

Diseño de una arquitectura de almacenamiento distribuido basada en Ceph RADOS Block Device para virtualización y nube privada institucional

Design of a distributed storage architecture based on Ceph RADOS Block Device for virtualization and institutional private cloud

Información del reporte:

Licencia Creative Commons



El contenido de los textos es responsabilidad de los autores y no refleja forzosamente el punto de vista de los dictaminadores, o de los miembros del Comité Editorial, o la postura del editor y la editorial de la publicación.

Para citar este reporte técnico:

Mares Mendoza, E. (2026). Diseño de una arquitectura de almacenamiento distribuido basada en Ceph RADOS Block Device para virtualización y nube privada institucional. *Cuadernos Técnicos Universitarios de la DGTIC*, 4(3), páginas (47 - 59). <https://doi.org/10.22201/dgtic.30618096e.2026.4.3.188>

Enrique Mares Mendoza

Dirección General de Cómputo y de Tecnologías
de Información y Comunicación
Universidad Nacional Autónoma de México

enrique.mares@unam.mx
ORCID: 0009-0008-5529-0080

Resumen

El aumento de los servicios tecnológicos universitarios llevó a plantear una solución de almacenamiento capaz de sostener ambientes de virtualización y servicios de nube privada con mayor estabilidad, crecimiento controlado y administración unificada. Con ese propósito, se diseñó una arquitectura basada en la plataforma Ceph orientada al almacenamiento en bloque, considerando las necesidades operativas de un centro de datos institucional.

La metodología consistió en revisar los requerimientos técnicos, calcular la capacidad disponible, definir la organización física y lógica del clúster, y analizar su conexión con plataformas de virtualización. Se consideraron servidores, unidades de almacenamiento de alto rendimiento, redes separadas para tráfico de acceso y operación interna, mecanismos de réplica, distribución de datos, control de permisos y criterios para mantener margen ante fallas o expansión.

Como resultado, se obtuvo una propuesta integrada por siete nodos de almacenamiento, con una distribución que permitió mejorar la tolerancia a fallas, ordenar el uso de recursos y separar responsabilidades por plataforma. También se identificaron aspectos que debían atenderse antes de una operación plena, como monitoreo continuo, protección de credenciales, políticas de crecimiento y revisión de escenarios de recuperación. Se concluyó que la arquitectura planteada ofreció una base técnica viable

para fortalecer los servicios institucionales de virtualización y nube privada, disponibles en el centro de datos, siempre que su adopción estuviera acompañada de documentación, seguimiento operativo y ajustes derivados del comportamiento real del entorno.

Palabras clave: Almacenamiento distribuido, almacenamiento en bloque, sistema de almacenamiento distribuido Ceph, tolerancia a fallas.

Abstract

The increase in university technological services led to the proposal of a storage solution capable of supporting virtualization environments and private cloud services with greater stability, controlled growth, and unified administration. For this purpose, an architecture based on Ceph and oriented toward block storage was designed, considering the operational needs of an institutional data center.

The methodology consisted of reviewing technical requirements, calculating available capacity, defining the physical and logical organization of the cluster, and analyzing its connection with virtualization platforms. Servers, high-performance storage drives, separate networks for access traffic and internal operations, replication mechanisms, data distribution, permission control, and criteria for maintaining margin in the face of failures or expansion were considered.

As a result, a proposal consisting of seven storage nodes was obtained, with a distribution that made it possible to improve failure tolerance, organize resource usage, and separate responsibilities by platform. Aspects that needed to be addressed before full operation were also identified, such as continuous monitoring, credential protection, growth policies, and the review of recovery scenarios. It was concluded that the proposed architecture offered a viable technical foundation for strengthening institutional virtualization and private cloud services, as long as its adoption was accompanied by documentation, operational follow-up, and adjustments derived from the real behavior of the environment.

Keywords: Distributed storage, block storage, Ceph distributed storage system, fault tolerance.

1. INTRODUCCIÓN

El crecimiento de los servicios digitales, la virtualización de servidores y las plataformas de nube privada incrementaron la importancia de contar con infraestructuras de almacenamiento capaces de ofrecer capacidad, disponibilidad, continuidad operativa y administración centralizada. En este escenario, el almacenamiento distribuido se consolidó como una alternativa relevante frente a esquemas tradicionales, debido a que permitió distribuir datos entre múltiples nodos, reducir dependencias de dispositivos aislados y acompañar el crecimiento progresivo de las cargas institucionales. Para una dependencia universitaria como la Dirección General de Cómputo y de Tecnologías de Información y Comunicación (DGTIC) de la Universidad Nacional Autónoma de México (UNAM), esta necesidad adquirió especial relevancia, ya que el Centro de Datos soportó servicios tecnológicos que demandaron resiliencia, escalabilidad y operación controlada para atender plataformas de virtualización y nube privada.

