

深層学習を用いた D2D キャッシュ制御方式

Device to Device Caching Method using Deep Learning

常清 陸¹

上山 憲昭²

Makoto Tsunekiyo

Noriaki Kamiyama

福岡大学大学院 工学研究科 電子情報工学専攻¹

Graduate School of Engineering, Fukuoka University

立命館大学 情報理工学部²

College of Information Science and Engineering, Ritsumeikan University

1. はじめに

移動端末で動画を視聴する形態が一般化したことで、セルラネットワーク (CN: cellular network) のバックホールのトラフィック負荷の急激な増大が懸念されている。バックホールの負荷を低減するため、基地局に設けられたキャッシュから動画コンテンツを配信するモバイルエッジコンピューティングが注目されているが、さらに負荷を軽減する方式として、移動端末 (MT: mobile terminal) でコンテンツをキャッシュして D2D (device-to-device) 通信で配信することが有効である。しかし MT のキャッシュ容量は有限であるため MT の移動経路上で高い需要が見込めるコンテンツを優先的にキャッシュすることが有効である。そこで本研究では、[1] で提案されている深層学習を用いてコンテンツの需要を推定することを D2D キャッシュ配信に応用し、深層学習のアルゴリズムの一つである長短期記憶 (LSTM: long short-term memory) ニューラルネットワークを用いて、移動経路上の他の MT が要求する可能性の高いコンテンツを推測し、MT にキャッシュするコンテンツを選択する方式を提案する。そして著名映画 10 タイトルに関するキーワード検索回数の時系列データを用いて、学習用・テスト用の時系列データを作成し、提案方式の需要推定部分の有効性を確認する。様々な場所、様々なコンテンツに対し個別に学習モデルを構築することは困難なことから、ある場所・あるコンテンツに対して作成した学習モデルを汎用的に様々な場所・コンテンツに適用することが望ましい。そこで異なる地点の実際の映画の検索回数の時系列データに対し LSTM を適用し、ある地点のデータを用いて構築した学習モデルの汎用性を確認する。

2. LSTM を用いた動画コンテンツの需要量推定

時間を固定長の Time Slot (TS) に分割した後、地域の最小エリア (町, 市, 区などの単位となる地域を表し、以下、DMA (Designated Market Area) と呼ぶ) の各 TS において発生した各コンテンツの視聴回数の時系列データから、各 DMA の各コンテンツの要求数の時系列データセットを作成する。そして LSTM を用いて作成した時系列データセットを対象に学習を行い、各都市 c の TS t におけるコンテンツ m の要求発生数 $\hat{x}_{c,m}(t)$ を予測する。

3. 評価条件

本稿では、著名な 10 個の映画のタイトルをキーワードとした検索回数の時系列データを、各映画の視聴回数の時系列データの代わりに用いる。Google Trends の提供する API を用いて、アメリカのカリフォルニア州 (CA: california) とニューヨーク州 (NY: new york) の DMA (市町区) それぞれある 1 箇所における著名映画 10 タイトルに関するキーワード検索回数を取得した。尚、調査期間を公開日から 100 日間とし、さらに各日の 8 分毎の検索回数を全てのデータの最大値を 100 とし、それ以外は最大値に対する比率を設定するように正規化した正規化検索回数を使用した。また、TS 一つを 8 分とし、入力データ x は連続する x タイムスロットの正規化検索回数、出力データ y を $x+1$ 番目のタイムスロットの正規化検索回数とする。今回は $x=5$ とし、各州、DMA、映画タイトルに対し、時系列データセット (x, y) を作成する。尚、入力と出力の TS が次のデータでは一つずつシフトするように作成する。学習に使用するデータを各映画タイトルの公開日から 99 日間とし、残りのおよそ 1 日分を評価データとなるようにデータセットを分割した。

