



HOKKAIDO UNIVERSITY

Title	Study on Robust Large-scale Neural Network Optimization [an abstract of dissertation and a summary of dissertation review]
Author(s)	張, 恩智
Degree Grantor	北海道大学
Degree Name	博士(情報科学)
Dissertation Number	甲第16502号
Issue Date	2025-03-25
Doc URL	https://hdl.handle.net/2115/94990
Rights(URL)	https://creativecommons.org/licenses/by/4.0/
Type	doctoral thesis
File Information	Enzhi_Zhang_abstract.pdf, 論文内容の要旨



学 位 論 文 内 容 の 要 旨

博士の専攻分野の名称 博士（情報科学） 氏名 張 恩智

学 位 論 文 題 名

Study on Robust Large-scale Neural Network Optimization
(堅牢な大規模ニューラルネットワーク最適化に関する研究)

Detecting and avoiding overfitting in deep neural networks is of high importance. By analyzing training regimes, especially the validation progression during training, experts developed techniques to solve the overfitting problem while improving training. Those methods include, but not limited to, learning rate schedulers, dropout, adversarial training, data augmentation, and Adam optimizers. Despite the experts intuition, theoretical, and/or empirical insight of the high dimensionality and non-linearity of the deep models, a natural question arises: Can we enable the gradient descent optimizer to decrease the validation loss, even if zero training loss indicates learning has stopped.

To address this problem, meta-learning and especially meta-HPO (Hyper Parameter Optimization) methods define a bi-level objective framework. Typically, there are two common approaches for such bi-level optimization. In the first approach, we optimize the outer target and let the controller search the parameter space. Then, we improve the controller using the meta-info from the inner optimization. This method is easy to implement, for example, grid search, yet is costly to compute due to the computationally intense inner optimization. In the second approach, we update the outer objective directly: a meta-gradient-based approach calculates the high-order gradient to the hyperparameters or settings. However, despite the impractical cost of the high-order Hessian matrix, the degradation of the meta-gradient is also serious. Thus, decoupling the meta-gradient while training the inner model directly and efficiently is needed.

Based on the practical limitations discussed above, this paper practically exploits the validation loss during the gradient descent process. To achieve this target, we mainly tried 3 directions: training process, data augmentation, and model structure. First in the context of training process, exploit means using the past validation knowledge in the form of selecting the action with the maximum value under the given state. Naturally, this notion of exploitation lends itself conceptually to Reinforcement Learning (RL). Baring similarity to the work, we commonly treat the gradient-based optimization as the RL task or MDP (Markov Decision Process). In particular, by treating the weight and gradient spaces as the state and action spaces, respectively. Second in the context of data augmentation, we inversed gradients are used to generate fake samples, which are then used for data augmentation during training. For the model structure side, especially for the transformer based pretrained models, we dynamically adjusts the image patches based on Gaussian smoothing and quadtree division techniques, significantly reducing the sequence length in the transformer model without losing segmentation accuracy.