Las investigaciones y experiencias técnicas recientes evidencian el uso de tecnologías abiertas en instituciones académicas y centros de investigación. Higuera-Balderrama *et al.* (2023) analizaron el caso del Instituto Tecnológico de La Paz (ITLP), donde la plataforma Ceph fue seleccionada como solución

de almacenamiento distribuido de código abierto por su flexibilidad, disponibilidad e integración con Proxmox VE. De forma complementaria, Zhao *et al.* (2024) propusieron una nube privada educativa basada en OpenStack y Ceph para la gestión unificada y autoservicio de recursos de cómputo, almacenamiento y red. Asimismo, Bocchi *et al.* (2024) documentaron el uso de Ceph en CERN para OpenStack, CephFS y S3, dentro de estrategias de continuidad operativa y recuperación ante desastres. Estas experiencias sustentan su pertinencia en entornos institucionales que requieren escalabilidad, resiliencia y administración centralizada.

A partir de esos antecedentes, en febrero de 2025, se identificó la necesidad de definir una arquitectura de almacenamiento distribuido basada en Ceph RBD para el Centro de Datos de la DGTIC-UNAM, orientada a plataformas consumidoras como Proxmox, OpenStack y Hyper-V. El problema técnico no se limitó únicamente a estimar la capacidad, sino que implicó establecer un diseño coherente entre nodos físicos, unidades NVMe, configuración de red, reglas de distribución de datos, políticas de acceso y criterios operativos. Por ello, el alcance del reporte se delimitó al diseño técnico y dimensionamiento de una arquitectura RBD, sin incluir CephFS, RADOS Gateway (RGW) ni Metadata Server (MDS) como componentes funcionales, con el propósito de concentrar el análisis en el almacenamiento de bloque para discos de máquinas virtuales y volúmenes persistentes.

El objetivo general fue diseñar una arquitectura de almacenamiento distribuido basada en Ceph RBD para plataformas de virtualización y nube privada institucional, considerando capacidad, disponibilidad, red, seguridad, distribución de datos e integración operativa como base para su implementación en producción en el Centro de Datos de la DGTIC-UNAM. Decido al carácter técnico de la propuesta, el análisis se orientó a personal vinculado con la operación de centros de datos, redes, almacenamiento y virtualización

2. DESARROLLO TÉCNICO

El desarrollo técnico se centró en el análisis y dimensionamiento conceptual de una arquitectura de almacenamiento distribuido basada en Ceph RBD en el Centro de Datos de la DGTIC-UNAM. Durante el diseño se examinaron componentes físicos, lógicos, operativos y de integración necesarios para proponer una plataforma de almacenamiento en bloque orientada a virtualización y nube privada institucional. Sin embargo, no se limitó a estimar la capacidad, sino que incorporó criterios de disponibilidad, separación de tráfico, seguridad, distribución de datos, dominios de falla y administración operativa.

Para ello, Ceph se consideró pertinente gracias a que su modelo de distribución mediante CRUSH (*Controlled Replication Under Scalable Hashing*) permite organizar la ubicación de los datos y sus réplicas entre los OSD (*Object Storage Daemon*). Estudios recientes señalan que la distribución equilibrada de datos en Ceph es crítica para evitar puntos calientes (tanto lógicos como físicos); además, reduce cargas desiguales y mejora la utilización de dispositivos en escenarios de gran volumen y alta concurrencia (Miao *et al.*, 2024). Con base en este enfoque, la propuesta se delimitó al uso de RBD como servicio principal de almacenamiento en bloque, debido a que este componente permite provisionar dispositivos de almacenamiento de este tipo que pueden ser utilizados directamente por plataformas de virtualización.

