

PERBANDINGAN KINERJA SISTEM PENYIMPANAN TERDISTRIBUSI CEPH DAN RAID DALAM LINGKUNGAN PROXMOX

Akbar Usamah¹, Jafaruddin Gusti Amri Ginting^{2*}, Bongga Arifwidodo³, Eko Fajar Cahyadi⁴

^{1,2,3,4}Program Studi Teknik Telekomunikasi, Telkom University

*e-mail: jafargustiamri@telkomuniversity.ac.id

Abstrak

Kemajuan teknologi komputasi awan memunculkan kebutuhan akan sistem penyimpanan yang mampu mengelola data besar secara efisien, seperti Ceph *cluster* dan sistem RAID. Penelitian ini bertujuan membandingkan performa kedua sistem dalam aspek kecepatan transfer data dan kemampuan pemulihan setelah terjadi bencana (*failure*). Penelitian ini dilakukan dengan menginstalasi Proxmox sebagai platform virtualisasi pada dua *node*, dimana *node* pertama digunakan untuk mengelola mesin virtual Ceph dan RAID. Alat uji FIO *tools* digunakan untuk mengukur performa dalam tiga parameter utama: IOPS, *latency*, dan *bandwidth*. Hasil penelitian menunjukkan bahwa sistem RAID unggul dalam kecepatan transfer data, dengan capaian tertinggi sebesar 171 IOPS, *bandwidth* 320 KiB/s, dan *latency* baca 742 ms, serta *latency* tulis 747 ms. Di sisi lain, Ceph *cluster* mencatatkan nilai 111 IOPS, *bandwidth* 210 KiB/s, *latency* baca 1.129 ms, dan *latency* tulis 1.135 ms. Pada skenario pemulihan bencana, Ceph *cluster* menunjukkan performa yang lebih baik dengan 206 IOPS dan *bandwidth* 206 KiB/s untuk operasi baca dan tulis, sementara RAID mencatatkan performa tertinggi dengan 259 IOPS dan *bandwidth* 259 KiB/s. Berdasarkan temuan ini, pemilihan antara Ceph *cluster* dan RAID harus disesuaikan dengan kebutuhan spesifik: RAID lebih cocok bagi yang memprioritaskan kecepatan transfer data, sedangkan Ceph *cluster* menawarkan mekanisme pemulihan bencana yang lebih baik meskipun performanya cenderung menurun signifikan dalam kondisi kegagalan OSD.

Kata kunci: ceph cluster; disaster recovery; RAID.

Abstract

The advancement of cloud computing technology has created a demand for storage systems capable of efficiently managing large-scale data, such as Ceph clusters and RAID systems. This study aims to compare the performance of both systems in terms of data transfer speed and disaster recovery capabilities. This study was conducted by installing Proxmox as a virtualization platform on two nodes, where the first node managed virtual machines for Ceph and RAID. FIO tools were used to measure performance across three main parameters: IOPS, latency, and bandwidth. The results indicate that RAID excels in data transfer speed, achieving a peak performance of 171 IOPS, 320 KiB/s bandwidth, with a read latency of 742 ms and a write latency of 747 ms. On the other hand, the Ceph cluster recorded 111 IOPS, 210 KiB/s bandwidth, with a read latency of 1,129 ms and a write latency of 1,135 ms. In disaster recovery scenarios, the Ceph cluster demonstrated better performance with 206 IOPS and 206 KiB/s bandwidth for read and write operations, while RAID achieved a peak performance of 259 IOPS and 259 KiB/s bandwidth. Based on these findings, the choice between Ceph clusters and RAID systems should align with specific user needs: RAID is better suited for scenarios prioritizing data transfer speed, whereas Ceph clusters offer superior disaster recovery mechanisms, despite experiencing significant performance degradation during OSD failure conditions.

Keywords: ceph cluster; disaster recovery; RAID, OSD.

1. PENDAHULUAN

Perkembangan teknologi informasi dan komunikasi saat ini membawa tren baru dalam pemanfaatan layanan teknologi, salah satunya adalah *cloud computing*. Teknologi ini berkembang pesat karena mampu memberikan akses ke sumber daya dan aplikasi melalui jaringan internet, sehingga keterbatasan infrastruktur komunikasi yang ada sebelumnya dapat diatasi [1]. Pemanfaatan teknologi yang baik akan mendukung berbagai aktivitas, termasuk pembelajaran dan pekerjaan sehari-hari, tanpa terhalang oleh ruang maupun waktu [2].

