

Analisis Performa GPU pada OpenStack Nova, Zun, dan Ironic Menggunakan Glmark2, Phoronix PyTorch, dan Phoronix NAMD pada Lingkungan Komputasi Awan Fasilkom UI = GPU Performance Analysis on OpenStack Nova, Zun, and Ironic Using Glmark2, Phoronix PyTorch, and Phoronix NAMD on the Fasilkom UI Cloud Computing Environment

Immanuel, author

Deskripsi Lengkap: <https://lib.ui.ac.id/detail?id=9999920551918&lokasi=lokal>

Abstrak

Penelitian ini menganalisis dampak dari penggunaan Virtual Machine (VM), container, dan bare-metal terhadap performa Graphics Processing Unit (GPU) dengan memanfaatkan VM pada OpenStack Nova, container pada OpenStack Zun, dan bare-metal pada OpenStack Ironic. Metode virtualisasi GPU yang digunakan pada penelitian ini adalah GPU passthrough. Pengukuran performa GPU dilakukan dengan menggunakan aplikasi Glmark2 untuk menguji performa graphic rendering, Phoronix NAMD untuk menguji performa simulasi molekuler, dan Phoronix PyTorch untuk menguji performa training model. Hasil analisis menunjukkan bahwa penggunaan VM pada OpenStack Nova mengakibatkan penurunan performa GPU sebesar 15.5% pada Glmark2, 44.0% pada Phoronix NAMD, dan 8.4% pada Phoronix PyTorch. Penggunaan container pada Open Stack Zun mengakibatkan penurunan performa GPU sebesar 5.8% pada Glmark2 dan 19.7% pada Phoronix NAMD, tetapi tak ada perbedaan signifikan pada Phoronix PyTorch jika dibandingkan dengan physical machine ($\hat{I} \pm = 0.05$). Penggunaan bare-metal pada OpenStack Ironic mengakibatkan penurunan performa sebesar 1.5% pada Phoronix NAMD dan peningkatan tak signifikan sebesar -6.2% pada Phoronix PyTorch. Pengujian Glmark2 pada OpenStack Ironic dengan perlakuan yang sama seperti benchmark lainnya menunjukkan adanya penurunan performa sebesar 8.7%. Namun, perlakuan khusus pada Glmark2 OpenStack Ironic menunjukkan peningkatan performa sebesar -1.0% pada resolusi 1920x1080 jika dibandingkan dengan physical machine. Perlakuan khusus ini berupa menjalankan dummy Glmark2 dengan resolusi yang sangat rendah dan Glmark2 utama secara bersamaan. Berdasarkan hasil penelitian, dapat disimpulkan bahwa urutan computing resource dengan penurunan performa GPU yang paling minimal adalah penggunaan bare-metal OpenStack Ironic, diikuti dengan penggunaan container OpenStack Zun, dan diikuti dengan penggunaan VM OpenStack Nova.

.....This research analyzes the effects of Virtual Machine (VM), containers, and bare-metal usage on Graphics Processing Unit (GPU) performance, using VMs provided by OpenStack Nova, containers provided by OpenStack Zun, and bare-metal provided by OpenStack Ironic. The GPU virtualization method employed in this paper is GPU passthrough. GPU performance is measured using multiple benchmark applications, those being Glmark2 to measure graphic rendering performance, Phoronix NAMD to measure molecular simulation performance, and Phoronix PyTorch to measure training model performance. The results of our analysis shows that the usage of OpenStack Nova's VMs causes GPU performance slowdown of up to 15.5% on Glmark2, 44.0% on Phoronix NAMD and 8.4% on Phoronix PyTorch. Using OpenStack Zun's containers also causes GPU performance slowdowns of up to 5.8% on Glmark2 and 19.7% on Phoronix NAMD, with no significant changes on GPU performance with Phoronix PyTorch compared to the physical machine ($\hat{I} \pm = 0.05$). In contrast, using OpenStack Ironic's bare-metal causes GPU performance

slowdown of 1.5% on Phoronix NAMD and an insignificant increase in performance on Phoronix PyTorch by 6.2%. Meanwhile the results of the Glmark2 benchmark on OpenStack IroniC following the normal procedures shows GPU performance slowdown of up to 8.7%. However, the same Glmark2 OpenStack IroniC benchmark with a special procedure shows an increase in GPU performance of up to 1.0% on the 1920x1080 resolution compared to the physical machine. This special procedure involves running a dummy Glmark2 process with a tiny resolution in parallel with the main Glmark2 process. Based on the results, we can conclude that the hierarchy of computing resources in terms of minimal GPU performance slowdown starts with the usage of OpenStack IroniC's bare-metal, followed by the usage of OpenStack Zun's containers, and lastly the usage of OpenStack Nova's VMs.