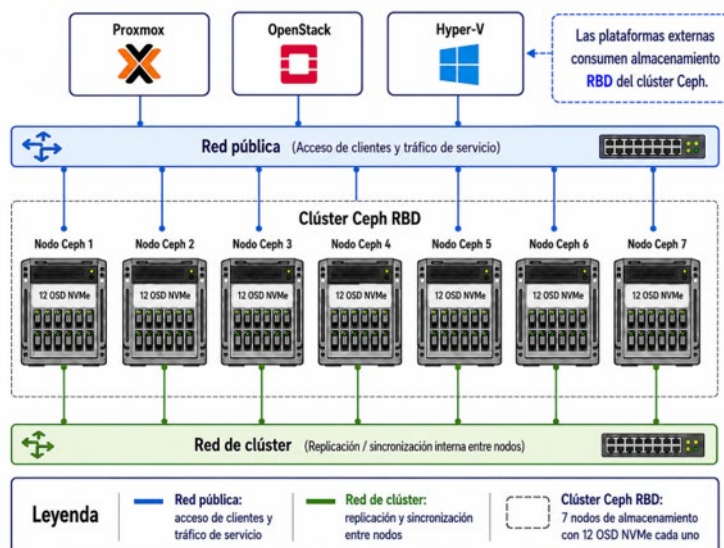
2.1 METODOLOGÍA

La metodología se estructuró en cuatro fases generales: análisis de requerimientos, dimensionamiento de capacidad, definición de arquitectura lógica y revisión de integración con plataformas como OpenStack, Proxmox y Hyper-V. Esta organización permitiría relacionar las necesidades institucionales del Centro de Datos de la DGTIC-UNAM con los componentes técnicos requeridos para una implementación productiva de almacenamiento en bloque, considerando capacidad, disponibilidad, red, seguridad, operación e integración con entornos de virtualización y nube privada.

En primer lugar, se definió una arquitectura física compuesta por siete nodos Ceph, cada uno con doce unidades NVMe de 15.36 TB, dos procesadores Intel Xeon 6747P y 2 TB de memoria RAM DDR5 por nodo. Bajo el criterio de asignar un OSD por cada unidad NVMe, se estableció un diseño conceptual de 84 OSD en total. Esta definición permitiría plantear una base física suficiente para distribuir las cargas de almacenamiento entre distintos servidores, reducir la dependencia de un único dispositivo, aislar fallas a nivel de nodo y proyectar crecimiento de manera ordenada. Además, cada nodo contempló dos tarjetas de red de cuatro puertos de 25 GbE, lo que permite disponer de ocho puertos de 25 GbE por servidor y facilita la separación del tráfico asociado al acceso de las plataformas consumidoras y a las operaciones internas del clúster, como se representa en la Figura 1.

Figura 1

Arquitectura física propuesta del clúster Ceph RBD y plataformas consumidoras



Con base en esta arquitectura, se realizó el dimensionamiento de capacidad para estimar el alcance operativo del clúster. La capacidad bruta por nodo se calculó en 184.32 TB y la capacidad bruta total en 1,290.24 TB. Bajo un esquema de réplica 3, la capacidad útil aproximada fue de 430.08 TB; sin embargo, este valor no debía considerarse como capacidad totalmente disponible para consumo.

Debido a que Ceph requiere espacio libre para procesos de recuperación, *backfill*, rebalanceo, crecimiento y tolerancia ante fallas, se propuso una ocupación operativa máxima de 80%, equivalente a 344.06 TB.

Esta estimación permitiría establecer un margen de operación más seguro para la planeación de la infraestructura, como se resume en la Tabla 1.

Tabla 1

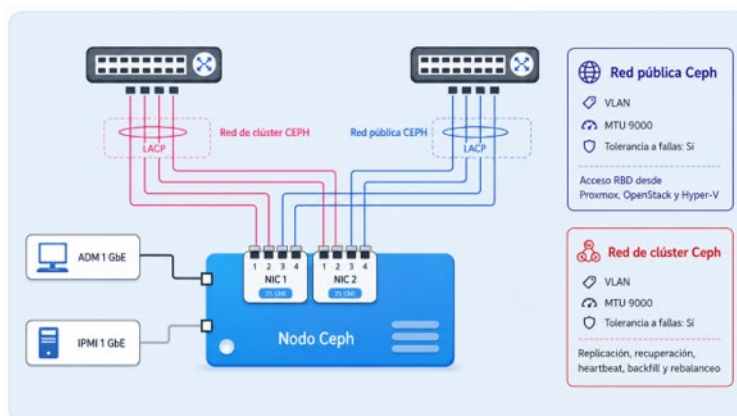
Dimensionamiento de capacidad bruta, replicada y operativa del clúster Ceph RBD

Concepto	Valor propuesto
Número de nodos	7
Discos por nodo	12 NVMe
Capacidad por disco	15.36 TB
OSD totales	84
Capacidad bruta por nodo	184.32 TB
Capacidad bruta total	1,290.24 TB
Factor de réplica	3
Capacidad útil aproximada	430.08 TB
Capacidad operativa al 80 %	344.06 TB

