

物体検出における学習データの品質改善手法

提出年月日 2022 年 9 月 16 日

代表執筆者
ながい しんご
長井 真吾
日本アイ・ビー・エム株式会社
テクノロジー・サービス

共同執筆者
わかばやし たけひろ
若林 丈紘
日本アイ・ビー・エム株式会社
テクノロジー・サービス

原稿量
本文 8,300 字
要約 1,200 字
図表 9 枚

<キーワード>

深層学習、物体検出、学習データ、ラベル、品質改善

<要約>

1. 背景と目的

深層学習による画像認識は、様々な分野において実用化への取り組みが増えてきているが、実用化にあたって、最も重要な基準の 1 つとなるのがモデルの予測精度である。モデルの精度向上には、学習データの品質が大きく影響するものの、これを改善する効率的な手法は確立されておらず、改善にあたっては膨大な時間と手間がかかるのが現状である。本論文では、深層学習の中でも適用範囲が広い物体検出において、モデル作成段階で品質の悪い学習データを効率的に特定し、修正するための手法の確立を目的とする。

2. 提案手法の概要

提案手法は 2 つのステップで構成される。まず、学習データセットから、モデルの精度に悪影響を及ぼす可能性のあるラベルを 3 種類の外れ値検出方法により特定し、修正候補とする。次に、それらの修正候補に対して、ラベル範囲の修正案を生成し、ユーザーに提示する。修正案はユーザーからのフィードバック情報を使って、ユーザーの意図に近づくように最適化を図る。

3. キーポイント

物体検出におけるラベリングは膨大な時間と手間のかかる作業であり、作業による品質のばらつきが生じるという課題があった。また、自動ラベリングする手法を用いた場合においても、ラベリングの間違いが発生するため、品質は担保できず、人手による事後確認と修正が必要となる。

提案手法では、学習データの物体位置を示すラベル範囲に着目し、ラベルの縦横サイズ、縦横比、ラベル内の色情報分布の 3 つの特徴量に対して外れ値検出を行い、大量のラベルから修正候補を自動的に絞り込むことを可能とした。さらに、修正候補に対して、平均値と過去のユーザーの修正結果に基づいた修正案を生成し、ユーザーに提示することにより、精度向上のためのラベル修正の時間と手間を大幅に軽減した。

4. 計画

提案手法の有効性を確認するため、実業務課題を想定して、血液の顕微鏡画像として公開されているデータセットを使用し、手動のラベリング作業で発生し得るラベルのノイズをシミュレートした上で、実証実験を計画した。ノイズを含む学習データセットとそれに対して提案手法により品質改善した学習データセットでそれぞれモデルを作成し、精度を比較評価することにより、提案手法がモデルの精度に及ぼす効果を検証した。

5. 評価

モデルの精度評価にあたり、物体検出で一般的に使用される **mAP (mean Average Precision)** を精度指標として採用した。学習データセットにノイズを含む状態で作成したモデルに対して、提案手法を適用して作成したモデルの精度は **mAP** で 4.2% 向上し、学習データの品質改善の効果が確認できた。その際、全 3898 個のラベルを確認することなく、提案手法により修正候補を 46 個のラベルに絞り込むことで、ラベル修正の時間と手間を大幅に軽減できた。

目 次

1. はじめに	4
2. 学習データの品質の重要性	4
2.1 物体検出における学習データ	4
2.2 学習データの品質の課題	4
2.3 従来技術	5
3. 学習データの品質改善手法	5
3.1 修正候補ラベルの特定	5
3.2 ラベル修正案の提示	6
4. 実証実験と考察	7
4.1 データセット	7
4.2 評価指標	7
4.3 実験方法	7
4.4 実験結果と考察	8
5. おわりに	9

1. はじめに

深層学習による画像認識は、製造業での外観検査や医療における自動診断など、様々な分野で実証実験を経て、実用化への取り組みが増えてきている[1]。深層学習のような先進的な技術を取り入れることは他社との競争優位性につながり、企業の DX 実現においても重要な役割を果たす。

実用化にあたって、最も重要な基準の1つがモデルの予測精度であり、モデルの精度は実用化後も継続的に向上させる必要がある。モデルの精度の向上には、適切なニューラルネットワークの選択、適切な前処理の適用、ハイパーパラメータの最適化、などが有効であることが知られており、すでに確立された手法が存在する。一方で、学習データの品質は精度に大きく影響するものの、これを改善する効率的な手法は確立されておらず、改善にあたっては膨大な時間と手間がかかるのが現状である。

本論文では、深層学習の中でも、特に適用範囲が広く、外観検査や自動診断でも使用されている物体検出に焦点を当て、物体検出における学習データの品質改善の新しい手法を提案し、実証実験により、モデル精度に及ぼす効果を検証する。

