

数理・情報のフロンティア
2019 年度採択研究代表者

2020 年度 年次報告書

TRUONG THAO NGUYEN

National Institute of Advanced Industrial Science and Technology
Postdoctoral Researcher

Exploring large-scale design of distributed deep neural networks.
分散型ディープニューラルネットワークの大規模設計の調査・研究

§ 1. 研究成果の概要

The target of this project is to enable training Deep Learning model on large-scale HPC system with a short time including study (1) new parallelism strategies (not only data parallelism) (2) the method to reduce the communication time and (3) the effect of different system architecture on training time. Based on this research, we aim to develop an estimating model as the basis for a utility we name it ParaDL, that aids end-users, framework developers and system builders in (i) identifying the optimal large-scale parallelism strategies, i.e. domain decomposition, when deploying a training job for a given DL model and data set on an HPC system; (ii) guiding the implementation and possible optimizations of different parallel strategies in existing DL frameworks; and (iii) advising system architects on the best co-design choices for their system depending on the workloads they plan to run. In this fiscal year, we conduct the following research topic:

Task 1 – Improving the accuracy of the Performance/Memory Estimation Model: Our proposed initial model last fiscal year did not have a good prediction accuracy in the comparison with the execution time of the empirical result on the ABCI supercomputer. In this fiscal year, we consider using (i) an empirical parameterization, (ii) self-contention modeling, and (iii) to detach the network congestion (caused by other applications running at the same times on a shared system). With that method, we archive a good prediction accuracy (86,7% correction on average) across all parallel strategies on multiple CNN models and datasets on up to 1K GPUs. A paper [2] related to this research topic has been accepted in HPDC2021 (A* conference, 1st author).

Task 2 – Extension of our Estimation Model: We demonstrated that our estimation model is helpful in solving different emergent scalability problem of large-scale training Deep Neural Network. Firstly, we propose a memory estimation model which is one of the components of the KARMA framework that helps to train very big Deep Learning models with out-of-core methods, when model sizes cannot fit into the local memories of computing nodes [1] (published in A* conference-SC2020, 3rd authors). We also extended our core estimated model to make it work with different network architectures [3] (accepted paper in a rank A-conference-CCGrid2021, 1st author). In the next fiscal year, we plan to extend our model to estimate the cost of I/O and staging during the training process (an on-going work). In this research topic, we collaborate with a researcher from RIKEN, who can help us to verify the accuracy of our model on a new supercomputer system, i.e., Fugaku.

【代表的な原著論文情報】

- [1] (A* – SC2020, 15 pages, , 3rd author) Mohamed Wahib, Haoyu Zhang, N. T. Truong, Aleksandr Drozd, Jens Domke, Lingqi Zhang, Ryousei Takano, Satoshi Matsuoka, “Scaling distributed deep learning workloads beyond the memory capacity with KARMA”, SC ’20: Proceedings of the International Conference for High Performance Computing, Networking,

Storage and Analysis, Article No 19, pp. 1-15, November 2020.

- [2] (A* - HPDC2021, 13 pages, 1st author) N. T. Truong, M. Wahib, R. Takano, A. Kahira, L. B. Gomez, R. M Badia “An Oracle for Guiding Large-Scale Model/Hybrid Parallel Training of Convolutional Neural Networks”, accepted at ACM Symposium on High-Performance Parallel and Distributed Computing (HPDC2021) - to be appeared in June 2021.
- [3] (A - CCGRID2021, 10pages, , 1st author) N. T. Truong, M. Wahib, “An Allreduce Algorithm and Network Co-design for Large-Scale Training of Distributed Deep Learning”, accepted at 2021 IEEE/ACM 21st International Symposium on Cluster, Cloud and Internet Computing (CCGrid2021) - to be appeared in May 2021.