD, which has an IP SLA failure detection time of one second. GLBP, meanwhile, achieved the fastest convergence time between active and standby after optimization, at 0.59 seconds. This benefited from GLBP's load balancing capabilities, which improved redundancy and failure response.

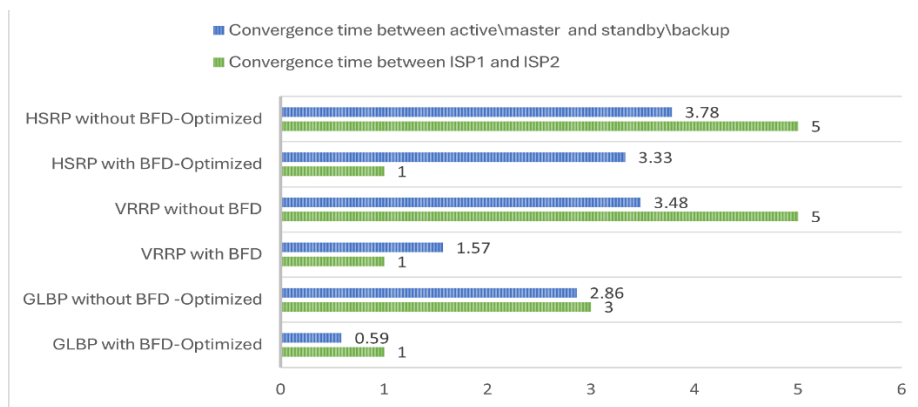


Fig. 11. Convergence Time Comparison.

7.2 Packet Loss Comparison

Figure 12 shows that during convergence, HSRP without BFD lost 4 packets in default due to IP SLA failure detection and hello packets sent every 3 seconds. With optimization, only 2 packets were lost as the hello interval improved to 1 second. GLBP without BFD lost 18 packets in default; this is caused by the forwarder preemption delay, but optimization reduced packet loss to just 1. While FHRP (HSRP, VRRP, and GLBP) with BFD, no packets were lost before or after optimization, thanks to BFD's fast failure detection time of 50 milliseconds.

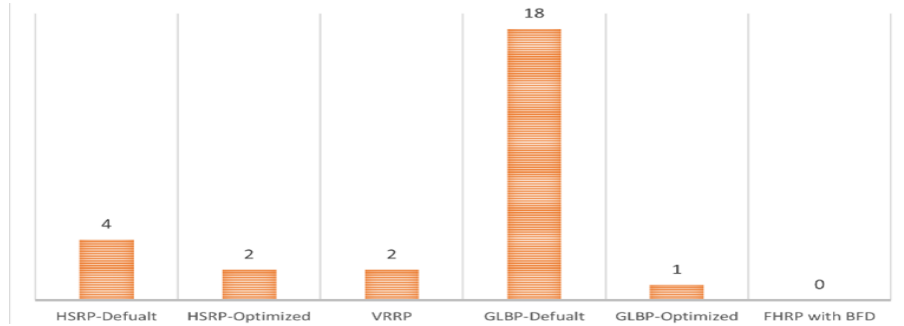


Fig. 12. Convergence Time Comparison.

7.3 CPU Utilization Comparison

Figure 13 shows the increase in CPU usage observed when using FHRP with BFD due to the high load resulting from sending BFD echo packets and BFD control packets, so it can be concluded that FHRP with BFD has the worst CPU usage compared to FHRP without BFD, and among the three protocols, GLBP exhibits the highest CPU consumption when using BFD, making it the least efficient in terms of resource utilization. In contrast, VRRP demonstrates the lowest CPU usage, making it the most efficient among the three protocols.

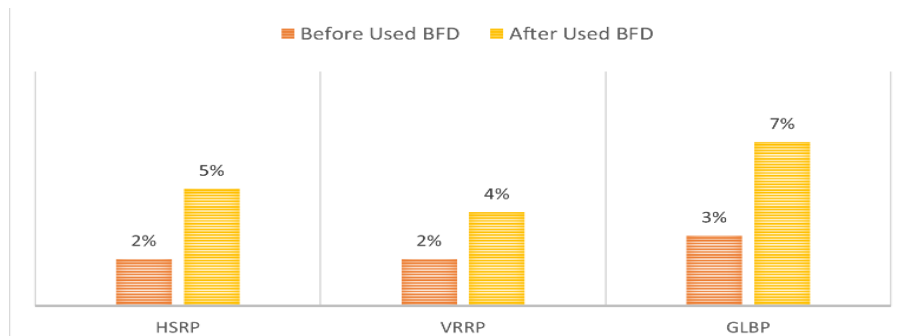


Fig. 13. CPU Utilization Comparison.

7.4 Bandwidth Consumption

Table 6 shows that BFD utilizes very high bandwidth, around 16.54 MB, compared to IP SLA, which was 562 KB, because of sending BFD echo packets every 50 milliseconds and BFD control packets every second, which we mentioned earlier. Additionally, GLBP has the highest bandwidth consumption compared to the other protocols. This is primarily due to the larger GLBP Hello packet size, which is 102 bytes, whereas HSRP and VRRP Hello packets are smaller in size. The increased packet size contributes to higher bandwidth usage, making GLBP the least efficient in terms of bandwidth consumption.