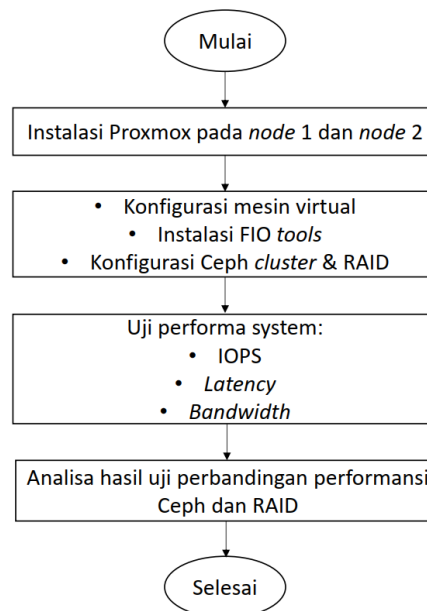
Teknologi *cloud computing* juga menjadi solusi alternatif bagi instansi pemerintah dalam mengelola data skala besar. Data tersebut membutuhkan sistem penyimpanan yang aman, mudah diakses, dan dapat diandalkan. Model *cloud computing* memungkinkan sumber daya seperti *server*, jaringan, penyimpanan, dan perangkat lunak diakses secara jarak jauh kapan saja [3]. Salah satu pendekatan yang digunakan adalah arsitektur Ceph *cluster* yang mengandalkan teknologi deduplikasi berbasis blok untuk mempercepat proses penyimpanan data tanpa mengganggu fungsi lainnya [4]. Selain itu, metode *grid* adaptif juga dikembangkan agar data dapat diproses secara paralel untuk meningkatkan efisiensi [5]. Penelitian sebelumnya telah

membahas performa sistem Ceph dan RAID dalam berbagai skenario menunjukkan bahwa deduplikasi berbasis blok pada Ceph dapat meningkatkan efisiensi penyimpanan, terutama untuk data dengan ukuran besar. Tang et al. [6] menyebutkan bahwa strategi penyimpanan berbasis Ceph cocok untuk data remote sensing dengan volume besar, meskipun terdapat tantangan pada pengelolaan metadata yang dapat memengaruhi performa.

Di sisi lain, penelitian terkait sistem RAID memberikan landasan penting dalam studi ini. Salah satu penelitian membahas evaluasi performa RAID *declustered* dibandingkan dengan ZFS RAIDZ, yang menunjukkan bahwa peningkatan performa pemulihan data berdampak pada reliabilitas penyimpanan yang lebih tinggi dibandingkan RAID tradisional [7]. Penelitian lain membahas arsitektur FreeRAID, yang mengoptimalkan interaksi antara RAID *controllers* dan SSD untuk memperpanjang umur penyimpanan berbasis *flash*. Arsitektur ini memisahkan ruang penyimpanan utama dan ruang eksploitasi data, yang keduanya berfungsi untuk melayani data normal dan data eksploitasi secara efektif [8].

Berdasarkan permasalahan tersebut, penelitian ini bertujuan untuk membandingkan performa Ceph *cluster* dan RAID dalam hal kecepatan transfer data dan kemampuan pemulihan dari bencana. Penelitian ini diharapkan memberikan kontribusi dalam pemilihan sistem penyimpanan yang tepat sesuai dengan kebutuhan spesifik pengguna, baik untuk kecepatan transfer data maupun mekanisme pemulihan yang handal. Struktur paper ini terdiri dari beberapa bagian, dimulai dengan pendahuluan, metode penelitian, hasil dan pembahasan, kesimpulan, serta daftar pustaka.