2. 学習データの品質の重要性

2.1. 物体検出における学習データ

物体検出は、画像内に複数の関心対象の物体が含まれ、それぞれの物体の位置とクラスを特定するタスクである。物体検出モデルを作成するためには、画像に加えて、関心対象の物体を矩形などでラベリングして生成されるアノテーションファイルが学習データとして必要となる。

図1に、画像をラベリングして生成されるアノテーションファイルの例を示す。例として使用した画像は、血液の顕微鏡画像で、BCCD Dataset[2]として公開され

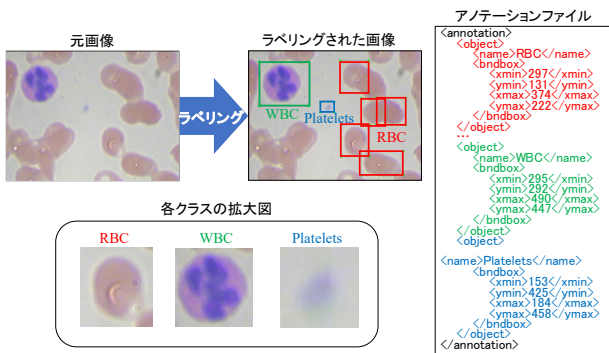


図1 ラベリングにより生成されるアノテーションファイルの例

ているデータセットに含まれるものである。画像内には、赤血球 (RBC: Red Blood Cell)、白血球 (WBC: White Blood Cell)、血小板 (Platelets) の3種類の細胞があり、これらの位置を示すラベルを付与している。ラベリング作業を補助するソフトウェアとしては、LabelImg[3]などがあり、これにより画像ファイルへ簡単にラベルを付与することができ、機械学習分野で標準的に使用されているフォーマットのアノテーションファイルが生成される。アノテーションファイルには画像のファイル名やサイズなどのメタ情報とともに、画像に含まれる物体ごとのクラス名と位置を示す座標が記録される。

物体検出モデルは、このようにしてラベリングした学習データを大量に集めたデータセットを準備し、学習データへの予測誤差が小さくなるように繰り返し学習して作成する。

2.2. 学習データの品質の課題

学習データへのラベリングは、対象となる物体に対して、一つ一つラベル範囲とラベル名を指定する必要があり、膨大な時間と手間のかかる作業である。そのため、複数人で作業を分担するケースが多く、さらに作業者も、データ提供者、モデル開発者、外部の専門業者、及びその組み合わせ、など様々なケースがあり、作業者による品質のばらつきが避けられない。

図2に、実際のラベリング作業でよく見られ、モデルの精度にも悪影響を与えるような品質の悪いラベルの例を示す。悪いラベルの例1や2では、ラベルが小さすぎたり、横幅が狭すぎて、WBCがラベルに入りきっておらず、周辺の背景との境界も分からない状態となっている。また、悪いラベルの例3では、WBCがラベルに入りきっているものの、ラベル内にWBCと周辺の背景以外の物体が映り込んでおり、この物体の特徴まで学習対象に入ってしまう状態となっている。

このようなラベリング工程で生じた品質の悪いラベルは、良いラベルと特徴量において想定外の差が生じ、ノイズとなる可能性がある。学習データのノイズはモデ

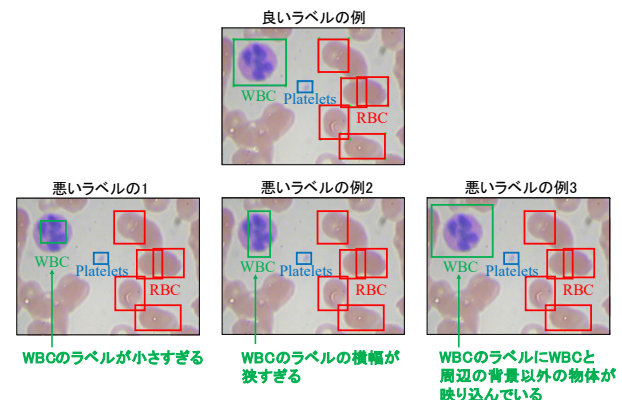


図2 モデルの精度に悪影響を与えるラベルの例

ルの精度に悪影響を及ぼすため、モデルの作成にあたっては学習データの品質を担保することが重要となる[4]。そのため、本論文では、モデル作成段階で品質の悪い学習データを効率的に特定し、修正するための手法の確立を目的とする。

2.3. 従来技術

深層学習の中でも、画像分類の分野では、品質の悪い学習データが混ざった中でモデルの精度を向上させる研究成果がいくつか報告されている[5]。画像分類は、画像全体からクラス分類する最も単純な部類のタスクである。画像全体に 1 つのクラスを割り当てて学習するため、個々の物体の位置情報を考慮する必要がない点で、物体検出とは異なる。従って、画像に割り当てるクラス名の間違いしかノイズの要因とはならず、ノイズへの対策が比較的容易なものとなっている。