Table 5. Bandwidth Consumption.

Protocols	Bandwidth Consumption
HSRP-Default	172KB
HSRP-Optimized	499KB
VRRP	240KB
GLBP-Default	270KB
GLBP-Optimized	800KB
IP SLA-Detection Time	562KB
BFD- Detection Time	16.54MB

8 Discussion.

The results indicate that to deliver an enhanced availability solution, different technologies could be integrated. In our study, we evaluated FHRPs as a redundancy and high availability solution and BFD for failure detection. To improve the FHRPs, the default timers were optimized to achieve faster response. For GLBP, the change in timers resulted in significant enhancement, whereas the packet loss in ICMP packets was reduced from 18 to 1. The BFD integration reduced the convergence time and packet loss. For HSRP with timer optimization, applying BFD reduced the convergence time between the active and standby routers from 3.78 to 3.33 seconds. For VRRP, the convergence time between the master and backup routers was reduced from 3.48 to 1.5 seconds, and for GLBP, the convergence time was reduced from 2.86 to 0.59 seconds. For all three protocols, the convergence time between ISP routers did not exceed 1 second. On the other hand, the CPU and bandwidth consumption increased, as BFD detection consumed about 16.54 MB, which is relatively high compared to the IP SLA detection mechanism, which consumed only 562 KB.

9 Conclusion.

Following the implementation and testing of FHRP (HSRP, VRRP, and GLBP) with and without BFD, and after evaluating the results based on convergence time, packet loss, CPU utilization, and bandwidth consumption, it is evident that integrating BFD with FHRP significantly enhances network performance by reducing convergence time to 1 second and eliminating packet loss during failures. However, this improvement comes at the cost of increased CPU usage and bandwidth consumption. Among the three protocols, GLBP, which is exclusive to Cisco, achieved the fastest convergence due to its load-balancing capability but exhibited the highest CPU and bandwidth consumption. VRRP demonstrated the lowest CPU utilization. Therefore, organizations must carefully consider the benefits of improved failover speed and reliability against the higher resource demands of BFD, ensuring that sufficient CPU and bandwidth resources are available to sustain its operation.

References

1. Chris Oggerino, "High Availability Network Fundamentals," Cisco Press, 1st edition, 2001.
2. S. Janardhanan and C. M. Machuca, "Availability modeling and evaluation of switches and data centers," in Proc. Int. Conf. Dependable Systems and Applications (DSA), Aug. 2023, doi: 10.1109/DSA59317.2023.00102.
3. Mansour, Mahmud, Ahmed Samood, and Najia Ben Saud. "Assessing Queue Management Strategies to Enhance Quality of Service in MPLS VPN Networks." *Libyan Journal of Informatics* 1.02 (2024): 29-48.
4. ITIC. (2022). Hourly Cost of Downtime Survey. ITIC. Retrieved from <https://itic-corp.com/tag/hourly-cost-of-downtime/>, last accessed 2025/05/21.
5. Mattias Thulin, "Measuring Availability in Telecommunications Networks", "Master's thesis report at Song Networks AB", 2004, pp 13-16.
6. Berrish, M.M., Mansour, M., Hassan, A.B.: Performance Analysis of Bidirectional Forwarding Detection (BFD) over the Hot Standby Router Protocol (HSRP). *International Journal of Computer Science & Security (IJCSS)* 18(4), 48–64 (2024).
7. Hassan, A.B., Mansour, M.: Performance Evaluation of Bidirectional Forwarding Detection (BFD) over the Virtual Router Redundancy Protocol (VRRP). *Procedia Computer Science*. 251, 256–264 (2024).
8. Niu, Y., Li, X.: Design and Implementation of VRRP and BFD Linkage Technology in Campus Information Service Platform Network. *ICMLCA '23: Proceedings of the 2023 4th International Conference on Machine Learning and Computer Application*. 197–201 (2023).
9. Najia Ben Saud and Mahmud Mansour , "Performance Evaluation of First Hop Redundancy Protocols in IPv4 and IPv6 Networks", 2023 IEEE 3rd International Maghreb Meeting of the Conference on Sciences and Techniques of Automatic Control and Computer Engineering (MI-STA) May 21-23 2023, Benghazi, Libya.
10. Mansour, M., Agomati, M., Alsaid, M., Berrish, M., Alasem, R.: Performance analysis and functionality comparison of First Hop Redundancy Protocol IPV6. *Procedia Computer Science*. 210, 19–27 (2022).
11. Julia, I.R., Suseno, H.B., Wardhani, L.K., Khairani, D., Hulliyah, K., Muharram, A.T.: Performance Evaluation of First Hop Redundancy Protocol (FHRP) on VRRP, HSRP, GLBP with Routing Protocol BGP and EIGRP. 2018 6th International Conference on Cyber and IT Service Management (CITSM). 1–5 (2020).
12. U. Anwar, J. Teng, H. Umair, and A. Sikander, "Performance Analysis and Functionality Comparison of FHRP Protocols," in Proc. IEEE Int. Conf. Communication Software and Networks (ICCSN), 2019, doi: 10.1109/ICCSN.2019.8905333.
13. A. Zemtsov, "Performance Evaluation of First Hop Redundancy Protocols for a Computer Networks of an Industrial Enterprise," in Proc. Int. Multi-Conf. Industrial Engineering and Modern Technologies (FarEastCon), 2019, doi: 10.1109/FarEastCon.2019.8934315.
14. Lee, C., Kim, S., Ryu, H.: FDVRRP: Router implementation for fast detection and high availability in network failure cases. *ETRI Journal*. 41, 473–482 (2019).
15. T. Li, B. Cole, P. Morton, D. Li, "Cisco Hot Standby Router Protocol", Request for Comments: 2281, 1998.
16. Huawei, "VRRP Technology White Paper", 2012.
17. Mansour, M., 2020. Performance evaluation of first hop redundancy protocols. *Procedia Computer Science*, 177, pp.330-337.
18. Cisco, glbp - gateway load balancing protocol , <https://community.cisco.com/kxiwq67737/attachments/kxiwq67737/5991-discussions-wan-routing-switching/29508/1/17533-GLBP.pdf>.