2. METODE

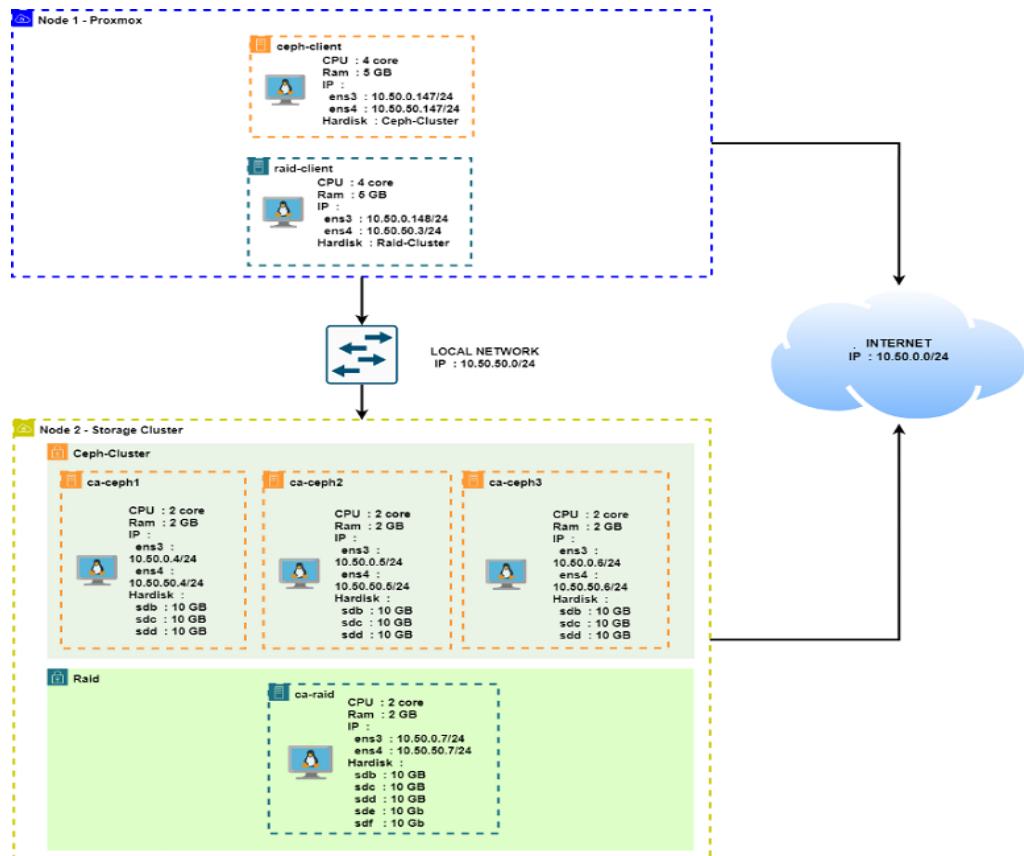


Gambar 1. *Flowchart* penelitian.

Proses penelitian ini diawali dengan instalasi Proxmox pada dua *node* yang digunakan sebagai *platform* virtualisasi. *Node* pertama dikonfigurasi untuk mengelola Ceph *client* dan RAID *client*, di mana kedua mesin virtual ini dilengkapi dengan alat uji FIO *tools* untuk mengukur performa sistem penyimpanan. Setelah konfigurasi selesai, pengujian dilakukan untuk mengukur tiga parameter utama, yaitu IOPS (*Input/Output Operations per Second*), *latency* sebagai waktu tunda operasi baca dan tulis, serta *bandwidth* yang menunjukkan kecepatan transfer data dari kedua sistem.

Hasil pengujian dari kedua sistem penyimpanan ini, Ceph *cluster* dan RAID, kemudian dianalisis untuk melihat performa masing-masing dalam hal kecepatan transfer data dan kemampuan pemulihan bencana. Analisis ini dilakukan secara komprehensif untuk memahami perbedaan signifikan antara kedua sistem.

Proses ini memastikan bahwa setiap tahapan, mulai dari instalasi, konfigurasi, pengujian, hingga analisis hasil, berjalan sistematis sesuai dengan tujuan penelitian.



Gambar 2. Topologi jaringan.

Topologi jaringan yang digunakan dalam penelitian ini terdiri dari dua *node* dengan konfigurasi sistem yang berbeda. *Node* pertama difungsikan sebagai sistem Proxmox yang bertugas mengelola mesin virtual Ceph *client* dan RAID *client*. Kedua mesin virtual ini dipasang dengan alat uji FIO sebagai perangkat pengujian performa sistem penyimpanan. Mesin virtual tersebut akan diarahkan ke sistem penyimpanan masing-masing (Ceph atau RAID) untuk mengukur parameter performa secara independen. Sementara itu, *node* kedua digunakan untuk instalasi sistem penyimpanan virtual, baik untuk Ceph *cluster* maupun RAID. Pada sistem Ceph *cluster*, terdapat tiga mesin virtual yang dikonfigurasi dengan kapasitas penyimpanan besar, mencapai 90 GB jika ketiganya digabungkan. Mesin virtual ini bertindak sebagai *server* penyimpanan untuk pengujian performa. Sedangkan pada sistem RAID, mesin virtual digunakan sebagai perbandingan utama dalam pengukuran performa penyimpanan, baik dalam hal kecepatan transfer data maupun kemampuan pemulihan data. Dengan desain ini, kedua sistem penyimpanan diuji dalam lingkungan yang sama untuk memberikan hasil perbandingan yang akurat dan objektif. Topologi ini dirancang agar pengujian dapat dilakukan secara sistematis dan menghasilkan data performa yang komprehensif.

3. HASIL PENELITIAN

Bagian ini menjelaskan hasil pengujian performa antara Ceph *cluster* dan sistem RAID berdasarkan dua parameter utama, yaitu kecepatan transfer data dan kemampuan pemulihan bencana (*disaster recovery*). Data hasil pengujian ditampilkan pada Tabel 1, Tabel 2, dan Tabel 3.

3.1 Hasil Kecepatan Transfer Data

Tabel 1 menunjukkan perbandingan kecepatan transfer data antara Ceph *cluster* dan RAID dalam lima percobaan terpisah untuk operasi baca dan tulis. Pada sistem Ceph *cluster*, kecepatan transfer data untuk operasi baca dan tulis pada percobaan pertama tercatat stabil di 210 KiB/s. Pada percobaan kedua, terjadi penurunan kecil pada kecepatan baca menjadi 195 KiB/s dan kecepatan tulis menjadi 180 KiB/s. Penurunan yang lebih signifikan terjadi pada percobaan ketiga, di mana kecepatan baca turun drastis menjadi 99 KiB/s dan kecepatan tulis menjadi 92 KiB/s. Tren penurunan ini berlanjut pada percobaan keempat, dengan kecepatan baca hanya mencapai 67 KiB/s dan tulis 70 KiB/s. Pada percobaan kelima, performa Ceph *cluster* sedikit membaik dengan kecepatan baca sebesar 94 KiB/s dan tulis 65 KiB/s.

