

# グラフ畳み込みを用いたカラー画像と距離画像による 高速な 3 次元物体検出

○高橋 正裕 (中央大学), Alessandro Moro (RITECS), 梅田 和昇 (中央大学)

## Fast 3D Object Detection with Color and Range Images Using Graph Convolutional Network

○Masahiro TAKAHASHI (Chuo University), Alessandro Moro (RITECS)  
and Kazunori UMEDA (Chuo University)

**Abstract:** In this paper, a lightweight 3D object detection model using color and depth images is proposed. In recent years, several studies have been done on the application of deep learning to 3D object detection. However, many of them are computationally expensive and difficult to run in real time because they deal with dense point clouds. In the proposed model, after feature extraction from the color image, a sparse point cloud is created from the range image to achieve fast object detection. Graph convolution for point clouds and feature extraction with depth information are also used. As a result, the proposed model achieved 56.4 fps when using ResNet34.

### 1. 研究背景

物体を検出・認識することは、防犯やマーケティング等、様々な分野に応用が可能である。そのため、これらを高速かつ正確に行うことができれば、カメラを用いたシステムの能力全体を向上させることができるので、重要なタスクであると考えられる。また、ステレオカメラを始めとした距離画像も計測可能なセンサが安価に手に入れることができるようになってきている。

近年、この分野に対しては、深層学習を 3 次元物体検出に応用する研究が多く行われている。その中でも PointNet[1]や PointNet++[2]は点群からの特徴抽出を可能とし、これらを物体検出や Semantic Segmentation タスクに応用した VoteNet[3]や YOLO3D[4]といったモデルが提案されている。しかし、これらの研究の多くは密な点群を扱うために計算コストが高く、リアルタイムでの運用が困難であり、求められる GPU の能力も高い。また、LiDAR 等によって得られた点群を用いることを前提にしているなど、大規模な環境を想定したものが多く、手軽に導入できるとは言えない。

そこで本研究では、RGB-D カメラから得られたカラー画像と距離画像を用いた軽量な 3 次元物体検出モデルを提案する。具体的には、カラー画像から特徴抽出を行ったのち、距離画像から疎な点群を作成することで、高速な物体検出を実現する。また、作成された点群に対してさらに GCN(Graph Convolutional Network)[5]を用い、奥行き情報を考慮した特徴抽出を行う。

本稿では、3 次元物体検出用データセットである SUN RGB-Dデータセット[6]を用いて提案モデルの物体検出

精度を評価する。また、処理速度についても Backbone Network 別に比較し、評価を行う。結果として、提案モデルの性能は VoteNet 等の物体検出手法に及ばないものの、非常に高速かつ軽量のモデルにより 3 次元物体検出を実現することができた。

### 2. 提案手法

#### 2.1 ネットワーク構造

本研究で提案するネットワーク構造を Fig. 1 に示す。ネットワークは大きく分けて、色特徴抽出部分 (Backbone Network)、グラフ作成部分、3 次元特徴抽出部分の 3 つに分けられる。

色特徴抽出部分：既存の特徴抽出モデルである VGG16[7]や ResNet[8]を用い、カラー画像から特徴抽出を行う。これにより得られた高次元の特徴をそのまま以降のモデルで用いることで、クラス分類や物体検出に必要なカラー画像からの特徴を、3 次元物体検出に活用する。また、この部分は一般的なカラー画像から得られた特徴を用いることが目的であるため、物体検出器[9-12]と同様の学習済みモデルを用いることが可能である。

グラフ作成部分：まず、カラー画像から得られた高次元の特徴を属性値として付与した点群を、KNN(K-Nearest Neighbor)によって 16 近傍に対してエッジを持つグラフ構造に変換する。この時、特徴マップに合わせてスパースな点群を生成し、そこからランダムサンプリングを行う。これにより、以降のモデルに入力する点の数を固定するとともに、ランダムサンプリングを行うことで Data Augmentation と同様の効果