それに対して、物体検出では位置情報もノイズ要因として加わるため、より複雑なものとなり、研究事例が少ないのが現状である。1 つのアプローチとしては、人手によるラベリングを最小限にして、ラベリング作業を自動化する技術が開発されている。例えば、画像認識の深層学習モデル作成を支援するソフトウェア製品である IBM Maximo Visual Inspection[6]では、そのような自動ラベリング機能を提供している。自動ラベリング機能では、ユーザーが手動で作成した最低数量の学習データをもとに、初期モデルを作成した後、残りの画像ファイルをモデルに予測させ、予測結果を利用してラベルを付与する。自動付与されたラベルの品質は、ラベリング用に作成した初期モデルの精度次第であり、多くの場合は、ラベル漏れ、ラベル付け違い、ラベル範囲のずれ、などのラベリングの間違いが発生する。そのため、自動化によりある程度の手間は軽減できるものの、品質を担保するためには、人手による事後確認と修正は必要となる。

3. 学習データの品質改善手法

提案手法は 2 つのステップで構成される。まず、学習データセットから、モデルの精度に悪影響を及ぼす可能性のあるラベルを 3 種類の外れ値検出方法により特定し、修正候補とする。次に、修正候補に対して、ラベル範囲の修正案を生成し、ユーザーに提示する。修正案はユーザーからのフィードバック情報を使って、ユーザーの意図に近づくように最適化を図る。

3.1. 修正候補ラベルの特定

最初のステップとして、学習データセット内の修正候補となるラベルを以下の 3 種類の外れ値検出方法により特定する。2.2 節で例示したように、物体検出にお

いては、ラベル範囲がモデルの精度に影響するため、ラベルの縦横サイズ、縦横比、ラベル内の色情報分布の 3 つの特徴量に着目し、外れ値検出を実施する。

(1) ラベルの縦横サイズにおける外れ値検出

アノテーションファイルから取得したラベルの縦横サイズにおける外れ値を検出し、該当するラベルを修正候補のリストに追加する。外れ値の検出には、スミルノフ・グラブス検定[7]のような統計手法を使用する。外れ値検出に使用する統計手法や、必要となるパラメーターは、ユーザーが任意に設定することもできる。

スミルノフ・グラブス検定は、平均値からのずれを標準偏差で割った値をもとに、外れが大きいものから順に 1 回の検定で 1 つの外れ値を除き、外れ値がなくなるまで検定を繰り返す手法である。検定統計量は以下の計算式で算出され、この値が有意点の値よりも大きければ外れ値と判定される。

$$T_i = \frac{|x_i - \mu|}{\sigma}$$

T_i : ラベル i の検定統計量

x_i : ラベル i の標本値

μ : ラベル全体の平均値

σ : 不偏標準偏差

図 3 に、WBC ラベルの縦横サイズの分布と縦横サイズの外れ値検出により外れ値と判定されたラベルの例を示す。この例では、極端に小さいサイズのラベルが分布の中心から離れており、有意点の値から外れていることが統計的検定で認められれば、修正候補のリストに追加されることになる。

(2) ラベルの縦横比における外れ値検出

アノテーションファイルから取得したラベルの縦横比

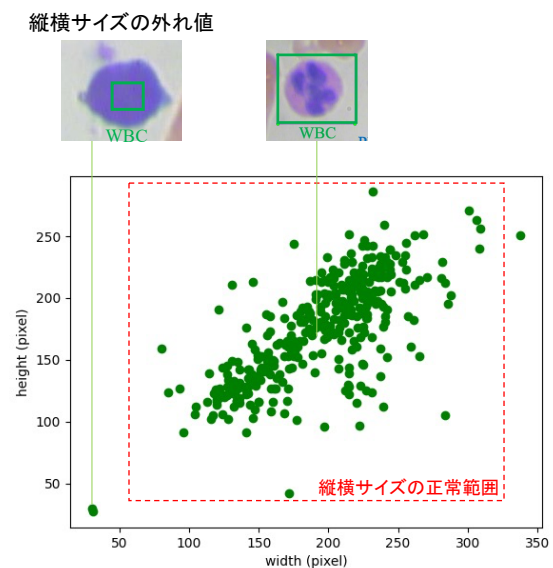


図 3 ラベルの縦横サイズにおける外れ値検出

における外れ値を検出し、該当するラベルを修正候補のリストに追加する。外れ値の検出には、(1)と同様に、統計手法を用いる。

図 4 に、WBC ラベルの縦横サイズの分布と縦横比の外れ値検出により外れ値と判定されたラベルの例を示す。この例では、極端に縦幅に比べて横幅の大きいラベルが分布の中心から離れており、有意点の値から外れていることが統計的検定で認められれば、修正候補のリストに追加されることになる。