Di sisi lain, sistem RAID menunjukkan performa yang lebih konsisten dan unggul dalam setiap percobaan. Pada percobaan pertama, kecepatan baca dan tulis RAID mencapai 320 KiB/s. Pada percobaan kedua, performa RAID masih stabil di 320 KiB/s untuk baca dan sedikit turun ke 330 KiB/s untuk tulis. Pada percobaan ketiga, meskipun terjadi sedikit penurunan, RAID tetap mencatat performa tinggi dengan kecepatan baca 340 KiB/s dan tulis 329 KiB/s. Namun, pada percobaan keempat, performa RAID menurun lebih signifikan menjadi 127 KiB/s untuk baca dan tulis. Pada percobaan kelima, performa RAID mencatat kecepatan transfer data 145 KiB/s untuk operasi baca dan tulis.

Table 1. Hasil *transfer rate*.

<i>Total Size</i>	<i>Num Job</i>	<i>Block Size</i>	<i>Ceph Cluster (read/ write)</i>	<i>RAID (read/ write)</i>
1	1	1	210/210 KiBs	320/320 KiBs
2	1	1	195/180 KiBs	320/330 KiBs
3	1	1	99/92 KiBs	340/329 KiBs
4	1	1	67/70 KiBs	127/127 KiBs
5	1	1	94/65 KiBs	145/145 KiBs

3.2 Hasil Pemulihan Bencana (*Disaster Recovery*)

Tabel 2 menunjukkan hasil pengujian pemulihan bencana pada sistem Ceph *cluster* yang menggambarkan bagaimana perubahan performa sistem terjadi saat mengalami berbagai kondisi kegagalan *object storage daemon* (OSD). Pengujian ini mencakup parameter kecepatan pemulihan (*recovery speed* - IOPS) dan *bandwidth* (KiB/s) untuk operasi baca dan tulis. Secara keseluruhan, hasil pengujian menunjukkan bahwa sistem Ceph *cluster* mampu mempertahankan performa yang cukup baik dengan penurunan kecil ketika hanya terjadi satu kegagalan OSD. Namun, peningkatan jumlah OSD yang gagal secara signifikan mengurangi performa sistem, meskipun terdapat mekanisme pemulihan yang dapat meningkatkan performa kembali pada kondisi kerusakan yang lebih parah.

Tabel 2. Hasil *recovery rate* Ceph *cluster*.

<i>Disk</i>	<i>System Condition</i>	<i>Recovery rate Read/ Write (IOPS)</i>	<i>Bandwidth Read/ Write (KiB/s)</i>
0	Health	206/206	206/206
1		205/205	205/205
2		93/93	93/93
3		90/90	90/90
4		115/115	115/115
5		188/188	188/188
0 and 3		163/163	164/164
1 and 4		95/95	95/95
2 and 5		181/181	181/181

Tabel 3 menunjukkan hasil pengujian pemulihan bencana pada sistem RAID yang menggambarkan perubahan performa sistem dalam berbagai kondisi kegagalan *disk*. Pengujian ini mencakup parameter kecepatan pemulihan (*recovery speed* - IOPS) dan *bandwidth* (KiB/s) untuk operasi baca dan tulis. Pada kondisi normal tanpa adanya kegagalan disk, sistem RAID mencapai performa optimal dengan kecepatan pemulihan tertinggi dan *bandwidth* terbaik, yaitu sebesar 259 IOPS dan 259 KiB/s untuk operasi baca dan tulis.

Tabel 3. Hasil *recovery rate* RAID.

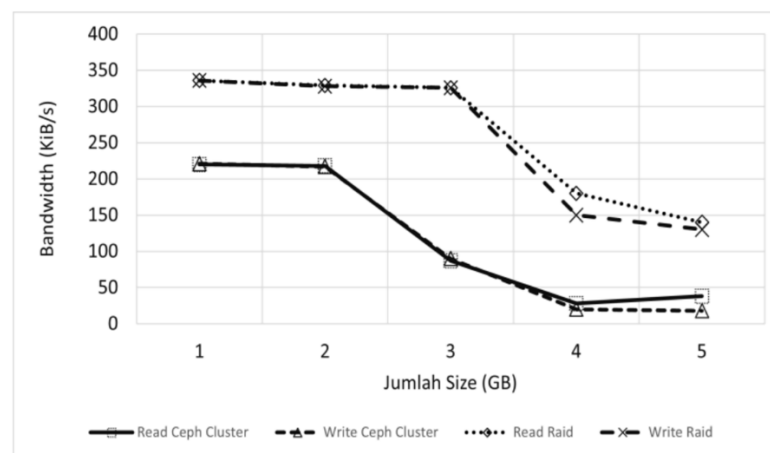
<i>Disk</i>	<i>System Condition</i>	<i>Recovery rate Read/ Write (IOPS)</i>	<i>Bandwidth Read/ Write (KiB/s)</i>
0	Health	259/259	259/259
1		250/249	250/250
2		196/197	197/197
3		219/219	220/219
4		238/237	239/237
5		161/169	169/169