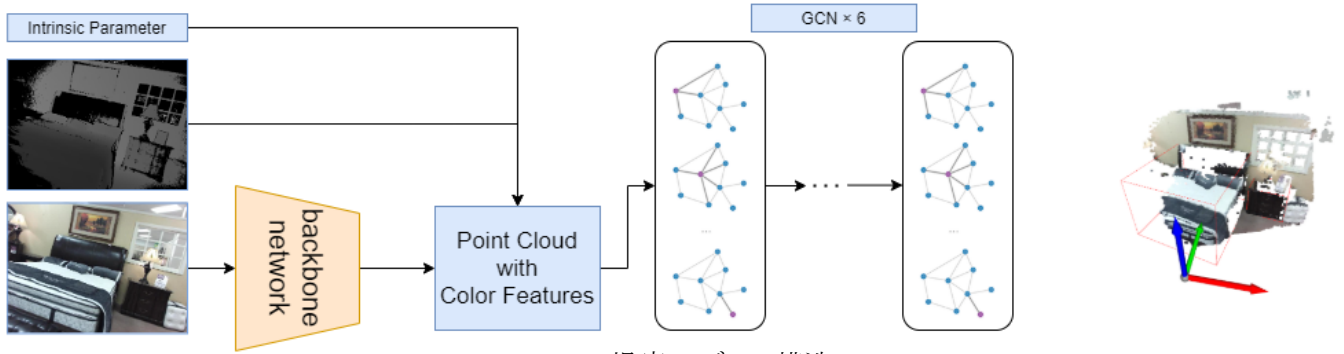


Fig.1 提案モデルの構造

を得ることができると考えられる。

3次元特徴抽出部分：グラフ構造となった点群を、6層のGCN(Graph Convolutional Network)により、近傍点の位置を考慮した、特徴抽出やバウンディングボックスの推定を行う。今回のモデルは点群の座標を扱うため、入力に負の値が存在する。そのため、活性化関数には、LeakyReLU[14]と同様に負の値を扱うことができるTanhExp[15]を用いる。出力形式はYOLO(You Only Look Once)やSSD(Single Shot Detection)に共通する手法であるUnified Detectionを元に作成した。Unified Detectionとは、バウンディングボックスの座標やクラス分類結果等をチャンネル毎に格納することで、クラス分類と領域特定を同時に出力することが可能となる手法である。Fig.2のように、提案モデルでは、出力の各3次元座標が、注目点からバウンディングボックスの中心点までのベクトル(dx, dy, dz)、バウンディングボックスの大きさ(width, height, depth)、バウンディングボックスの角度(phi, theta)、出力ベクトルの信頼度、クラス分類結果をチャンネル別で保持するように出力する。よって、出力サイズは、バッチサイズ×3次元点の数×(9+クラス数)となる。これにより、モデルに大きな分岐がなく、よりコンパクトなモデルとなるので、高速な物体検出が可能となる。

以上のようなネットワーク構造により、提案モデルはカラー画像と距離画像から3次元のバウンディングボックスを出力する。

## 2.2 学習と推論

提案モデルでは、Backbone Networkに用いるVGGやResNetは、ImageNet[16]によって事前学習を行ったものを用いる。これにより、画像中で事前に特徴のある領域がある程度提示された状態で3次元物体の推定を行うことができる。

提案モデルの出力形式はUnified Detectionであるため、誤差関数もYOLOに類似したものを学習に用いる。この誤差関数は、バウンディングボックスの中心座標へ

向かうベクトルのMSE(Mean Squared Error)、バウンディングボックスの大きさのMSE、バウンディングボックスの回転角度のMSE、ベクトルの信頼度のBCE(Binary Cross Entropy)誤差の合計で表される。よって、誤差関数は式(1)ようになる。

Loss =

$$\begin{aligned} & \lambda_{bb}\{MSE((dx, dy, dz)_{out}, (dx, dy, dz)_{tar}) \\ & + MSE((\phi, \theta)_{out}, (\phi, \theta)_{tar}) \\ & + BCE(\cos\theta_{out}, \cos\theta_{tar}) + BCE(cls_{out}, cls_{tar})\} \\ & + \lambda_{nobb}\{BCE(\cos\theta_{out}, \cos\theta_{tar})\} \end{aligned} \quad (1)$$

ここで、下付き文字 out はモデルの出力、tar は教師データを表している。cosθ は、バウンディングボックスの中心点へ向かうベクトルと正解のベクトルの角度θから得られ、式(2)によって与えられる。

cosθ =

$$\frac{dx_{out}dx_{tar} + dy_{out}dy_{tar} + dz_{out}dz_{tar}}{\sqrt{dx_{out}^2 + dy_{out}^2 + dz_{out}^2}\sqrt{dx_{tar}^2 + dy_{tar}^2 + dz_{tar}^2}} \quad (2)$$

