

Small-Scale Object Detection Label Reassignment Strategy

Jung-In An*, Yoon Kim*, Hyun-Soo Choi**

*Student, Dept. of Computer Science and Engineering, Kangwon National University, Chuncheon, Korea

*Professor, Dept. of Computer Science and Engineering, Kangwon National University, Chuncheon, Korea

**Professor, Dept. of Computer Science and Engineering, Kangwon National University, Chuncheon, Korea

**Professor, Dept. of Computer Science and Engineering, Seoul National University of Science and Technology, Seoul, Korea

[Abstract]

In this paper, we propose a Label Reassignment Strategy to improve the performance of an object detection algorithm. Our approach involves two stages: an inference stage and an assignment stage. In the inference stage, we perform multi-scale inference with predefined scale sizes on a trained model and re-infer masked images to obtain robust classification results. In the assignment stage, we calculate the IoU between bounding boxes to remove duplicates. We also check box and class occurrence between the detection result and annotation label to re-assign the dominant class type. We trained the YOLOX-L model with the re-annotated dataset to validate our strategy. The model achieved a 3.9% improvement in mAP and 3x better performance on AP_S compared to the model trained with the original dataset. Our results demonstrate that the proposed Label Reassignment Strategy can effectively improve the performance of an object detection model.

▶ **Key words:** Artificial Intelligence, Object Detection, Dataset Refinement, Dataset Reassignment

[요 약]

본 논문은 객체 위치식별 알고리즘의 성능을 향상하기 위한 레이블 재할당 방법을 제안한다. 제안한 방법은 추론 단계와 재할당 단계로 구분한다. 추론 단계에서는 학습된 모델로부터 사전 지정된 크기에 따라 다중 스케일 추론을 수행한 뒤, 이를 마스킹한 영상을 다시 한번 추론하여 강인한 클래스 종류의 추론 결과를 얻는다. 재할당 단계에서는 박스간의 IoU를 계산하여 중복 박스를 제거하고, 박스와 클래스의 빈도를 계산하여 지배적 클래스를 다시 할당하였다. 제안한 방법을 검증하기 위하여 공사현장 안전장비 인식 영상 데이터 세트에 레이블 재할당 방법을 적용하고 이를 YOLOX-L 객체 탐지 모델에서 학습하였다. 실험 결과 적용 전 대비 mAP가 3.9% 향상하여 51.07%를 달성하였으며 AP_S를 3배 이상 향상하여 14.53%를 달성하였다. 실험 결과를 통해 레이블 재할당 알고리즘이 더 우수한 성능의 모델을 훈련해 낼 수 확인하였다.

▶ **주제어:** 인공지능, 객체 탐지, 데이터셋 정제, 데이터셋 재할당

-
- First Author: Jung-In An, Corresponding Author: Yoon Kim, Hyun-Soo Choi
 - *Jung-In An (jungin500@kangwon.ac.kr), Dept. of Computer Science and Engineering, Kangwon National University
 - *Yoon Kim (yooni@kangwon.ac.kr), Dept. of Computer Science and Engineering, Kangwon National University
 - **Hyun-Soo Choi (choi.hyunsoo.1112@gmail.com), Dept. of Computer Science and Engineering, Kangwon National University and Dept. of Computer Science and Engineering, Seoul National University of Science and Technology
 - Received: 2022. 11. 28, Revised: 2022. 12. 23, Accepted: 2022. 12. 23.

I. Introduction

고용노동부가 발표한 2020 산업 재해 현황분석[8]의 연도별 산업 재해 지표 추이에 따르면, 2011~2020년간 요양 재해자가 지속해서 발생하는 양상을 나타낸다. Fig. 1은 이러한 지표 추이를 나타내고 있다.

Year	2011	2012	2013	2014	2015	2016	2017	2018	2019	2020
Establishments	1,738,196 (100)	1,825,296 (105)	1,977,057 (114)	2,187,391 (136)	2,367,186 (136)	2,457,225 (141)	2,507,364 (144)	2,654,107 (153)	2,680,874 (154)	2,719,308 (155)
Workers	14,362,372 (100)	15,548,423 (108)	15,449,228 (108)	17,062,308 (119)	17,968,931 (125)	18,431,716 (128)	18,560,142 (129)	19,073,438 (133)	18,725,160 (133)	18,974,513 (132)
Injuries	93,292 (100)	92,256 (99)	91,824 (98)	90,909 (97)	90,129 (97)	90,656 (97)	89,848 (96)	102,305 (110)	109,242 (117)	108,379 (116)

