

Neural Architecture Search を取り入れた時系列予測モデル運用手法の検討

高橋 佑里子†

田宮 豊‡

小口 正人†

† お茶の水女子大学

‡ 富士通株式会社

1 はじめに

近年のクラウドサービスでは、物理サーバ (Physical Machine: PM) の低 CPU 使用率が課題となっており、これを改善すべく、事業者では、サーバを仮想化することで CPU 使用率を向上させ、PM 数を削減する取り組みが行われている。この取り組みでは、PM が自身の CPU 資源を超えた仮想 CPU を割り当てられるオーバーコミット状態に陥ることで、仮想マシン (Virtual Machine: VM) の性能が低下する可能性がある [1]。これを防ぐことを目的として、図 1 のように VM の CPU 使用率を予測モデルによって予測し、その結果に基づいて制御を行う技術が知られている [2]。

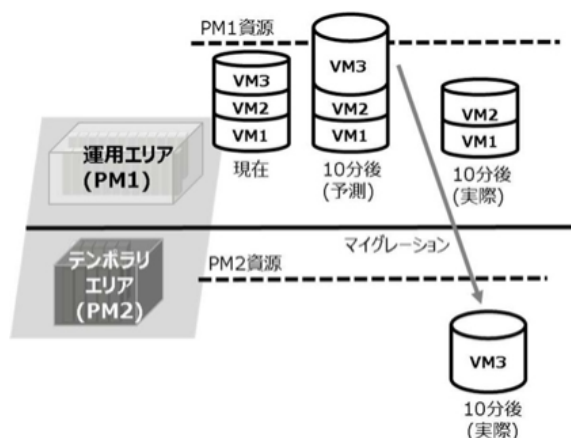


図 1: VM 制御のイメージ

VM 上で多く実行されるアプリケーションは時間が経つにつれ変化していくため、それに伴って実行される VM のワークロードも変化していく。そのため、環境の変化に合わせて予測モデルを継続的に学習し、モデルを更新することで予測精度を担保することが望ましいが、従来は学習させるデータを変えるのみで、予測モデルのネットワーク構造を変えることはなかった。しかし、学習させるデータによって最適な予測モデルのネットワーク構造は異なる。

A Study on Operational Methods of Time Series Prediction Models Incorporating Neural Architecture Search

†Yuriko Takahashi

‡Yutaka Tamiya

†Masato Oguchi

†Ochanomizu University

‡FUJITSU LIMITED

そこで本研究は、Neural Architecture Search(NAS) を取り入れることで、VM の状態変化に応じて予測モデルの構造を変化させながら予測モデルを運用する手法の検討を行う。NAS とは、ニューラルネットワークの構造自体を最適化することである [3]。この NAS を予測モデルの継続学習時のフローとして適切なタイミングで組み込むことで、より高精度な VM の CPU 使用率予測が可能になると考えた。

2 提案手法

VMCPU 使用率予測モデル運用における提案手法のフローチャートを図 2 に示す。黄色部分が準備段階、オレンジ色部分が運用段階であり、パラメータはクラスタリングの基準作成時のクラスタ数の $n_clusters$ 、類似度の最大値の $max_similarity$ の 2 種類である。

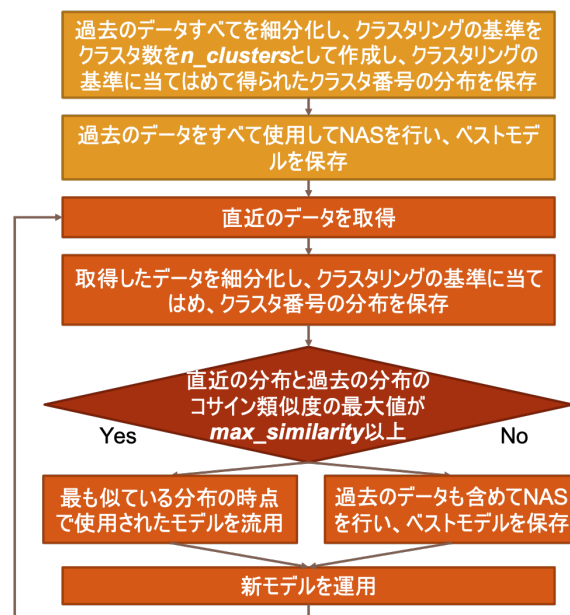


図 2: フローチャート

準備段階では、過去のデータすべてを細分化し、クラスタリングの基準の作成を行い、クラスタリングの基準に当てはめて得られたクラスタ番号の分布を保存する。その後、過去のデータすべてを使用して NAS を行い、ベストモデルを保存する。