19. Cisco, VRRPV3 Protocol Support., https://www.cisco.com/c/en/us/td/docs/ios-xml/ios/ipapp_fhrp/configuration/15-e/fhrp-15-e-book/VRRPv3-Protocol-Support.pdf
20. Cisco, HSRP for IPV6. 2016. . https://www.cisco.com/c/en/us/td/docs/ios-xml/ios/ipapp_fhrp/configuration/15-s/fhrp-15-s-book/ip6-fhrp-hsrp.pdf
21. First Hop Redundancy Protocols Configuration Guide, Cisco IOS Release 15M&T - VRRPv3 Protocol Support [Cisco IOS 15.5M&T] - Cisco, last accessed 2025/05/11.
22. Routing Configuration Guide, Cisco IOS XE Everest 16.6.X (Catalyst 9500 Switches) - Configuring Bidirectional Forwarding Detection, https://www.cisco.com/c/en/us/td/docs/switches/lan/catalyst9500/software/release/16-6/configuration_guide/b_166_rtng_9500_cg/b_166_rtng_9500_cg_chapter_00.html, last accessed 2025/02/17.

تقييم التكرار واكتشاف الأعطال

دراسة لتوفر الشبكة المعتمدة على بروتوكولات BFD و FHRP

مي البعباع ، أحمد بن حسن ، محمود منصور، ناجية بن سعود

كلية تقنية المعلومات، جامعة طرابلس، طرابلس، ليبيا
n.ben_saud@uot.edu.ly

الملخص: إن النمو الهائل للإنترنت واندماجه في الحياة اليومية يؤكدان الأهمية القصوى للشبكات المرنة. يمكن أن تتسبب انقطاعات الخدمة في خسائر مالية كبيرة وتضرر بالسمعة. تُستخدم بروتوكولات التكرار في القفزة الأولى (FHRPs) بشكل شائع لتعزيز مرونة البوابة الافتراضية وتقليل وقت التوقف عن العمل، ولكنها قد تعاني من بطء في اكتشاف الأعطال، مما يؤدي إلى فقدان الحزم. يوفر اكتشاف التوجيه ثنائي الاتجاه (BFD) آلية سريعة لاكتشاف أعطال الارتباط ومراقبة الاتصال.

تستكشف هذه الورقة المشهد المعقد لموثوقية الشبكة، وتبحث في فوائد الجمع بين BFD وثلاثة بروتوكولات FHRPs بارزة HSRP و VRRP و GLBP لتحسين أداء الشبكة وزيادة التوفر وتقليل وقت التوقف عن العمل. يعتمد التقييم على مقاييس وقت التقارب convergence time ، وفقدان الحزم packet loss ، واستخدام وحدة المعالجة المركزية CPU utilization ، واستهلاك النطاق الترددي bandwidth consumption ، تشير النتائج المستخلصة من محاكاة PNETLAB إلى أن استخدام BFD يسرع بشكل كبير من اكتشاف الأعطال ويقلل من فقدان الحزم لجميع البروتوكولات الثلاثة. حقق GLBP أسرع تقارب، بينما أظهر VRRP أقل استخدام لوحدة المعالجة المركزية. تشير النتائج إلى أن دمج اكتشاف التوجيه ثنائي الاتجاه (BFD) مع بوابات بروتوكول التكرار في القفزة الأولى (FHRP) يعزز بشكل كبير أوقات تقارب الشبكة، وبالتالي يحسن الموثوقية والاستقرار الكليين للشبكة.

الكلمات المفتاحية: توافرية الشبكة، الوفرة العالية، بروتوكولات التكرار في القفزة الأولى، اكتشاف التوجيه ثنائي الاتجاه.