データセットを学習する際の LSTM のハイパーパラメータを隠れ層の数を 15、バッチサイズを 55、エポックタイムを 100、活性化関数を linear 関数、損失関数を平均二乗誤差、最適化アルゴリズムを RMSprop に各々設定して学習を行った。また、本研究では予測モデルの評価を測る指標として次式で与えられる平均絶対誤差 (MAE: mean absolute error) を採用した。MAE は各 TS における実測値 y_i と予測値 y_i^p の差の絶対値の総和で

構成され、0 に近いほど高精度な予測ができていることを表す。

$$MAE(y_i, y_i^p) = \frac{1}{n} \sum_{i=1}^n |y_i - y_i^p| \quad (1)$$

4. 性能評価

作成した時系列データセットに対して LSTM を適用し、測定された予測値と実測値の比較により提案方式の需要推定部分の有効性を評価する。また、学習モデルの汎用性を CA および NY の学習データで構築した学習モデルに対して、CA および NY の両方で作成したテストデータを与えた時の MAE を各々比較することで評価する。

図 1 は NY における Aladdin で作成したデータセットに対して LSTM を適用した時の各 TS での予測値と実測値をプロットしたグラフである。グラフの縦軸は正規化検索回数、横軸は時間を表している。各 TS の予測値と実測値の誤差に着目すると、最大 2.5 回程程度の誤差が散見されるものの、測定期間全体で見るとおおその人気傾向は掴めていることから、提案方式の需要推定部分の有効性が確認できる。

図 2 は CA および NY の各映画で構築した学習モデルに対し、CA および NY の各映画のテストデータを与えた時の MAE を算出し、1 つの学習モデルに対して与えるテストデータを 10 タイトル分変化させた時の MAE の平均値を示す。グラフの縦軸は 10 個のテストデータにおける平均 MAE、横軸は学習で使った映画タイトルを表している。図 2 上部の緑線と赤線はそれぞれ CA で作成したテストデータに対して適用する学習モデルを NY と CA のデータで構築したものとして変化させた場合の MAE の平均値の推移であり、どのタイトルに関しても誤差は最大 0.05 以下に抑えられている。図 2 下部の青線と橙線は NY のテストデータに対して同様の操作を行った場合のグラフであり、先程と同様にどのタイトルに関しても誤差は最大 0.05 以下に抑えられていることが分かる。以上のことから、異なる地点および映画タイトルにおけるテストデータと学習モデルの依存関係は低く、CA あるいは NY の映画タイトルの近い将来の需要推定に NY あるいは CA のどちらの学習モデルでも適用可能と言える。したがって、学習モデルの汎用性を示している。

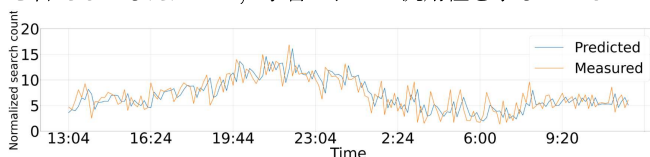


図 1: LSTM による予測値と実測値の比較 (Aladdin(NY))

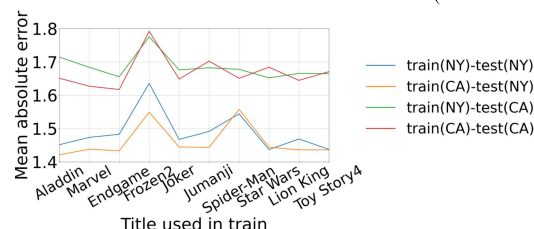


図 2: 学習モデルの汎用性

謝辞 本研究成果は、KDDI 財団研究助成寄付金 190051 の助成を受けたものである。ここに記して謝意を表す。

参考文献

[1] Arvind Narayanan, Saurabh Verma, Eman Ramadan, Pariya Babaie, and Zhi-Li Zhang, Making Content Caching Policies ‘Smart’ using the DeepCache Framework, Volume 48 Issue 5, October 2018.