(3) ラベルの色情報分布における外れ値検出

画像ファイルから取得したラベル内の色情報分布における外れ値を検出し、該当するラベルを修正候補のリストに追加する。外れ値の検出には、ヒストグラムの類似度であるヒストグラムインターセクション[8]を使用し、類似度の閾値を下回るラベルを外れ値とする。閾値のようなパラメーターはユーザーが任意に設定することができる。

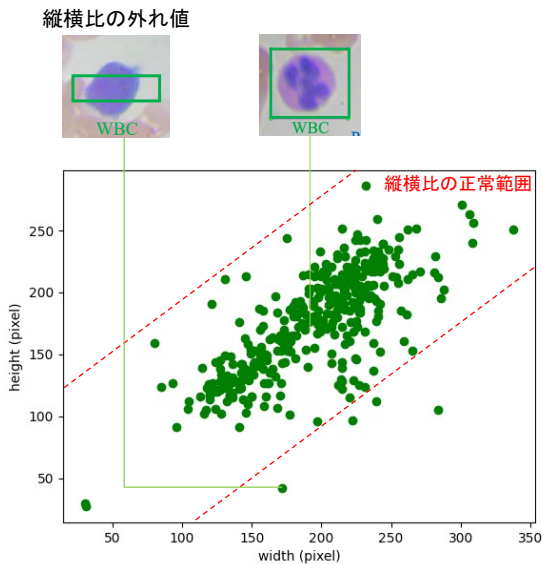


図 4 ラベルの縦横比における外れ値検出

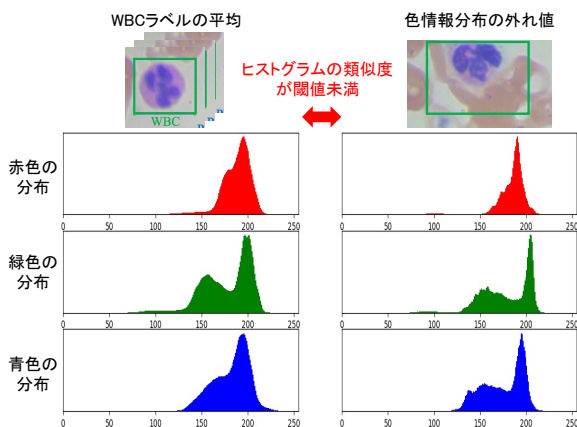


図 5 ラベルの色情報分布における外れ値検出

ヒストグラムインターセクションは 2 つのベクトルを比較し、次元ごとに最小値を求めることで類似度を測る手法である。RGB ヒストグラム同士の類似度は、RGB それぞれ 8 ビットを 2 ビットに減色し、64 次元のベクトルとした上で、以下の計算式により算出する。

$$D = \sum_{i=0}^{63} \min(a_i, b_i)$$

D: ヒストグラム・インターセクション (類似度)

a_i : 全ラベルのカラー値 i のピクセルの個数の平均

b_i : 比較対象ラベルのカラー値 i のピクセルの個数

図 5 に、WBC ラベルの RGB それぞれの色情報分布と外れ値と判定されたラベルの例を示す。この例では、ラベル範囲に対象物と背景以外の物体が含まれており、ラベル全体の平均の色情報分布との類似度が閾値を下回れば、修正候補のリストに追加されることになる。

3.2. ラベル修正案の提示

前節で示した外れ値検出により特定された修正候補ラベルに対して、ラベル範囲の修正案を生成し、ユーザーに提示する。修正案は、ラベル全体における平均値のような統計情報と、修正案に対するユーザーからのフィードバック情報に基づいて生成される。フィードバック情報として、提示された修正案をもとにユーザーがラベルを実際に修正したかどうかを使用する。また、ユーザーのフィードバック情報をもとに、検定の有意水準のような、外れ値検出に使うパラメーターを更新する。この理由は、学習データセットのラベル範囲は対象物の性質やモデルで検出したい範囲など様々な要件から決まることから、外れ値を決める最適なパラメーターがデータセットごとにそれぞれ異なるためである。

図 6 に、外れ値のラベルに対して、ラベル範囲の修正案をユーザーに提示する処理の流れを示す。図 6 は、ラベルの縦横比の分布における有意水準を下回ったものを外れ値として検出し、ラベル全体の平均値をラベル修正案として提示する例である。各ステップにおける処理内容は以下の通りである。