また、λ<sub>bb</sub>とλ<sub>nobb</sub>は、それぞれ物体が存在する、もしくはしないときに計算される誤差に対する係数で、今回はλ<sub>bb</sub>=1、λ<sub>nobb</sub>=10を用いる。ここでλ<sub>nobb</sub>を大きめに設定する理由としては、ここをλ<sub>bb</sub>と同じ値を設定すると誤検出が多く発生してしまうためである。clsはクラス分類結果であり、one-hotベクトルで出力されるため、BCEによって誤差計算を行う。

得られた候補から最適なバウンディングボックスを選択する代表的な手法としては、R-CNNに用いられているNMS(Non Maximum Suppression)がある。これはIoU(Intersection over Union)と呼ばれる領域の重なり度合いを表すスコアをもとに、同じ物体に対して推定されたバウンディングボックスを消去する方法である。YOLO3D等の3次元情報を扱う物体検出器は、センサ正面方向と垂直方向の2方向から2次元に対するIoUを計算し、NMSによるバウンディングボックスの選択を行っている。しかし、この方法では同じ3次元のバウンディングボックスに対して2回IoUを計算していることとなり、GPUによる並列計算を行う上で非効率

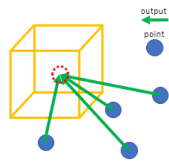


Fig.2 提案モデルの出力形式

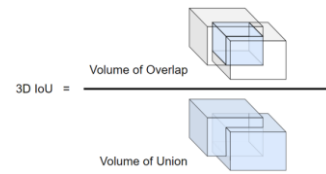


Fig.3 3D IoU

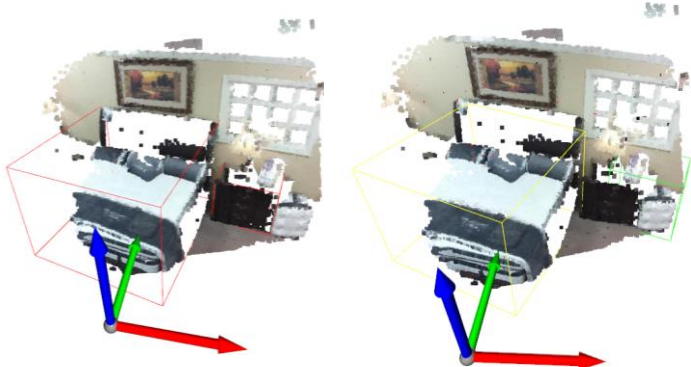


Fig.4 成功例の Ground Truth と提案モデルの出力結果

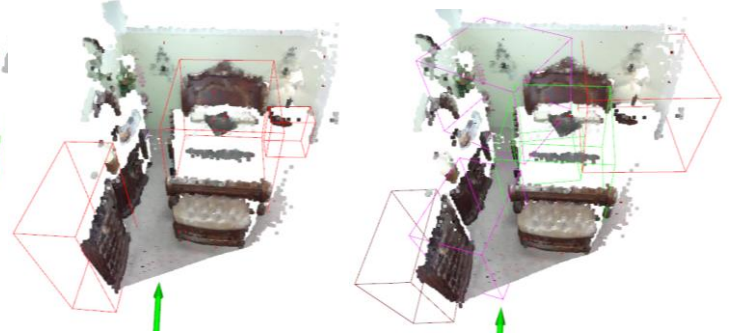


Fig.5 失敗例の Ground Truth と提案モデルの出力結果

的である。そこで、提案モデルでは、Fig.3のような体積の重なり度合いを表す 3D IoU を定義し、これをもとに NMS を実行する。これにより、IoU 計算は GPU 上で 1 回しか行わず、処理速度の改善が期待できる。NMS における IoU のしきい値は、YOLO 等においては 0.5 が用いられていたが、提案モデルでは最も良いスコアであった 0.6 を用いた。

### 3. 物体検出実験

提案モデルの有効性を確認するため、SUNRGB-D データセットを用いて精度を検証した。また、リアルタイム性についても検証するため、Backbone Network を変えたときの処理時間の比較を行った。検証に用いたマシンのスペックは、CPU が Intel Core i7 8770K、GPU が RTX2080 であった。