Posteriormente, se definió la configuración conceptual de red como un componente esencial para la operación del clúster. La red pública se destinó al acceso de las plataformas consumidoras hacia RBD, mientras que la red de clúster se reservó para procesos internos de Ceph, como replicación, recuperación, *heartbeat*, *backfill* y rebalanceo entre OSD. Asimismo, se consideraron criterios de *bonding*, LACP (*Link Aggregation Control Protocol*), MTU (*Maximum Transmission Unit*), VLAN, rutas, latencia, tolerancia a fallas y separación de dominios de *broadcast*. Esta separación permitió reducir el riesgo de que las operaciones internas del clúster afectaran el acceso de las plataformas de virtualización, especialmente en escenarios de recuperación, crecimiento o redistribución de datos, como se muestra en la Figura 2.

Figura 2

Diseño lógico de red de un nodo Ceph: separación entre tráfico público RBD y tráfico interno del clúster



En continuidad con el diseño físico y de red, se propuso desplegar cinco monitores para favorecer el *quorum* y mejorar la resiliencia ante fallas. Esta distribución permitió considerar la continuidad de los servicios de control del clúster, ya que los monitores son necesarios para mantener los mapas del estado general de Ceph, incluyendo OSD, *pools*, *placement groups*, monitores y CRUSH.

Tabla 2

Distribución propuesta de servicios MON y MGR

Nodo	Función MON	Función MGR	Rol propuesto	Observaciones de disponibilidad
ceph-node01	Sí	Sí	MON + MGR activo inicial	Nodo participante en <i>quorum</i> y punto inicial de administración del clúster
ceph-node02	Sí	Sí	MON + MGR <i>standby</i>	Nodo participante en <i>quorum</i> y respaldo del servicio MGR
ceph-node03	Sí	Sí	MON + MGR <i>standby</i>	Nodo participante en <i>quorum</i> y segundo respaldo del servicio MGR
ceph-node04	Sí	No	MON	Nodo participante en <i>quorum</i>
ceph-node05	Sí	No	MON	Nodo participante en <i>quorum</i>
ceph-node06	No	No	OSD	Nodo dedicado principalmente a servicios de almacenamiento
ceph-node07	No	No	OSD	Nodo dedicado principalmente a servicios de almacenamiento

Asimismo, se consideraron tres instancias MGR (Ceph Manager) distribuidas en nodos físicos distintos, una activa y dos en espera, debido a que este servicio aporta funciones de administración, monitoreo y exposición de métricas. Con ello, se buscó evitar la concentración de servicios críticos en un único servidor físico y fortalecer la disponibilidad operativa de la arquitectura, como se sintetiza en la Tabla 2.

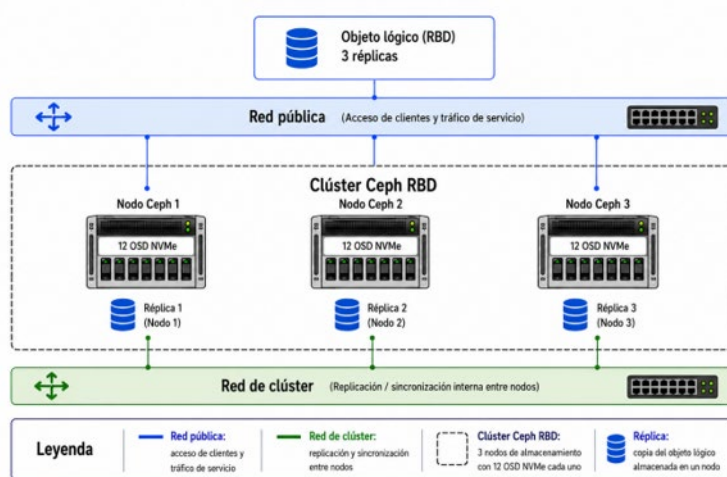
2.2 DESARROLLO DE LA ARQUITECTURA PROPUESTA

El *backend* de almacenamiento se definió sobre BlueStore, debido a que se consideró una arquitectura de *object store* diseñada para la capa RADOS de Ceph. Esta elección fue pertinente porque BlueStore permite almacenar los datos directamente en dispositivos de bloque y gestionar metadatos mediante RocksDB, evitando la dependencia de un sistema de archivos convencional en la ruta principal de datos y reduciendo limitaciones asociadas a FileStore, como la sobrecarga por *journaling* (Lee *et al.*, 2017). A partir de ello, se diseñó la distribución lógica mediante CRUSH *map*, mecanismo que permite organizar la dispersión de datos, reglas de replicación y dominios de falla sin depender de un servicio centralizado

de localización de objetos (Goggin *et al.*, 2025). En consecuencia, se propuso una estructura por *buckets* de *host* y una regla *replicated_rule* con dominio de falla por *host*, para ubicar tres réplicas en nodos físicos distintos y fortalecer la resiliencia de la arquitectura, como se ilustra en la Figura 3.