運用段階では、まず直近のデータを取得して細分化

を行った後、準備段階で作成したクラスタリングの基準に当てはめ、得られたクラスタ番号の分布を保存する。そして、直近の分布と過去の分布のコサイン類似度の最大値が $max_similarity$ 以上だった場合は、NASを行う必要がないため、その時点で使用されたモデルを流用する。そうでない場合は、過去のデータも含めてNASを行い、ベストモデルを保存して使用する。これを一定間隔で繰り返すというものである。

3 実験

3.1 データセットの作成と実験準備

本研究では、Microsoft 社が提供している Azure の VM トレースデータセット [4][5] の一部の、"avg cpu" 項目を使用した。異なる傾向を持つデータセットを作成する目的で、データセットを正規化した後 360 点に細分化し、その中から類似度の低い 100 種類の波形を抽出し、これらを異なる割合で含む傾向を 20 種類作成し、これらを準備段階の 10 種類と運用段階の 10 種類に分け、準備段階で必要な作業を行った。

3.2 実験方法

実験では、提案手法を $n_clusters$ を 30、類似度の最大値の $max_similarity$ を 0.92 として行った。上記のようにパラメータを設定したところ、NAS を行う回数は 3 回となった。NAS は、ライブラリとして AutoKeras[6][7] を使用して行い、モデルの評価指標には RMSE を使用した。また、提案手法以外に、過去のデータで学習したモデルを使用する場合、準備段階で作成した過去のデータで NAS 行い出力されたベストモデルを使用する場合、各区間で毎回学習を行う場合、各区間で毎回 NAS を行う場合 (過去データ有無で 2 種類)、1 つ前の値を繰り返しただけのデータ (repeat 値) との比較も行った。

3.3 結果

実験結果は図 3 のようになった。縦軸は RMSE の平均値、横軸は区間番号である。この結果から、NAS を行う場合は過去のデータを使用することが重要であること、NAS を行うことで通常の学習よりもモデルの精度を向上させられること、提案手法が少ない NAS の回数で高精度を維持していることが読み取れる。

4 まとめと今後の予定

NAS を取り入れた時系列予測モデル運用に向けて提案手法を示し、それに沿って実験を行った。その結

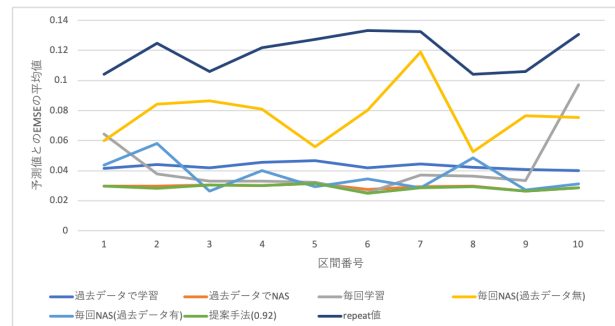


図 3: 実験結果

果、NAS を取り入れることでより高精度な予測が可能になることが判明した。

今後は、提案手法のパラメータを変えた実験や、パラメータの再利用の検討、フローチャートの改善に取り組みたいと考えている。

謝辞

本研究の一部はお茶の水女子大学と富士通株式会社との共同研究契約に基づくものであり、JST CREST JPMJCR1503 の支援を受けたものである。

参考文献

- [1] Rahul Ghosh and Vijay K Naik. Biting off safely more than you can chew: Predictive analytics for resource over-commit in iaas cloud. In *2012 IEEE Fifth International Conference on Cloud Computing*, pages 25–32. IEEE, 2012.
- [2] 児玉宏喜, 鈴木成人, 福田裕幸, 吉田英司, et al. マイグレーションを利用したデータセンタの高効率運用手法の提案とオーバコミット時における vm の性能評価. *研究報告システムソフトウェアとオペレーティング・システム (OS)*, 2018(13):1–7, 2018.
- [3] Thomas Elsken, Jan Hendrik Metzen, Frank Hutter, et al. Neural architecture search: A survey. *J. Mach. Learn. Res.*, 20(55):1–21, 2019.
- [4] Mohammad Shahradd, Rodrigo Fonseca, Íñigo Goiri, Gohar Chaudhry, Paul Batum, Jason Cooke, Eduardo Laureano, Colby Tresness, Mark Russinovich, and Ricardo Bianchini. Serverless in the wild: Characterizing and optimizing the serverless workload at a large cloud provider. *arXiv preprint arXiv:2003.03423*, 2020.
- [5] Azure/azurepublicdataset: Microsoft azure traces. <https://github.com/Azure/AzurePublicDataset>.
- [6] Haifeng Jin, Qingquan Song, and Xia Hu. Auto-keras: An efficient neural architecture search system. In *Proceedings of the 25th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining*, pages 1946–1956. ACM, 2019.
- [7] Autokeras. <https://autokeras.com/>.