#### 3.1 物体検出精度検証

物体検出精度の検証として、SUNRGB-D データセットのうち、bed, table, sofa, chair, toilet, desk, dresser, night stand, bookshelf, bathtub の計 10 クラスを 100 エポック分学習させた。バッチサイズは 3、入力画像サイズは  $224 \times 224$  [pixel]、学習係数は 0.0001 に設定し、最適化手法には Adam (Adaptive moment estimation) を用いた。

物体検出の Ground Truth と提案モデルによる物体検出結果の成功例を Fig. 4 に示す。このシーンでは、bed と night stand が正解として与えられているが、どちらもかなり正確に検出できているといえる。特に bed に関しては角度や物体のサイズも正確に検出できており、大きい物体に関しても検出が可能であることがわかる。night stand に関しては物体の中心推定に誤差が生じているが、サイズを正しく推定できていることがわかる。

続いて、失敗例を Fig. 5 に示す。このシーンでは、それぞれの物体は検出できているものの、各物体の中間に誤検出が発生していることがわかる。これら 2 つの例を比較すると、点群の質に差があることがわかる。成功例の点群は各物体に対して得られている点群にあまりノイズがなく、点群がまとまっており、その結果近傍点を考慮した GCN による特徴抽出がうまくいったと言える。しかし失敗例においては物体毎に得られている点群にまとまりがなく、物体が存在しない部分に関しても点群のばらつきが多いため、近傍点探索によるグラフの作成がうまくいかなかったと考えられる。

Backbone Network を変えたときの 3D IoU による結果の平均値を Table 1 に示す。この結果によると、ResNet34 を Backbone として用いたときの結果が最も良いことがわかる。ResNet34 より深い層を持つモデルでうまくいかなかった理由としては、今回用いた入力画像のサイズが小さく、出力サイズが足りなかったことが考えられる。3D IoU のスコアとしては、0.5 を超えているため、物体をある程度正確に検出ができているといえる。しかし同じ 3 次元物体検出用モデルである VoteNet のスコアの 0.83 と比較すると、精度面では及ばなかったことがわかる。この理由としては、VoteNet では密な点群を用いて検出しているのに対し、提案モデルではスパースな点群を用いた検出であるため、その分 3 次元推定が困難になっていると考えられる。しかし、Fig. 5 のように誤検出が多いものの物体位置の特定はできているため、学習エポック数の増加や学習係数の調整次第では解決可能であると考えられる。

#### 3.2 処理速度検証

続いて、提案モデルの処理速度の検証を行った。検

Table 1 モデルの物体検出精度の結果

| Backbone name | Mean 3D IoU |
|---------------|-------------|
| VGG16         | 0.586       |
| ResNet18      | 0.603       |
| ResNet34      | 0.627       |
| ResNet50      | 0.606       |
| ResNet101     | 0.610       |
| ResNet152     | 0.581       |

証方法としては、合計 200 フレームを推論させ、その結果を用いて統計量を算出した。

Table 2 は、Backbone Network を変えたときの処理速度の結果である。この結果によると、最も軽量の VGG16 を用いたモデルは、処理速度の中央値が 85[fps]を超えており、物体検出精度検証において一番良い結果であった ResNet34 を用いたモデルも、56.4[fps]を出せているため、十分にリアルタイム性を保持できていると言える。ここで、中央値をベンチマークとした理由としては、中央値や標準偏差からわかるように、最小値が外れ値を取っているためである。

3次元物体検出を行うことが可能な他のモデルと比較しても、処理速度に関しては、提案モデルがシンプルなネットワークで構成されていることもあり、高速な検出を実現することができた。最も高速な YOLO3D も、TITAN Xを用いて 40[fps]であるため、処理速度に関してはこれらを大きく上回ったといえる。処理速度を維持したまま精度向上を行う方法としては、シンプルな数層の GCNを追加することでバウンディングボックスの候補数を絞ることが考えられる。また、これによりバウンディングボックス中心点の推定精度の向上も見込める。

#### 4. 結言

本研究では、RGB-Dから3次元のバウンディングボックスを出力可能で軽量のモデルを作成した。提案モデルでは、カラー画像から得た特徴を点群とともにGCNへ入力することで、バウンディングボックスの奥行き方向の位置と長さを取得した。