Figura 3

Distribución conceptual de réplicas RBD mediante CRUSH map con dominio de falla por host



Con base en el uso exclusivo de RBD, se propusieron *pools* separados por plataforma o función, a fin de facilitar la operación, el monitoreo y la administración del almacenamiento. Esta segmentación permitiría aplicar permisos, cuotas y políticas diferenciadas según la plataforma consumidora, además de mejorar el aislamiento operativo, la trazabilidad y el control del crecimiento. Los parámetros considerados fueron *size 3*, *min_size 2*, *pg_autoscale_mode on* y *aplicación rbd*, como se muestra en la Tabla 3.

Tabla 3

Diseño conceptual de pools RBD por plataforma consumidora

Pool	Propósito	Plataforma	Réplica	Autoscaling	Aplicación
rbd_proxmox	Discos de VM	Proxmox	3 max – 2 min	on	rbd
rbd_openstack_volumes	Volúmenes persistentes	OpenStack Cinder	3 max – 2 min	on	rbd
rbd_openstack_images	Imágenes	OpenStack Glance	3 max – 2 min	on	rbd
rbd_openstack_vms	Discos de instancias	OpenStack Nova	3 max – 2 min	on	rbd
rbd_hyperv	Discos para Hyper-V	Hyper-V	3 max – 2 min	on	rbd

Respecto a los *placement groups*, no se definió una asignación fija, ya que su cantidad depende del número de OSD, los *pools* creados, el consumo esperado, el tamaño relativo de cada *pool* y el comportamiento real del clúster. Por ello, se consideró *pg_autoscale_mode on*, puesto que la documentación oficial de Ceph indica que el *autoscaling* permite ajustar o recomendar el número de PG por *pool* con base en la utilización observada y esperada (Ceph Documentation, s. f.). Esta decisión evitaría una configuración inicial arbitraria y favorecería una administración flexible del almacenamiento en escenarios productivos institucionales, como se muestra en la Tabla 4.

Tabla 4

Estimación de placement groups por pool RBD

Pool	Uso estimado	Autoscaling	PG inicial sugerido	Criterio de ajuste	Observaciones
rbd_proxmox	Alto	on	Validar por autoscaler	Uso real y crecimiento	No fijar arbitrariamente
rbd_openstack_volumes	Alto	on	Validar por autoscaler	Capacidad asignada a Cinder	Considerar cuotas por proyecto
rbd_openstack_images	Medio	on	Validar por autoscaler	Cantidad y tamaño de imágenes	Puede crecer menos que volúmenes
rbd_openstack_vms	Alto	on	Validar por autoscaler	Uso de Nova con RBD	Considerar impacto de arranque
rbd_hyperv	Alto	on	Validar por autoscaler	Integración no nativa	Considerar impacto de arranque

Como parte de los criterios operativos, se consideraron umbrales de capacidad *nearfull*, *backfillfull*, *full* y *failsafe_full*, junto con acciones preventivas y correctivas orientadas a conservar margen de operación y evitar condiciones críticas de llenado.

En materia de seguridad, se contempló el uso de CephX con usuarios separados por plataforma o servicio, bajo el principio de privilegios mínimos. Esta decisión mejoraría el control de acceso, protegería los *keyrings*, separaría responsabilidades y mantendría trazabilidad sobre el uso del almacenamiento, como se muestra en la Tabla 5.

Tabla 5

Matriz de control de acceso CephX por plataforma

Usuario CephX	Plataforma	Pool autorizado	Permisos	Alcance
client.proxmox	Proxmox	rbd_proxmox	R/W sobre <i>pool</i> asignado	Discos VM
client.cinder	OpenStack Cinder	rbd_openstack_volumes	R/W sobre volúmenes	Volúmenes persistentes

Usuario CephX	Plataforma	Pool autorizado	Permisos	Alcance
client.glance	OpenStack Glance	rbd_openstack_images	R/W o R según flujo	Imágenes
client.nova	OpenStack Nova	rbd_openstack_vms	R/W sobre <i>pool</i> de instancias	Discos de VM
client.openstack	OpenStack general	<i>Pools</i> requeridos	Restringido por servicio	Uso administrativo limitado
client.hyperv	Hyper-V	rbd_hyperv	R/W sobre <i>pool</i> asignado	Integración Hyper-V

Finalmente, se revisaron los mecanismos de integración con las plataformas consumidoras. Proxmox se planteó mediante RBD nativo; OpenStack, a través de Cinder, Nova y Glance; e Hyper-V, mediante *rbd-wnbd*.

