

TSUBAME 共同利用 令和6年度 学術利用 成果報告書

TSUBAME4.0 における AI 環境の利用技術の習得
Learning of technique to use AI programing environments in TSUBAME4.0

浅見 暁

一般財団法人高度情報科学技術研究機構
<https://www.rist.or.jp/>

近年、HPCI の資源提供機関においても GPU を搭載したマシンが増えつつあり、AI を活用した業務や研究が盛んに行われている。最新の GPU (NVIDIA 社製 H100) を搭載したスーパーコンピュータ TSUBAME4.0 と富岳プリポストノードを利用し、PyTorch、Tensorflow のベンチマークを行った。データサイズによるが、富岳プリポストノードの GPU (V100) と TSUBAME4.0 の GPU (H100) では、PyTorch (YOLOv5) で約 3 倍、行列積 (fp32) では、約 10 倍 (CPU 比では 50 倍以上) の性能比であることが確認できた。

In recent years, the number of machines equipped with GPUs has been increasing at HPCI's resource-providing institutions, and AI-based work and research have been actively conducted. We benchmarked PyTorch and Tensorflow using the TSUBAME4.0 supercomputer equipped with the latest GPU (NVIDIA's H100) and the Fugaku pre-post node. Depending on the data size, we confirmed that the performance ratio of the GPU (V100) on the Fugaku pre-post node and the GPU (H100) on TSUBAME4.0 was about 3 times higher for PyTorch (YOLOv5) and about 10 times higher for matrix products (fp32) (more than 50 times higher than CPU ratio).

Keywords: 5つ程度 HPCI、TSUBAME4.0、富岳、PyTorch、TensorFlow

背景と目的

近年、HPCI の資源提供機関においても GPU を搭載したマシンが増えつつあり、AI を活用した業務や研究が盛んに行われている。また AI 分野以外の既存プログラムの GPU 化対応も進んでいる。そこで、最新の GPU (NVIDIA 社製 H100) を搭載したスーパーコンピュータ TSUBAME4.0 において、AI フレームワーク (PyTorch、Tensorflow) や AI プログラム (Resnet 等) の環境構築、性能分析、性能評価を行うことを目的とする。

性能比であることが確認できた。以降、測定結果の詳細を述べることにする。

種別	システム名	内容
GPU	TSUBAME4-h100 [T4-h100]	東京科学大(旧東工大)のTSUBAME-4.0システムのGPU-h100(資源タイプ:gpu_1)
	Fugaku-v100 [fg-v100]	理研富岳のprepostサーバのGPU-v100搭載サーバ
	Merie-A100 [Mr-A100]	RIST所有のGPU-A100搭載サーバ(称:Merie)
CPU	Fugaku-Xeon [fg-Xeon]	理研富岳のprepostサーバでのXeon(gold6240)のみ使用
	TSUBAME4-EPYC [T4-EPYC]	同TSUBAME-4.0資源タイプgpu_1でのEPYC(9654)のみ使用

図1 測定システム

概要

図1にベンチマークに用いたシステム、アプリケーション、図2に測定をおこなったベンチマークアプリケーションを示す。

TSUBAME4.0 (CPU、GPU)、富岳プリポストノード (CPU、GPU) についてベンチマークを行った結果、富岳プリポストノードの GPU (V100) と TSUBAME4.0 の GPU (H100) では、PyTorch (YOLOv5) で約 3 倍、行列積 (fp32) では、約 10 倍 (CPU 比では 50 倍以上) の

グループ	BMT名	内容	備考
pytorch	Yolo.v5	物体認識:YOLO(You Look Only Onse)のpythonコードで、サンプルデータ(bus.jpg)を使用	https://www.tankobucreate.com/post-299/
tensorflow	matmul	tensorflow実装の行列積(fp32)性能測定カーネル	
	mnist	画像解析処理でデータセットとして、手書き文字、ファッション画像、CNN(cifar10) の3種類で測定	

図2 測定ベンチマークアプリケーション

結果および考察

図3に YOLOv5 の実行時間のグラフを示す。左側 5 本が CPU (1c、4c などコア数を示す)、右側 4 本が GPU を使用した場合の測定時間となる。本ベンチマークにおいては、小規模データのためか GPU の優位性は 1.2 倍にとどまった。

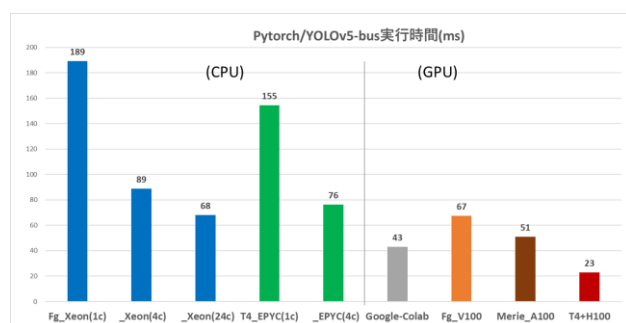


図3 PyTorch/YOLOv5-bus 実行時間 (ms) (小規模データ bus)