Step1: ラベルの縦横比が有意水準を下回ったラベルが外れ値として検出される。

Step2: 外れ値として検出されたラベルに対して、縦横比を全体の平均値に近づけるようなラベル修正案が、ユーザーへ提示される。ラベル範囲をどの程度平均値に近づけるかは、ユーザーからフィードバックされる重みスコアによって調整される。重みスコアは、0 から 1 の範囲の値をとるパラメーターであり、0 に近いとき

は元のラベルから変更量が少なくなるように、1に近いときは平均値に近づけ変更量が大きくなるようにラベル修正案を提示する。重みスコアが1のときは、平均値をラベル修正案として提示する。

Step3: ユーザは提示された修正案をもとに、ラベル範囲の修正をするか選択する。ユーザーが修正案をそのまま採用してラベル範囲を修正した場合、重みスコアを増加させる (Step3a)。ユーザーが提示された修正案とは異なる形でラベル範囲を修正した場合、重みスコアを減少させる (Step3b)。ユーザーがラベル修正せず、元のラベルをそのまま使用することにした場合、外れ値検出の結果が妥当ではなかったとみなし、外れ値検出に使用する有意水準を正常値の範囲が広がる方向に更新する (Step3c)。ユーザーがラベルを削除した場合は、重みスコアや外れ値検出に使用するパラメータは変更しない。

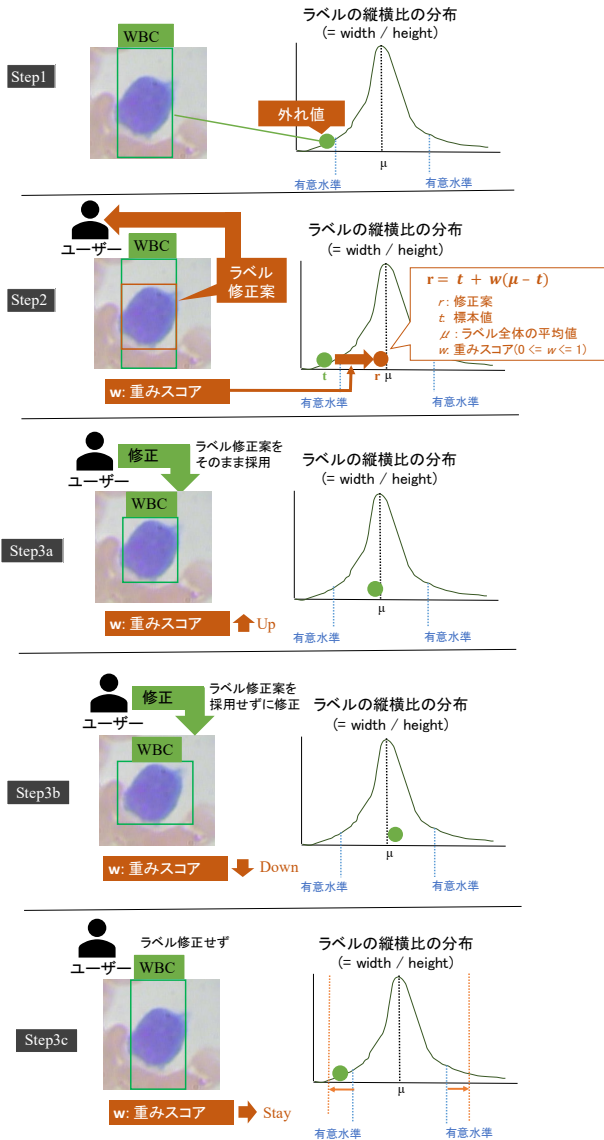


図6 外れ値のラベルに対する修正案の提示

4. 実証実験と考察

4.1. データセット

実証実験として、2.1 節で紹介した 3 クラスのラベルを含む BCCD Dataset を使用して、提案手法の効果を検証した。このデータセットは、ラベル対象物の周辺様相に様々なバリエーションがあり、適切なラベル境界を決めるのが難しい点において、実業務課題に相当するものであると考えられる。このデータセットを用いて評価を行うために、学習用とテスト用で 8:2 の割合となるようランダムに分割し、学習用により作成したモデルをテスト用で精度評価を行った。表 1 に、学習用とテスト用のクラスごとのラベル数の内訳を示す。

4.2. 評価指標

モデルの精度評価にあたり、物体検出で使用される mAP (mean Average Precision) を精度指標として採用した。mAP は、モデルが予測したものの中で正しく予測できたものの比率である Precision、及び正解のラベルに対して正しく予測できたものの比率である Recall の 2 つを合わせた精度指標である。Precision と Recall を単一指標で評価するため、確信度閾値を変えて Recall を変化させたときの Precision の平均である AP (Average Precision) を計算し、クラスごとの AP の平均をとることにより、mAP を求める。