Tabla 6

Matriz comparativa de integración de Ceph RBD con plataformas de virtualización y nube privada

Plataforma	Servicio	Uso	Integración	Ventajas
Proxmox	QEMU/KVM	Discos de VM	RBD nativo en Proxmox	Integración directa y almacenamiento distribuido
OpenStack	Cinder	Volúmenes persistentes	Driver Ceph RBD	Volúmenes distribuidos y escalables
OpenStack	Nova	Discos de instancias	libvirt/QEMU con RBD	Reduce dependencia de disco local
OpenStack	Glance	Imágenes	<i>Backend</i> RBD opcional	Uso eficiente con Cinder/Nova
Hyper-V	VM	Discos para VM	<i>rbd-wnbd</i>	Exposición de RBD a Windows

Esta revisión permitió identificar las condiciones técnicas necesarias para incorporar el almacenamiento en bloque a los entornos de virtualización y nube privada considerados en el diseño, como se compara en la Tabla 6.

3. RESULTADOS

Se logró, como resultado principal, definir un clúster de siete nodos con doce unidades NVMe de 15.36 TB por nodo, equivalente a 84 OSD bajo el criterio de un OSD por dispositivo. Este dimensionamiento permitió estimar 1,290.24 TB de capacidad bruta y 430.08 TB útiles con réplica 3; no obstante, se identificó una capacidad operativa cercana a 344.06 TB para conservar margen ante recuperación, *backfill*, rebalanceo

y crecimiento. El análisis consolidó una propuesta de arquitectura Ceph RBD para virtualización y nube privada institucional. Los resultados deben interpretarse como producto de un diseño conceptual y de dimensionamiento técnico.

Desde la perspectiva de disponibilidad, el uso propuesto de réplica 3 y una regla CRUSH con dominio de falla por *host* permitiría distribuir copias en nodos físicos distintos, por lo que el beneficio esperado sería reducir el riesgo ante fallas de disco o servidor. Este planteamiento es coherente con antecedentes que señalan la importancia de evitar dispositivos aislados y presentar almacenamiento distribuido como recurso único para plataformas de virtualización (Higuera-Balderrama *et al.*, 2023). Asimismo, la separación entre red pública y red de clúster se identificó como una decisión de diseño orientada a disminuir interferencias entre el acceso de las plataformas consumidoras y el tráfico interno de replicación o recuperación.

Finalmente, los beneficios planteados, como tolerancia a fallas, administración por *pools*, separación de responsabilidades y crecimiento controlado, deben entenderse como beneficios esperados y no como mejoras cuantitativamente demostradas. Por ello, antes de una implementación productiva será necesario validar rendimiento con cargas reales, ajustar PG mediante *autoscaling*, proteger credenciales CephX, revisar umbrales de llenado y comprobar la integración de Hyper-V.

4. CONCLUSIONES

El análisis desarrollado permitió establecer que una arquitectura de almacenamiento distribuido basada en Ceph RBD representa una alternativa técnicamente viable para fortalecer la infraestructura de virtualización y nube privada institucional del Centro de Datos de la DGTIC-UNAM. Más que centrarse únicamente en la capacidad disponible, el diseño evidenció la importancia de articular de manera coherente la disponibilidad, la distribución de datos, la separación de tráfico, la seguridad y la operación continua como elementos indispensables para una plataforma de almacenamiento institucional.

Asimismo, el trabajo identificó que la solidez de la propuesta no depende de un único componente, sino de la correcta integración entre la arquitectura física, la lógica de distribución mediante CRUSH, la administración de *pools*, el control de acceso y la planeación operativa. La propuesta contribuyó a definir una base técnica ordenada para orientar decisiones futuras sobre crecimiento, integración con plataformas de virtualización y continuidad del servicio.