図4に TensorFlow (fp32 行列積) の結果を示す。左側 4 本が CPU での測定値、右側 2 本が GPU での測定値となる。CPU での測定値においては、TF_INTRA_OP_PARALLELISM_THREADS の設定によるマルチスレッドの効果を確認することができた。また GPU の測定においては、CPU のマルチスレッドをしのぐ大幅な高速化の効果を得ることができ、V100 と H100 の比較でも約 10 倍の性能差となることが分かった。また CPU (EPYC48c) との比では約 50 倍の性能差となることが分かった。

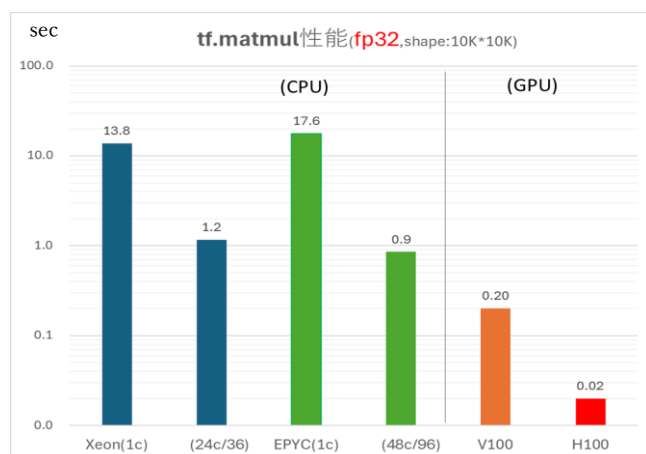


図4 Tensoeflow2.10 行列積 (fp32)

図5に TensorFlow(mnist)のベンチマーク結果を示す。単純な手書き文字(mnist)やその拡張の Fashion 画像の識別では GPU の優位性は少なく(MNIST の V100 は CPU 比で逆転)、CNN (Convolutional Neural Network)での cifar10(物体カラー画像)において、ようやく H100 の優位性が確認できた。

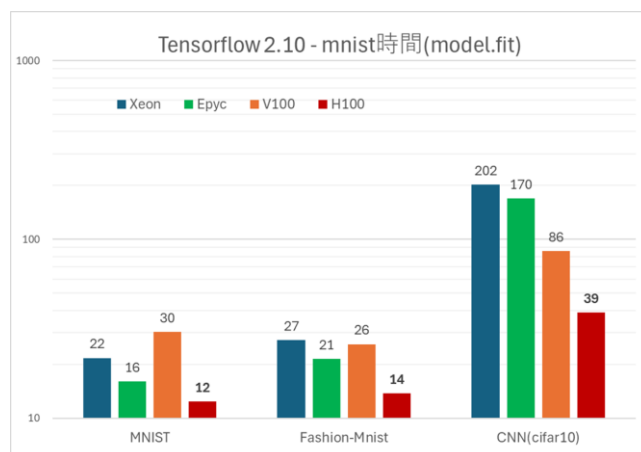


図5 Tensorflow 2.10 mnist

まとめ、今後の課題

TSUBAME-4.0 での GPU-H100 に関して、各種ベンチマーク (① PyTorch- 物体検出 YOLOv5, ② TensorFlow-matmul 行列積, ③ TensorFlow-画像認識 mnist 系 3 種) での性能を測定、他 GPU (V100, A100) や CPU (Xeon, EPYC) 等との比較を実施した。

YOLOv5 では、小規模デモデータのためか、CPU 系と比較して GPU の優位性は僅かの結果であったが、その中でも H100 の高速性は、V100 比で 3 倍弱と際立っている。

H100 の tensor 演算の高速性の確認のため、② TensorFlow-matmul(fp32:10Kx10K) の測定を実施した。その結果、CPU (Xeon, EPYC) と比較して、50 倍以上の高速性を確認できた。同様に、③画像識別 mnist 系の比較でも、手書き文字や Fashion 画像の識別処理の性能では、データ規模(演算量)が足りないためか、GPU の有利性は少なく、物体カラー画像化での CNN データにおいて、漸く GPU 系が 4-5 倍程度の高速性を確認する結果となった。

以上の結果から、GPU の優位性に関しては、②の

行列積のようなものでないとカタログ性能がなかなか出ないが、しっかりと演算処理のあるアプリであれば CPU に比べて 1 桁程度の優位性が確認できることが判明、その中でも H100 の高速性を再確認できた。

今後の課題としては、より実践的な問題での技術習得、性能測定を行っていきたいと考える。

以上