各精度指標の計算式は以下の通りである。

$$Recall = \frac{TP}{TP + FN}$$

$$Precision = \frac{TP}{TP + FP}$$

TP: 正しく検出できた物体の数

FN: 未検出の物体の数

FP: 誤検出した物体の数

$$AP = \int_0^1 P(r) dr$$

$$mAP = \frac{1}{C} \sum_{i=1}^C AP_i$$

$P(r)$: Recall が r のときの Precision の値

AP_i : クラス i の AP の値

C : クラス数

4.3. 実験方法

BCCD Dataset は、深層学習の知見を持つ作業者によって作成されており、配布されているラベルは正確

で、品質の高いものである。そのため、実業務環境で見られるような品質の悪いデータが混ざった状況をシミュレートするため、学習データ全体の 5% のラベルに対して、ランダムノイズを加えた。具体的には、乱数により各ラベルの変更量を決定し、ラベルの位置とサイズが元のラベルから変わるように、ラベルの頂点座標に変更を加えた。

この実業務環境を想定したノイズを含むデータセットから作成したモデルと提案手法を適用して学習データの品質を改善した上で作成したモデルの精度を比較し、提案手法の効果を検証した。また、ラベルのノイズが精度にどの程度影響を及ぼしたかを確認するため、元のデータセットで作成したモデルの精度もベースラインとして比較対象に加えた。モデルは物体検出でよく使われる Faster-RCNN[9] をベースモデルとして採用し、汎用的なデータセットで事前学習されたモデルに対して、学習用データセットを使ってファインチューニングを実施した。学習時のハイパーパラメータは、ベースラインモデルにおける最適値に設定し、各モデルの作成において同じパラメータを使用した。

4.4. 実験結果と考察

表 2 に、元のデータセット (Baseline)、全体の 5% のラベルにノイズを加えたデータセット (Noise5%)、それに対して提案手法を適用したデータセット (提案手法)、それぞれから作成したモデルの精度結果を示す。Baseline と Noise5% を比較すると、mAP が 80.3% から 71.4% へ 8.9% 低下しており、ラベルのノイズが精度に大きな影響を及ぼすことが裏付けられている。クラスごとの AP で見ると、WBC は 99.9% に対して 99.8% とほとんど影響が見られなかったが、RBC は 84.7% に対して 82.9% と 1.8% 低下し、Platelets では 56.2% から 31.6% へと 24.6% 低下と顕著な差が見られた。次に、Noise5% と提案手法を比較すると、mAP が 71.4% から 75.6% へ 4.2% 向上しており、提案手法によるラベルの品質改善の効果が確認できた。クラスごとの AP で見ると、ノイズの影響を受けた RBC は 82.9% から 83.7% へと 0.8% 向上、特に影響が大きかった Platelets では 31.6% から 43.1% へと 11.5% の精度向上が見られた。

図 7 では、クラスごとの詳細な精度傾向を確認するため、縦軸に Precision、横軸に Recall をプロットした曲線である PR 曲線をクラスごとに示し、3 種類のモデルで比較した。PR 曲線をクラス間で比較した場合、クラスごとに精度傾向が大きく異なっており、それぞれのクラスの検出難易度やラベルの品質が精度へ及ぼす影響に差が大きいことがわかる。PR 曲線は右上に膨らむ形、つまり Recall も Precision も同時に高い値を示すものほど予測ミスが少ないといえるが、グラフからは Platelets, RBC, WBC の順に予測ミスが多く、検出

表 1 データセットのクラスごとのラベル数

	学習用	テスト用
RBC	3309	852
WBC	304	71
Platelets	285	76
合計	3898	999

表 2 モデルの精度結果 - mAP

	AP % (RBC)	AP % (WBC)	AP % (Platelets)	mAP %
Baseline	84.7	99.9	56.2	80.3
Noise5%	82.9	99.8	31.6	71.4
提案手法	83.7	100	43.1	75.6

難易度も高いものと考えられる。WBC は Recall が 100% のときに Precision も 100% に近く、ほとんど予測ミスのないクラスである一方、Platelets は確信度をどれだけ下げても Recall が 100% に届かず、検出が難しいクラスであることがわかる。次に、PR 曲線をモデル間で比較した場合、Platelets, RBC, WBC の順にモデル間の差が大きく、学習データの品質の影響を受けやすいと考えられる。特に、Platelets は PR 曲線がモデルによって大きく異なり、Recall の上限値がラベルの品質が悪いほど低下するため、品質低下が検出漏れに直結する結果となっている。