Disk	System Condition	Recovery rate Read/ Write (IOPS)	Bandwidth Read/ Write (KiB/s)
0 and 3		172/172	172/172
1 and 4		164/164	164/165
2 and 5		150/150	150/151

4. DISKUSI

A. Analisa Performansi *Transfer Rate Ceph Cluster* dan RAID

Gambar 3 menunjukkan perbandingan performa *bandwidth* antara Ceph *cluster* dan RAID untuk operasi baca dan tulis dengan ukuran data yang bervariasi. Pada percobaan pertama dengan ukuran data 1 GB, Ceph *cluster* mencatatkan *bandwidth* sebesar 220 KiB/s (baca) dan 221 KiB/s (tulis), sedangkan sistem RAID mencapai 336 KiB/s untuk kedua operasi. Performa Ceph *cluster* terlihat baik pada ukuran data kecil, namun perbedaannya mulai signifikan pada percobaan berikutnya.



Gambar 3. *Transfer rate Ceph cluster* dan RAID.

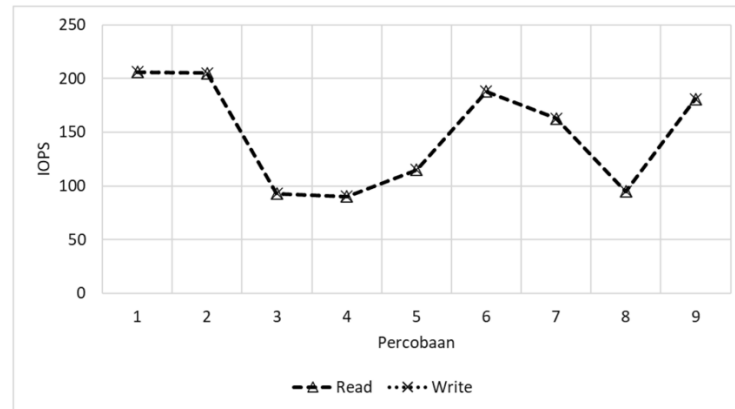
Pada ukuran data 2 GB, RAID tetap unggul dengan *bandwidth* 329 KiB/s (baca) dan 328 KiB/s (tulis), sementara Ceph *cluster* mengalami penurunan kecil menjadi 218 KiB/s (baca) dan 217 KiB/s (tulis). Penurunan drastis pada performa Ceph *cluster* terjadi pada ukuran data 3 GB, dimana *bandwidth* turun menjadi 87 KiB/s (baca) dan 90 KiB/s (tulis). Sebaliknya, RAID masih stabil dengan *bandwidth* 326 KiB/s untuk kedua operasi. Hal ini disebabkan oleh arsitektur Ceph *cluster* yang menggunakan sistem penyimpanan terdistribusi, sehingga *overhead* komunikasi antar *node* dan pengelolaan metadata semakin membebani performa ketika ukuran data bertambah besar. Pada percobaan keempat dengan ukuran data 4 GB, Ceph *cluster* mencatatkan performa terendahnya, yakni 28 KiB/s (baca) dan 20 KiB/s (tulis), sedangkan RAID tetap unggul dengan *bandwidth* 180 KiB/s (baca) dan 150 KiB/s (tulis). Pada percobaan kelima dengan ukuran data 5 GB, Ceph *cluster* mengalami sedikit peningkatan pada operasi baca menjadi 38 KiB/s, namun performa tulis tetap rendah di 18 KiB/s. Sementara itu, performa RAID menurun namun masih jauh lebih baik, dengan *bandwidth* 140 KiB/s (baca) dan 130 KiB/s (tulis).

Secara keseluruhan, hasil pengujian menunjukkan bahwa RAID memiliki performa yang lebih baik dan konsisten dibandingkan Ceph *cluster*, terutama untuk ukuran data yang lebih besar. Keunggulan ini disebabkan oleh arsitektur penyimpanan RAID yang menggunakan teknik *striping*, sehingga dapat mengoptimalkan *throughput* untuk operasi baca dan tulis. Sementara itu, Ceph *cluster* mengalami penurunan performa signifikan akibat *overhead* komunikasi antar *node* dan kompleksitas pengelolaan metadata, yang semakin terasa seiring bertambahnya ukuran data.

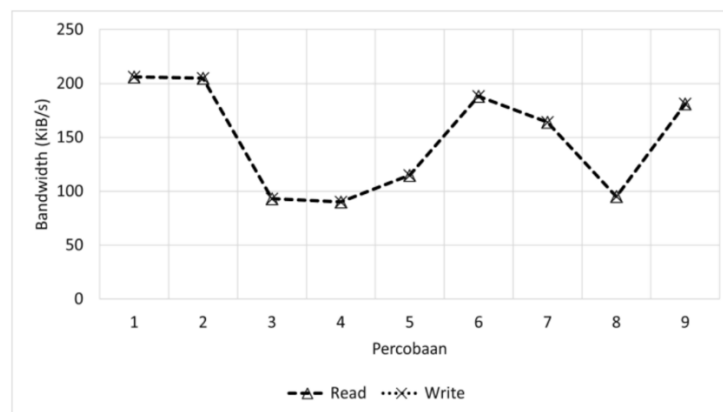
B. Analisa *Disaster Recovery Data Ceph Cluster* dan RAID