Se concluyó que la arquitectura propuesta debe acompañarse de una planeación operativa gradual, documentación técnica, monitoreo constante y políticas claras de administración. Para ello, se recomendó revisar el comportamiento del clúster ante cargas reales, escenarios de falla, ajustes de red, políticas de llenado, monitoreo y mecanismos de integración, con el propósito de reducir riesgos operativos y consolidar una implementación segura, escalable y administrable.

REFERENCIAS

- Bocchi, E., Lekshmanan, A., Valverde, R., & Goggin, Z. (2024). Enabling storage business continuity and disaster recovery with Ceph distributed storage. *EPJ Web of Conferences*, 295, 01021. <https://doi.org/10.1051/epjconf/202429501021>

- Ceph Foundation. (s. f.). *Ceph Block Device*. Ceph Documentation. <https://docs.ceph.com/en/squid/rbd/>
- Goggin, Z., Lekshmanan, A., Valverde, R., & Bocchi, E. (2025). Ceph at CERN in the multi-datacentre era. *EPJ Web of Conferences*, 337, 01331. <https://doi.org/10.1051/epjconf/202533701331>
- Higuera-Balderrama, D., Morejón-López, D., Canosa-Reyes, R. M., Alonso-Labrada, R., & Marín-Hernández, F. (2023). Diseño e implementación de una red de almacenamiento distribuido basada en CEPH. *Pädi Boletín Científico de Ciencias Básicas e Ingenierías del ICBI*, 11(Especial 2), 55–60. <https://doi.org/10.29057/icbi.v11iEspecial2.10839>
- Jelten, J., Wollek, A., Frank, D., & Lasser, T. (2023). *Equilibrium: Optimization of Ceph cluster storage by size-aware shard balancing*. arXiv. <https://doi.org/10.48550/arXiv.2310.15805>
- Lee, D.-Y., Jeong, K., Han, S.-H., Kim, J.-S., Hwang, J.-Y., & Cho, S. (2017). *Understanding write behaviors of storage backends in Ceph object store*. MSST 2017. <https://msstconference.org/MSST-history/2017/Papers/CephObjectStore.pdf>
- Miao, Y., Fan, Z., Zhang, M., & Yang, L. (2024). A data balanced distribution algorithm based on Ceph storage. *Microelectronics & Computer*, 41(3), 90–97. <https://doi.org/10.19304/J.ISSN1000-7180.2023.0180>
- Zhao, L., Hu, G., & Xu, Y. (2024). Educational resource private cloud platform based on OpenStack. *Computers*, 13(9), 241. <https://doi.org/10.3390/computers13090241>

ANEXO A. GLOSARIO DE TÉRMINOS TÉCNICOS

El presente glosario reúne los principales conceptos técnicos utilizados en el reporte con el propósito de facilitar la comprensión de la arquitectura de almacenamiento distribuido basada en Ceph RBD para plataformas de virtualización y nube privada institucional. Los términos incluidos corresponden a componentes, servicios, procesos y criterios operativos abordados en el diseño propuesto para el Centro de Datos de la DGTIC-UNAM.

Almacenamiento distribuido. Modelo de almacenamiento en el que los datos se reparten entre varios nodos interconectados por una red, es decir, no se guardan en un solo servidor o dispositivo.

Backfill. Proceso mediante el cual Ceph completa o redistribuye datos hacia otros OSD después de una falla, expansión o cambio en la distribución del clúster.

backfillfull. Umbral que puede limitar procesos de recuperación o redistribución cuando no existe suficiente espacio disponible para realizar *backfill* de manera segura.

BlueStore. Es el *backend* de almacenamiento utilizado por Ceph para guardar datos directamente en los discos administrados por los OSD.

Bonding. Técnica que permite agrupar varias interfaces de red físicas para mejorar redundancia, disponibilidad o capacidad agregada.

Bucket. Elemento jerárquico dentro del CRUSH *map* que agrupa dispositivos o nodos. Puede representar niveles como *host*, *rack*, sala o centro de datos.

Ceph. Plataforma de almacenamiento distribuido de código abierto que permite integrar múltiples servidores y discos en un clúster unificado, capaz de ofrecer almacenamiento de bloque, archivos u objetos.

CephFS. Es el sistema de archivos distribuido de Ceph que permite almacenar y acceder a datos en forma de archivos y carpetas, similar a un sistema de archivos tradicional, pero distribuido entre varios nodos del clúster.

CephX. Mecanismo de autenticación de Ceph que permite controlar el acceso de usuarios y servicios a los recursos del clúster.

Cinder. Servicio de OpenStack encargado de administrar volúmenes persistentes.

