

ディープラーニングによるダムポップアウトの自動検出手法の提案

八千代エンジニアリング株式会社 正会員 ○嶋本ゆり 天方匡純 藤井純一郎 安野貴人 大久保順一

1. はじめに

インフラの老朽化が顕在化しており、点検・評価・記録および適切な補修が必要不可欠である。ダムにおいては、堤体の劣化調査は主にダム技術者の目視およびスケッチによって行われており、双眼鏡でポップアウト等の損傷を1つ1つ確認していく作業は大変な労力を要する。また、目視点検では見落としや、評価に個人差が発生する可能性がある。そこで本提案は、ダム堤体をドローンにより撮影した画像に対し、AI技術を適用することで、ポップアウトの位置及び形状を自動で抽出する。この自動検出技術によって、大幅な労力の削減が期待できる。それに加え、ディープラーニングは非常に高精度に解析を行うことができるため、ポップアウトの見落としや評価のバラツキが減少する。

今回検出対象としたポップアウトは、コンクリートの劣化現象の1つで、コンクリート表面が皿状に剥がれ落ちることを指す。

2. ディープラーニングについて

ディープラーニング[Alex Krizhevsky 2012]は、人工知能(AI)技術の1つであり、従来の機械学習と比較し高精度な結果が得られるため、今日に至るまで世界中の企業や研究機関でディープラーニング方式の研究が盛んに行われている。従来手法では、あるデータの特徴量は人が定義しモデルに学習させることで学習を進めていたが、このプロセスは多くの作業量を伴い、特徴量の定義には限界がある。一方で、ディープラーニングモデルでは、機械が自ら特徴を見つけ出し学習を進めるため、人間が気づかないような複雑な特徴を見つけることができるため、非常に高精度な検出を行うことができる。

本研究で扱うようなポップアウトの検出にディープラーニングを採用すれば、ポップアウトを撮影した画像と、画像内の正しいポップアウトの位置をモデルに与えることで、ポップアウトの形状的な特徴量を自ら見つけ出すことができる。そして、特徴を

学習したモデルに、学習に用いていない新規の画像を入力すると、ポップアウトの位置を自動で抽出することができる。

3. ディープラーニングモデルの作成

ディープラーニングによる検出手法は様々あるが、今回はセマンティックセグメンテーションを用いて、ポップアウトの自動検出を行う。この手法では、画像全体や画像の一部の検出ではなく、各ピクセルに対して、そのピクセルのクラスを判定していく。そのため、細かな検出を行うことができ、正確なポップアウトの位置や大きさを判定することができる。

(1) 教師画像

本手法では、ドローンでダム堤体を撮影した画像を100枚選定し、学習データを作成した。選定した学習データの作成対象とする堤体画像についてポップアウトの領域を人手の入力により赤くマーキングし教師画像を作成する。教師画像はダム技術者がフィードバックを行い精度の高いものを作成した。



図-1 撮影画像

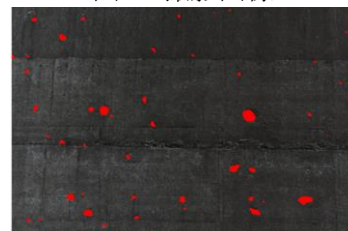


図-2 着色後の画像

(2) 最適なネットワークの選定

学習データを入力として、特徴を表現するために最適なネットワーク構造を選定するため、40枚の堤体画像および教師ラベル画像を入力画像に用いて、試行的に計算をすることで検出精度を以下の3種類のネットワークで比較検討する。最も精度の高かったネットワークを本格計算で用いる。また、少ない