Pada Gambar 4 dan 5, pengujian pemulihan bencana pada Ceph *cluster* menunjukkan dampak dari penonaktifan OSD terhadap performa IOPS dan *bandwidth*. Pada percobaan dengan satu OSD yang dinonaktifkan, performa IOPS dan *bandwidth* terlihat relatif stabil di kisaran 205–206 pada percobaan awal. Namun, mulai dari percobaan ketiga hingga keenam, terjadi penurunan signifikan, dimana IOPS turun hingga 90 dan *bandwidth* menurun menjadi 90 KiB/s. Meskipun sempat ada perbaikan pada percobaan

kelima dan keenam, penurunan performa kembali terjadi ketika dua OSD dinonaktifkan secara bersamaan. Pada percobaan ketujuh, performa menurun ke angka 163 untuk IOPS dan 164 KiB/s untuk *bandwidth*, sedangkan percobaan kedelapan menunjukkan nilai terendah. Pada percobaan kesembilan, terjadi sedikit perbaikan performa dengan IOPS mencapai 181 dan *bandwidth* 181 KiB/s. Secara keseluruhan, hasil ini menunjukkan bahwa Ceph *cluster* mampu mempertahankan performa dengan cukup baik ketika satu OSD mengalami kegagalan, namun degradasi performa menjadi jauh lebih signifikan saat dua OSD dinonaktifkan secara bersamaan.



Gambar 4. Ceph *cluster* IOPS *disaster recovery*.



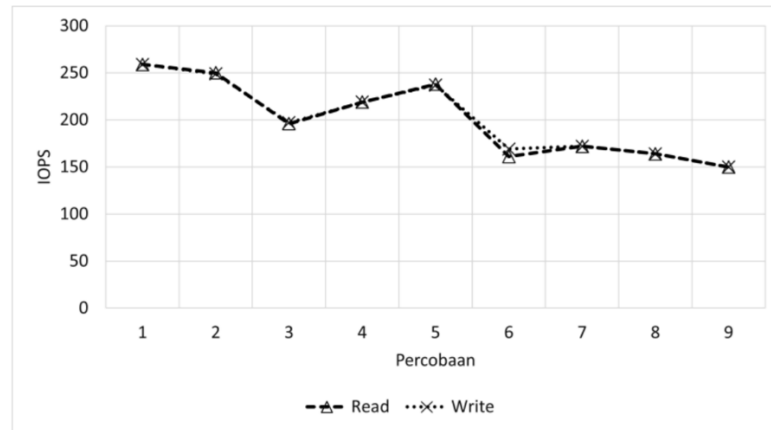
Gambar 5. Ceph *cluster* *disaster recovery* *bandwidth*.

Penurunan performa ini disebabkan oleh mekanisme *replication* dan *recovery* pada Ceph *cluster*, di mana sistem secara otomatis mereplikasi data ke OSD yang aktif untuk memastikan keandalan data [14]. Proses ini membutuhkan sumber daya komputasi tambahan, yang menyebabkan penurunan performa baca/tulis. Ketika dua OSD dinonaktifkan secara bersamaan, beban kerja sistem meningkat drastis karena proses *recovery* harus berjalan secara paralel untuk mendistribusikan ulang data yang hilang. Hal ini menciptakan *overhead* komunikasi antar *node* dan mengakibatkan performa IOPS dan *bandwidth* turun tajam, seperti yang terlihat pada percobaan ketujuh dan kedelapan. Selain itu, keterbatasan sumber daya sistem juga menyebabkan terjadinya *bottleneck* yang semakin memperlambat operasi baca dan tulis. Namun, perbaikan performa pada percobaan kesembilan menunjukkan bahwa sistem Ceph *cluster* masih mampu beradaptasi setelah proses *recovery* berjalan lebih stabil, meskipun performanya belum sepenuhnya pulih ke kondisi normal.

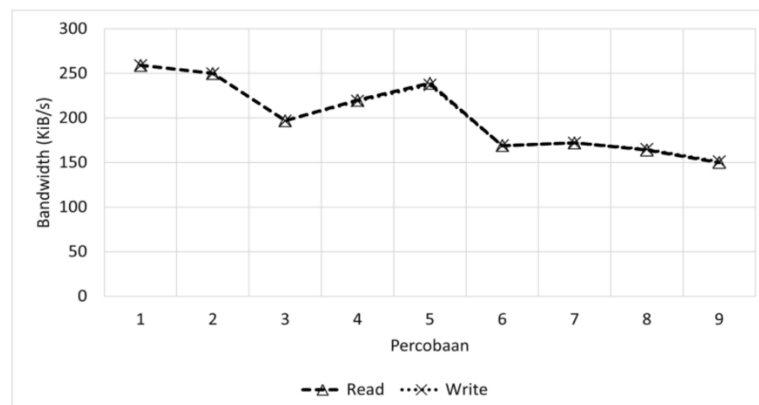
Secara keseluruhan, hasil ini menegaskan bahwa meskipun Ceph *cluster* memiliki keunggulan dalam menjaga keandalan data melalui mekanisme redundansi, performa sistem akan mengalami degradasi signifikan jika terjadi kegagalan pada lebih dari satu OSD. Oleh karena itu, perencanaan kapasitas penyimpanan dan manajemen redundansi yang baik menjadi hal penting untuk meminimalkan dampak kegagalan OSD dalam sistem Ceph *cluster*.