CRUSH. Algoritmo de Ceph que determina la ubicación de los datos y sus réplicas dentro del clúster, sin depender de una tabla centralizada de localización.

CRUSH map. Mapa lógico que describe la jerarquía del clúster, incluyendo dispositivos, *hosts*, *buckets* y reglas de distribución de datos.

Failsafe_full. Umbral adicional de protección para evitar condiciones extremas de llenado que puedan comprometer la operación del clúster.

FileStore. *Backend* anterior de Ceph que dependía de un sistema de archivos convencional y de mecanismos de *journaling*. Se menciona como antecedente frente a BlueStore.

Full. Estado crítico en el que un OSD o el clúster alcanza un nivel de llenado que puede impedir nuevas escrituras.

Glance. Servicio de OpenStack encargado de administrar imágenes de máquinas virtuales.

Heartbeat. Comunicación interna entre componentes del clúster utilizada para verificar disponibilidad, detectar fallas y actualizar estados operativos.

Keyring. Archivo que contiene credenciales de autenticación CephX. Debe protegerse adecuadamente para evitar accesos no autorizados al almacenamiento.

LACP (Link Aggregation Control Protocol). Protocolo utilizado para la agregación dinámica de enlaces de red, comúnmente empleado en configuraciones de *bonding*.

Libvirt. Capa de administración de virtualización utilizada para gestionar hipervisores como KVM/QEMU.

MDS (Metadata Server). Es el servicio de Ceph encargado de administrar los metadatos de CephFS, como nombres de archivos, carpetas, permisos y rutas. No almacena directamente los datos del usuario, sino la información necesaria para organizar y acceder correctamente a los archivos dentro del sistema de archivos distribuido.

MGR (Manager). Es el servicio de Ceph encargado de recopilar información operativa del clúster y proporcionar funciones de administración, monitoreo y visualización. Complementa a los monitores al ofrecer métricas, módulos de gestión, panel web y herramientas que facilitan la supervisión del estado, rendimiento y capacidad del clúster.

MON (Monitor). Servicio de Ceph encargado de mantener los mapas del clúster y participar en el *quorum*. Su función es esencial para conservar la consistencia del estado general del clúster.

MTU (Maximum Transmission Unit). Tamaño máximo de paquete que puede transmitirse por una interfaz de red. Su correcta configuración ayuda a evitar problemas de fragmentación o rendimiento.

Nearfull. Umbral de advertencia que indica que un OSD o el clúster se aproxima a un nivel alto de ocupación. Permite tomar acciones preventivas antes de llegar a una condición crítica.

Nova. Servicio de OpenStack encargado de la gestión de instancias de cómputo.

NVMe (Non-Volatile Memory Express). Es una tecnología de almacenamiento de alta velocidad utilizada principalmente en unidades de estado sólido.

OSD (Object Storage Daemon). Es el servicio de Ceph encargado de almacenar los datos en los discos del clúster. Cada OSD administra un disco o unidad de almacenamiento y participa en la distribución, réplica, recuperación y balanceo de los datos, permitiendo que Ceph mantenga la disponibilidad y tolerancia a fallos.

PG (Placement Group). Unidad lógica que agrupa objetos dentro de un *pool* y facilita su distribución entre los OSD. Su cantidad influye en el balance, desempeño y administración del clúster.

Pool. Agrupación lógica de almacenamiento dentro de Ceph. Permite separar datos por plataforma, servicio o función, por ejemplo, Proxmox, Cinder, Glance, Nova o Hyper-V.

Quorum. Mecanismo de consenso utilizado por los monitores de Ceph para mantener una visión consistente del estado del clúster. Requiere la disponibilidad de la mayoría de los monitores configurados.

RADOS (Reliable Autonomic Distributed Object Store). Es el servicio de Ceph que permite acceder al almacenamiento de objetos mediante interfaces compatibles con S3 y Swift.

RBD (RADOS Block Device). Proporciona almacenamiento en bloque y permite presentar discos virtuales o volúmenes persistentes a plataformas como Proxmox, OpenStack e Hyper-V.

Rrbd-wnbd. Herramienta que permite exponer dispositivos RBD de Ceph en sistemas Windows, y facilita su uso en entornos como Hyper-V.

RGW (RADOS Gateway). Es el servicio de Ceph que permite acceder al almacenamiento de objetos mediante interfaces compatibles con S3 y Swift.

RocksDB. Es una base de datos interna utilizada por Ceph, especialmente con BlueStore, para almacenar metadatos del OSD, tales como información sobre objetos, ubicaciones y estado del almacenamiento.