キーワード 人工知能 ディープラーニング 維持管理 ダム ポップアウト

連絡先 〒111-8648 東京都台東区浅草橋 CS タワー 八千代エンジニアリング株式会社 TEL:03-5822-6843

画像で高い精度を得るため、学習済みのネットワークから一部層を取り出し、再学習を行う、転移学習の手法を用いた。

- ① Alex ネット (FCN-Alex) : 全層を畳み込む Alex ネットを基礎としたモデルで、画像 1,000 クラス 100 万枚から学習を行っている。層の深さは 8 層である。
- ② VGG ネット (FCN-VGG16) : 深いニューラルネットワーク層の末尾まで、畳み込みニューラルネットワークを導入しており、同様に 100 万枚の画像で学習を行っている。層の深さは 16 層である。
- ③ セグネット (SegNet-VGG16) : 入力画像の情報を圧縮して畳み込み (Encoder)、逆の操作で元に戻す層 (Decoder) を持ち、パラメータが層すべてに届くようになっている。Encoder は VGG モデルの一部を用いて物体の特徴情報を抽出する。それぞれの予測精度を表 1 に示す。

表 1 各ネットワークの精度比較

領域検出モデル	予測精度	予測と現実の誤差損失	計算時間
① Alex ネット	mIoU 0.5763	loss 0.0764	468 分
② VGG ネット	mIoU 0.5886	loss 0.0398	1,525 分
③ セグネット	mIoU 0.5967	loss 0.1603	832 分

ポップアウトの予測精度は、セグネットが最も高い精度を発揮している。一方、VGG ネットは、誤差損失が最小値であるものの、予測精度は中間値であり、計算時間が最長である。Alex ネットは、学習時間の面では最短であるが、予測精度が最低となっており、推奨しがたい。よって本研究ではセグネットを採用することが適切であると考えた。

(3) ネットワークの本格計算

セグネットを用いて、堤体画像 100 枚の入力画像を対象に本格計算を行う。学習の結果、表 2 のような結果が得られた。学習データ枚数が増えたこともあり、セグネットの試行モデルよりも 3 分の 1 程度に誤差損失が抑えられている。

表 2 セグネットによる精度

領域検出モデル	予測精度	予測と現実の誤差損失
① セグネット	mIoU 0.6465	loss 0.0549

また、学習中の予測精度、誤差損失の推移をみると、学習の反復とともに十分に抑制され、最終的には青線で示す予測精度が 98% 程度の高い水準に安定している。

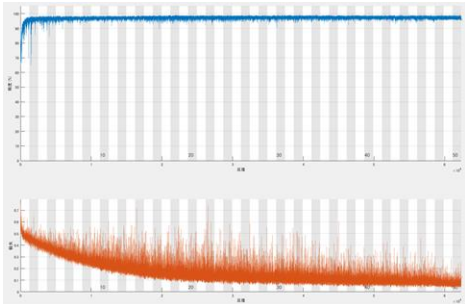


図-3 学習過程 (青線：予測精度、オレンジ：誤差)

4. 結果の出力

学習済みのモデルに、ダム堤体の撮影画像を入力し、ポップアウトの位置を自動検出する。ポップアウトの位置を概ね正しく検出することができている。



図-4 検出結果 (赤色：予測出力)

5. おわりに

本手法では、ダム堤体をドローンによる撮影した画像に対し、3 種のネットワークを比較検討し、最適なものを選定した。ネットワーク構造については世界中で研究が進められており、今後、最新のものを適用することでより高性能なモデルが作成できる。また、モデルの学習においては、最適化手法、学習回数、学習率など調整するパラメータが無数にあり、学習の結果に大きく影響を及ぼすため、様々なパターンで試行することで、より高い精度を達成できると考える。

6. 参考文献

[Alex Krizhevsky 2012] Alex Krizhevsky, Ilya Sutskever, Geoffrey E,Hinton : ImageNet Classification with Deep Convolutional Neural Network, Advances in neural information processing systems, 1097-1105 , 2012.

[Karen Simonyan 2015] Karen Simonyan, Andrew Zisserman : Very Deep Convolutional Networks for Large-Scale Image Recognition, arXiv preprint arXiv:1409.1556, 2015.

7. 謝辞

様々な新しい試みに理解を示して頂き、データやダム管理上必要なアドバイスを頂いた東北地方整備局鳴子ダム管理所の皆様には謝意を表します。