Berdasarkan Gambar 6 dan 7, hasil pengujian pemulihan bencana (*disaster recovery*) pada sistem RAID menunjukkan dampak dari penonaktifan disk terhadap performa IOPS dan *bandwidth*. Pada percobaan 1 hingga 5, di mana satu disk dinonaktifkan, performa baca/tulis IOPS mengalami penurunan bertahap dari 259 IOPS menjadi 196 IOPS, dengan sedikit peningkatan pada percobaan ke-4 dan ke-5. Penurunan serupa

juga terjadi pada parameter *bandwidth*, di mana nilainya turun dari 259 KiB/s menjadi 197 KiB/s pada percobaan ke-3 sebelum mengalami kenaikan kembali pada percobaan ke-4 dan ke-5. Namun, ketika dua disk dinonaktifkan secara bersamaan pada percobaan 6 hingga 9, penurunan performa menjadi lebih signifikan. Nilai IOPS turun drastis hingga mencapai 161 IOPS pada percobaan ke-6 dan mencapai titik terendah di 150 IOPS pada percobaan ke-9. Untuk *bandwidth*, pola serupa terlihat dengan penurunan yang lebih tajam. Dari nilai awal 259 KiB/s, *bandwidth* terus menurun dan mencatat angka terendah 150 KiB/s pada percobaan ke-9.



Gambar 6. RAID IOPS *disaster recovery*.



Gambar 7. RAID *disaster recovery bandwidth*.

Penurunan performa ini dapat dijelaskan oleh mekanisme kerja RAID yang mengandalkan redundansi data melalui teknik *striping* dan *parity* untuk menjaga keandalan data. Ketika satu disk dinonaktifkan, sistem RAID masih dapat beroperasi dengan baik karena data dapat dipulihkan dari disk yang tersisa. Penurunan performa yang terjadi relatif bertahap karena proses *rebuild* atau pemulihan data hanya melibatkan satu disk, sehingga tidak membebani sistem secara signifikan. Namun, ketika dua disk dinonaktifkan secara bersamaan, performa sistem menurun drastis karena proses pemulihan memerlukan sumber daya yang jauh lebih besar. RAID harus membaca dan merekonstruksi data dari beberapa disk yang masih aktif untuk menggantikan data yang hilang. Proses ini menciptakan *overhead* yang tinggi dan menyebabkan *bottleneck* pada sistem penyimpanan, sehingga performa IOPS dan *bandwidth* menurun signifikan [15]. Selain itu, penonaktifan dua disk melampaui toleransi kegagalan pada sebagian besar konfigurasi RAID, yang mengakibatkan sistem tidak dapat mempertahankan performa optimalnya.

Secara keseluruhan, hasil pengujian menunjukkan bahwa sistem RAID memiliki ketahanan yang baik ketika hanya satu disk mengalami kegagalan, dengan penurunan performa yang relatif terkendali. Namun, penonaktifan dua disk secara simultan menyebabkan degradasi performa yang lebih tajam, menunjukkan batas toleransi kegagalan pada sistem RAID. Oleh karena itu, dalam implementasi nyata, perencanaan redundansi yang baik dan pemantauan kesehatan disk secara berkala sangat penting untuk meminimalkan risiko kegagalan ganda yang dapat mengganggu performa sistem secara signifikan.

5. KESIMPULAN

Berdasarkan hasil penelitian, dapat disimpulkan bahwa sistem RAID secara keseluruhan menunjukkan performa yang lebih baik dibandingkan Ceph *cluster* dalam hal transfer data dan ketahanan terhadap kegagalan. Hasil pada Tabel 2 menunjukkan bahwa performa transfer data Ceph *cluster* mengalami penurunan signifikan dari percobaan pertama hingga kelima, dengan *bandwidth* baca dan tulis yang turun drastis dari 210 KiB/s menjadi 67 KiB/s dan 65 KiB/s. Sebaliknya, RAID, seperti yang ditunjukkan pada Gambar 4.1, mampu mempertahankan performa yang lebih stabil dan tinggi pada semua ukuran data yang diuji, bahkan pada percobaan terakhir dengan ukuran data 5 GB, di mana performa RAID tetap jauh lebih baik dibandingkan Ceph *cluster*. Dalam hal pemulihan bencana, Tabel 3 serta Gambar 4 dan 5 menunjukkan bahwa Ceph *cluster* mampu mempertahankan performa yang cukup baik dengan penurunan yang stabil ketika satu OSD gagal. Namun, penurunan performa menjadi lebih signifikan ketika dua OSD dinonaktifkan secara bersamaan, meskipun terdapat beberapa mekanisme pemulihan yang membantu meningkatkan performa. Di sisi lain, RAID, seperti yang ditunjukkan pada Gambar 6 dan 7 menunjukkan penurunan performa yang lebih terkendali ketika satu disk mengalami kegagalan, namun penurunan menjadi lebih tajam ketika dua disk dinonaktifkan, menyoroti batas toleransi RAID terhadap kegagalan disk ganda. Secara keseluruhan, meskipun Ceph *cluster* memiliki beberapa keunggulan dalam hal pemulihan bencana, performa transfer data RAID yang lebih konsisten dan stabil menjadikannya pilihan yang lebih unggul untuk berbagai kondisi operasional dan kegagalan.