これらの結果から、WBC のようなサイズが大きく、様相の色変化も多い物体はラベルの品質が精度に及ぼす影響は少ない傾向にあると考えられる。それに対して、Platelets のようなサイズが小さく、様相の色変化の少ない物体はラベル品質が精度に及ぼす影響が大きく、提案手法が特に有効に働く対象であると考えられる。また、作業効率の観点では、提案手法を適用することにより、大幅にラベル修正の時間と手間を削減できた。具体的には、提案手法では修正候補が 46 個まで自動的に絞り込まれ、さらにそれらに対する修正案が提示されたため、従来の全 3898 個のラベルを精査する時間と手間が大幅に抑制された。仮に、従来の全ラベルを精査する方法と提案手法で同じ数のラベルを修正したと仮定すると、従来の方法では、ラベル 1 個当たりの確認に要する時間から概算すると、12 時間 20 分程度かかる見込みとなる。それに対して、提案手法では、修正候補の特定から修正までに要した全体の時間は実測で 13 分であった。従って、この実証実験においては、ラベル修正にかかる時間を 50 分の 1 以下に短縮した上で、精度向上を実現できた結果となり、より大規模なデータセットに対してはさらに大きな作業時間の削減が可能となると考えられる。

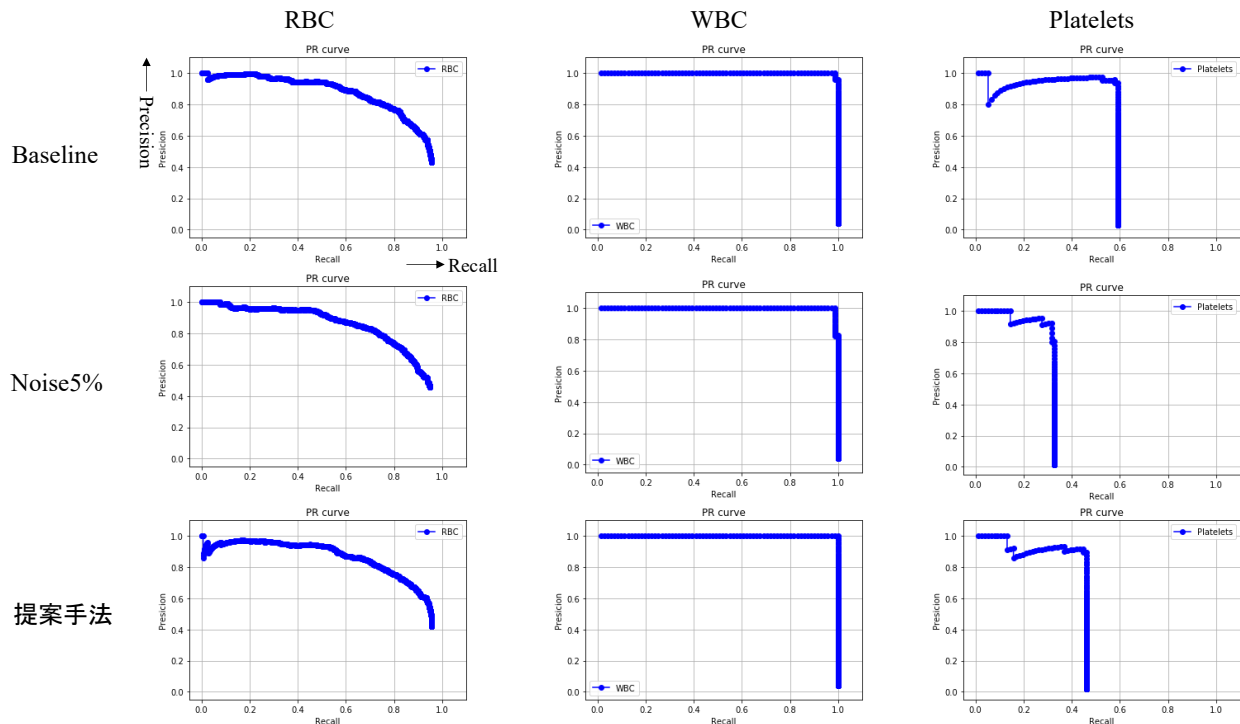


図7 モデルの精度結果 - PR 曲線

5. おわりに

本論文では、深層学習による画像認識の実用化にあたって大きな課題となるモデルの精度向上について、学習データの品質の課題に着目し、効率的に品質改善する手法を提案した。物体検出において、モデル精度に悪影響を及ぼすラベルの特定と修正を効率的に行う手法を確立し、実証実験によりその効果を確認した。効果測定にあたっては、実業務課題を想定して、血液の顕微鏡画像のデータセットを使用し、手動のラベリング作業で発生し得るラベルのノイズをシミュレートした上で、提案手法を適用した結果、mAPで4.2%の精度向上が示された。