今後の展望としては、更なる深層化やパラメータ調整を行うことで、精度向上を図る。

#### 参考文献

[1] C. R. Qi, H. Su, K. Mo, and L. J. Guibas, "Pointnet: Deep learning on point sets for 3d classification and segmentation," in Proc. of

Table 2 処理速度検証結果

| Backbone name | Speed [fps] |     |      |        | Standard deviation |
|---------------|-------------|-----|------|--------|--------------------|
|               | max         | min | mean | median |                    |
| VGG16         | 91.6        | 4.6 | 79.4 | 86.7   | 13.6               |
| ResNet18      | 69.5        | 4.6 | 60.1 | 64.7   | 9.7                |
| ResNet34      | 60.7        | 4.6 | 53.0 | 56.4   | 7.6                |
| ResNet50      | 21.9        | 4.0 | 20.6 | 21.1   | 1.5                |
| ResNet101     | 19.0        | 2.7 | 17.9 | 18.1   | 1.3                |
| ResNet152     | 16.8        | 4.0 | 16.1 | 16.3   | 1.0                |

- the IEEE Conference on Computer Vision and Pattern Recognition (CVPR), pp. 652–660, 2017.
- [2] C. R. Qi, H. Su, and L. J. Guibas, "Pointnet++: Deep hierarchical feature learning on point sets in a metric space," Conference on Neural Information Processing Systems (NeurIPS), pp. 5099–5108, 2017.
- [3] C. R. Qi, O. Litany, K. He, and L. J. Guibas, "Deep Hough Voting for 3D Object Detection in Point Clouds," in arXiv preprint arXiv:1904.09664v2, 2019.
- [4] W. Ali, S. Abdelkarim, M. Zahran, M. Zidan, and A. E. Sallab, "YOLO3D: End to end real time 3D Oriented Object Bounding Box Detection from LiDAR Point Cloud," arXiv preprint arXiv:1808.02350, 2018.
- [5] T. N. Kipf and M. Welling, "Semi-Supervised Classification with Graph Convolutional Networks," in 5<sup>th</sup> International conference on Learning Representations (ICLR), 2016.
- [6] S. Song, S. Lichtenberg, and J. Xiao, "Sun rgb-d: A rgb-d scene understanding benchmark suite," in Proc. of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR), pp. 567–576, 2015.
- [7] K. Simonyan and A. Zisserman, "Very Deep Convolutional Networks for Large Scale Image Recognition," in arXiv preprint arXiv:1409.1556, 2014.
- [8] K. He, X. Zhang, S. Ren, and J. Sun, "Deep Residual Learning for Image Recognition," in Proc. of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR), 2016.
- [9] R. Girshick, J. Donahue, T. Darrell, and J. Malik, "Rich feature hierarchies for accurate object detection and semantic segmentation," in Proc. of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR), pp. 580–587, 2014.
- [10] S. Ren, K. He, R. Girshick, and J. Sun, "Faster R CNN: Towards real time object detection with region proposal networks," in Proc. of Conference on Neural Information Processing Systems (NeurIPS), 2015.
- [11] J. Redmon, S. Divvala, R. Girshick, and A. Farhadi, "You Only Look Once: Unified, Real

- Time Object Detection,” in Proc. of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR), pp. 779–788, 2016.
- [12] J. Redmon and A. Farhadi, “YOLOv3: An Incremental Improvement,” in arXiv preprint arXiv:1804.02767, 2018.
- [13] W. Liu, D. Anguelov, D. Erhan, C. Szegedy, S. Reed, C. Y. Fu, and A. C. Berg, “SSD: Single Shot MultiBox Detector,” in arXiv preprint arXiv:1512.02325, 2015.
- [14] A. L. Maas, A. Y. Hannun, and A. Y. Ng, “Rectifier nonlinearities improve neural network acoustic models,” in International Conference on Machine Learning (ICML), p. 3, 2013.
- [15] X. Liu and X. Di, “TanhExp: A Smooth Activation Function with High Convergence Speed for Lightweight Neural Networks,” in arXiv preprint arXiv:2003.09855v2, 2020.
- [16] J. Deng, W. Dong, R. Socher, L. J. Li, K. Li, and L. F. Fei, “ImageNet: A large-scale hierarchical image database,” in Proc. of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR), pp. 248–255, 2009.