UCAPAN TERIMA KASIH

Terima kasih penulis ucapkan kepada rekan-rekan dosen program studi Teknik Telekomunikasi, Fakultas Teknik Elektro, Telkom University yang telah membantu penulis dalam melaksanakan penelitian ini. Semoga hasil analisis dan pembahasan pada jurnal ini dapat memberi manfaat bagi banyak pihak terutama di bidang telekomunikasi.

DAFTAR PUSTAKA

- [1] J. Jamaluddin, "Utilization of Cloud Computing Facilities for Document and Presentation Creation," Maj. Ilm. Methoda, vol. 5, no. 2, pp. 63–68, 2015, doi: 10.17605/OSF.IO/UE9DW.
- [2] W. Hartanto, "Cloud Computing in System Development," J. Educ. Econ. J. Ilm. Educational Sciences, Econ. and Social Sciences, vol. 10, no. 2, pp. 1 –10, 2017, [Online]. Available: <https://jurnal.unej.ac.id/index.php/JPE/article/view/3810>
- [3] K. Wardani, "Cloud Computing Implementation in Government Agencies," pp. 1 –5, 2007.
- [4] H. P. Ginanjar and A. Setiyadi, "Application of Cloud Computing Technology in Product Catalog in Balatkop West Java," Komputa J. Ilm. Comput. and Inform., vol. 9, no. 1, pp. 25–33, 2020, doi: 10.34010/komputa.v9i1.3722.
- [5] J. Wang et al., "Towards Cluster-wide Deduplication Based on Ceph," 2019 IEEE Int. Conf. Networking, Archit. Storage, NAS 2019 - Proc., pp. 1–8, 2019, doi: 10.1109/NAS.2019.8834729.
- [6] X. Tang et al., "A Ceph-based storage strategy for large gridded remote sensing data," Big Earth Data, vol. 6, no. 3, pp. 323–339, 2022, doi: 10.1080/20964471.2021.1989792.
- [7] Z. Qiao, S. Liang, S. Fu, H. B. Chen, and B. Settlemeyer, "Characterizing and Modeling Reliability of Declustered RAID for HPC Storage Systems," Proc. - 49th Annu. IEEE/IFIP Int. Conf. Dependable Syst. Networks - DSN 2019 Ind. Tracks, pp. 17–20, 2019, doi: 10.1109/DSN-Industry.2019.00011.
- [8] Z. Xiao, G. E. Lester, Y. Luo, Z. Xie, L. Yu, and Q. Wang, "Effect of light exposure on sensorial quality, concentrations of bioactive compounds and antioxidant capacity of radish microgreens during low temperature storage," Food Chem., vol. 151, pp. 472–479, 2014, doi: 10.1016/j.foodchem.2013.11.086.
- [9] A. Idrus, "Designing a Owncloud Storage Server Based," J. Pinter, vol. vol 4., pp. 1 – 4, 2020. [10] Y. P. Santi, Delvia, R. Rumani M, "implementation and performance analysis of raid on data storage infrastructure as service (IAAS) cloud computing," ISSN.1412-0100, vol. 14, p. no 2, 2013.
- [10] A. Widarma and Y. Handika Siregar, "Virtualization Technology Design for Server Optimization at Asahan University," CESS (Journal Comput. Eng. Syst. Sci., vol. 4, no. 2, pp. 313–319, 2019, [Online]. Available: <https://jurnal.unimed.ac.id/2012/index.php/cess/article/view/14356>

- [11] S. Purwantoro E.S.G.S, “Performance Comparison of Clustered File System on Cloud Storage using GlusterFS and Ceph,” *INOVTEK Polbeng - Seri Inform.*, vol. 7, no. 2, p. 319, 2022, doi: 10.35314/isi.v7i2.2753.
- [12] S. B. Abdulazeez, “RAID-based Storage Systems,” no. September 2020, 2018.
- [13] A. G. Addaffi, “Implementation and Comparative Analysis of Server Virtualization Performance Using Proxmox and Openstack,” 2016.
- [14] D. Gudu and M. Hardt, “Evaluating the Performance and Scalability of Ceph Distributed Storage Systems,” pp. 177–182, 2014.
- [15] H. P. Ginanjar and A. Setiyadi, “Application of Cloud Computing Technology in Product Catalog in Balatkop West Java,” *Komputa J. Ilm. Comput. and Inform.*, vol. 9, no. 1, pp. 25–33, 2020, doi: 10.34010/komputa.v9i1.3722.