なお、提案手法は、特許として出願した上で[10]、すでに様々な実証実験や実用化プロジェクトでの活用を進めている。この手法は物体検出タスクであれば、どのようなデータセットでも汎用的に適用可能なものであり、今後さらに広く利用されることが期待される。深層学習による画像認識は、スマートファクトリーやスマートホスピタルなどにおいて要となるソリューションであり、構想の実現にあたっては様々な知見や技術が必要となるが、提案手法もそのようなDX推進の一助になれば幸いである。

参考文献

- [1] 独立行政法人情報処理推進機構, "DX 白書 2021 第4部 DXを支える手法と技術" p.292, 2021

- [2] BCCD Dataset (under MIT license), https://github.com/Shenggan/BCCD_Dataset
- [3] LabelImg, <https://github.com/heartexlabs/labelimg>
- [4] Dingwen Zhang, Junwei Han, Gong Cheng, Ming-Hsuan Yang, "Weakly Supervised Object Localization and Detection: A Survey", IEEE Transactions on Pattern Analysis and Machine Intelligence, 2021
- [5] Gökem Algan, Ilkay Ulusoy, "Image Classification with Deep Learning in the Presence of Noisy Labels: A Survey", arXiv: 1912.05170, 2019
- [6] IBM Maximo Visual Inspection documentation, <https://www.ibm.com/docs/en/maximo-vi>
- [7] Frank E. Grubbs, "Sample Criteria for Testing Outlying Observations", Annals of Mathematical Statistics Vol. 21 No. 1 pp. 27-58, 1950
- [8] Annalisa Barla, Francesca Odone, Alessandro Verri, "Histogram intersection kernel for image classification", Proceedings 2003 international conference on image processing (Cat. No. 03CH37429). Vol. 3. IEEE, 2003.
- [9] Shaoqing Ren, Kaiming He, Ross Girshick, Jian Sun, "Faster R-CNN: Towards Real-Time Object Detection with Region Proposal Networks", Advances in Neural Information Processing Systems 28, 2015.
- [10] "IMAGE RECOGNITION TRAINING DATA QUALITY ENHANCEMENT", US Patent Application No. 17/539330, 2021

本論文の著作権は、日本アイ・ビー・エム株式会社 (IBM Corporation を含み、以下、IBM といいます。) に帰属します。

ワークショップ、セッション、および資料は、IBM またはセッション発表者によって準備され、それぞれ独自の見解を反映したものです。それらは情報提供の目的のみで提供されており、いかなる参加者に対しても法律的またはその他の指導や助言を意図したものではなく、またそのような結果を生むものでもありません。本論文に含まれている情報については、完全性と正確性を期するよう努力しましたが、「現状のまま」提供され、明示または暗示にかかわらずいかなる保証も伴わないものとします。本論文またはその他の資料の使用によって、あるいはその他の関連によって、いかなる損害が生じた場合も、IBM またはセッション発表者は責任を負わないものとします。本論文に含まれている内容は、IBM またはそのサプライヤーやライセンス交付者からいかなる保証または表明を引き出すことを意図したものでも、IBM ソフトウェアの使用を規定する適用ライセンス契約の条項を変更することを意図したものでもなく、またそのような結果を生むものでもありません。

本論文で IBM 製品、プログラム、またはサービスに言及していても、IBM が営業活動を行っているすべての国でそれらが使用可能であることを暗示するものではありません。本論文で言及している製品リリース日付や製品機能は、市場機会またはその他の要因に基づいて IBM 独自の決定権をもっていつでも変更できるものとし、いかなる方法においても将来の製品または機能が使用可能になると確約することを意図したものではありません。本論文に含まれている内容は、参加者が開始する活動によって特定の販売、売上高の向上、またはその他の結果が生じると述べる、または暗示することを意図したものでも、またそのような結果を生むものでもありません。パフォーマンスは、管理された環境において標準的な IBM ベンチマークを使用した測定と予測に基づいています。ユーザーが経験する実際のスループットやパフォーマンスは、ユーザーのジョブ・ストリームにおけるマルチプログラミングの量、入出力構成、ストレージ構成、および処理されるワークロードなどの考慮事項を含む、数多くの要因に応じて変化します。したがって、個々のユーザーがここで述べられているものと同様の結果を得られると確約するものではありません。

記述されているすべてのお客様事例は、それらのお客様がどのように IBM 製品を使用したか、またそれらのお客様が達成した結果の実例として示されたものです。実際の環境コストおよびパフォーマンス特性は、お客様ごとに異なる場合があります。

IBM、IBM ロゴは、米国やその他の国における International Business Machines Corporation の商標または登録商標です。他の製品名およびサービス名等は、それぞれ IBM または各社の商標である場合があります。現時点での IBM の商標リストについては、ibm.com/trademark をご覧ください。