

BigGPU

A DISTRIBUTED PTX COMPILE AND EXECUTION SYSTEM ON HYBRID CPU/GPU CLUSTERS



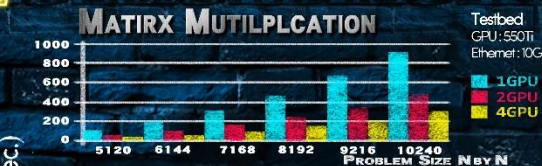
INTRODUCTION

BigGPU is a distributed PTX compilation and execution system. From the view point of users, BigGPU can be regarded as a single CUDA-compatible GPU with lots of virtual stream processors and a virtual large device memory space which physically are emulated by integrating CPUs and GPUs distributed over networks.

BigGPU can support

- Single GPU programming for hybrid CPU/GPU clusters
- Execution of large-scale kernel functions
- Dynamic load balance among heterogeneous processors

EXPERIMENTAL RESULTS



FRAMEWORK

1 REGISTER KERNEL FUNCTION



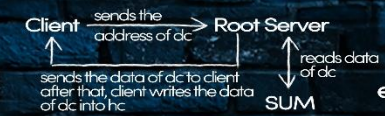
2 DEVICE MEMORY ALLOCATION AND COPY



3 LAUNCH KERNEL



4 GET EXECUTION RESULT



CONTRIBUTION

- Simplify the CUDA programming because of large-scale kernels.
- Reduce the programming complexity of GPU or hybrid CPU/GPU clusters because of no need of hybrid MPI/CUDA or MPI/OpenMP programming.
- Helpful for moving problems to hybrid CPU/GPU clusters in order for saving energy and obtaining high performance.



REFERENCES

- [1] Gregory Diamos, Andrew Kerr, Sudhakar Yalamanchili, Nathan Clark, "Ocelot: a dynamic optimization framework for bulk-synchronous applications in heterogeneous systems", PACT 2010, pp. 353-364.
- [2] CUDA Programming Guides, <http://docs.nvidia.com/cuda/#programming-guides>.
- [3] Parallel Thread Execution ISA, <http://docs.nvidia.com/cuda/parallel-thread-execution>.
- [4] Tyng-Yeu Liang, Chun-Yi Wu, Jyh-Biau Chang, Ce-Kuen Shieh, "Teamster-G: A Grid-enabled Soft-ware DSM System", CCGrid 2005, pp. 905-912.

sponsor: National Science Council (Project No. NSC-102-2221-E-151-028)

HUNG-FU LI | YU-JIE LIN | BI-SHING CHEN | TYNG-YEU LIANG



Matrix Multiplication

```

1 global __void MM(int *A, int
2 *B, int *C, int, N){
3     int i=blkDim.x*blkIdx.x+thIdx.x;
4     int r=i/N, c=i%N;
5     for(int n=0; n<N; ++n)
6         C[r*N+c]+=A[r*N+n]*B[n*N+c];
7 }
8 static void mklnts(int **h, int **d){
9     *h=(int**)malloc(65536*4);
10    cudaMalloc(d, 65536*4);
11    cudaMemcpy(d, h, h*4, cudaMemcpyHostToDevice);
12 }
13 int main(){
14     int *ha, *hb, *hc, *da, *db, *dc;
15     mklnts(&ha, &da);
16     mklnts(&hb, &db);
17     mklnts(&hc, &dc);
18     MM(<<256, 256>>>)(da, db, dc, 256);
19     cudaMemcpy(dc, hc, 65536*4, cudaMemcpyDeviceToHost);
20     free(ha); free(hb); free(hc);
21 }

```

- Software Unified Memory (SUM): a software distributed shared memory based on eager-released consistency.
- MPU: Multiple Processor Unit
- In BigGPU a device can be a CPU or GPU.