Fig. 1. Trend Of Occupational Injuries By Year

이러한 배경 아래 산업 재해 피해 감소를 위한 여러 장치가 등장하고 있으며 최근 인공지능 기술을 활용하여 산업 재해를 사전에 방지하고자 하는 노력이 계속되고 있다. 인공지능을 기반으로 하는 재해 발생 탐지 시스템은 위험 객체 탐지 또는 위험 상황 판별을 통해 산업작업자의 위험 요소를 사전에 판단하고 예방함으로써 추가적인 재해 피해를 방지할 수 있다. 안전사고를 효과적으로 예방하는 재난 알림 시스템은 높은 재현율과 낮은 지연시간이 요구되며, 이에 부합하기 위하여 수많은 데이터를 이용하여 학습한 인공지능 모델이 요구된다. 일반적인 인공지능 기반의 재난 알림 시스템은 작업장의 위험 상황 알림 서비스를 제공하여 안전 장구 미착용과 쓰러짐을 알려 작업자의 안전 사고를 예방하고 빠르게 대처하는 데 도움을 준다.

최근 재난 알림 시스템의 성능 개선을 위하여 새로운 인공지능 모델이 지속해서 연구되고 있다. 이러한 모델의 성능을 정량적으로 비교하기 위하여 오픈 데이터 세트를 주로 사용하고 있으며, 객체 탐지 모델의 성능 측정을 위한 데이터 세트로 PASCAL VOC[1], MS COCO[2]가 널리 사용된다.

객체 탐지 문제는 영상 분류 문제에 비하여 인공지능 모델과 데이터의 복잡도가 높다. 이에 따라 데이터 세트의 레이블을 만드는 과정에서 잘못된 레이블을 생성하는 경우 성능이 크게 하락한다. 데이터 세트를 생성하기 위하여 사람이 직접 정지 영상에서 해당 영상에 존재하는 객체의 위치와 해당하는 객체의 클래스를 레이블로 기록하기 때문에, 레이블 기록자의 기록 실수나 전처리 과정의 실수로 인하여 잘못된 레이블을 포함할 가능성이 크다. 따라서 레이블의 검증이 필요하며 이를 자동화하는 방안 관련 연구가 대두되고 있다.

데이터의 양이 늘어날수록 데이터 레이블의 일관성을 검증하는 데에 필요한 소요는 더 커진다. 본 논문에서는 이를 해결하고자 일부 객체 레이블이 잘못 기재된 데이터 세트 레이블을 정제하는 레이블 재할당 기법을 소개하고, 방법을 한국지능정보사회진흥원 AIHub에서 제공하는 공사 현장 안전 장비 인식 데이터 공개 데이터셋[7]과 YOLOX[5] 모델에 적용하여 결과를 입증하였다. 본 기법은 영상에 실제로 존재하는 객체이나 레이블에 다른 클래스로 기재된 이상 클래스를 정상 클래스로 다시 할당하고, 기존 레이블에서 놓친 작은 크기의 객체들에 대한 새로운 박스 레이블을 생성한다.

본 논문의 구성으로 2장에서는 이전 연구들에 관하여 소개하고, 3장에서는 제안하는 데이터 세트 재할당 기법을 소개한다. 4장에서는 해당 기법을 활용한 실험 과정을 보이며 5장에서는 그 결과를 제시하고 마지막으로 6장에서 결론을 맺는다.

II. Related Works

인공지능 기반의 재난 알림 탐지 시스템에 관한 연구가 활발히 진행되고 있다. ITLM[9]과 INCV[10]에서는 YOLOv5 모델과 공개 데이터 세트를 활용하여 개인 안전 장구의 착용 여부를 탐지하고 이를 화면에 표시하는 시스템을 구현하였다. [9]에서는 실시간 카메라 영상을 이용하여 구현한 시스템의 동작을 검증하였으며, [10]에서는 Jetson Nano 임베디드 시스템을 활용하여 실제 공장 현장에 배치하기 쉽도록 저전력의 관제 시스템을 구성하였다.

레이블 교정에 관한 연구는 방대한 데이터 세트에 내재한 이상 레이블을 선별하고 해당 레이블을 정상 레이블로 재할당하는 방법을 고안한다. [15], [16]에서는 데이터 세트 내 잘못된 레이블을 분류하기 위하여 단계별 학습을 적용하였다. 현재 단계에서는 일정 부분의 데이터를 선택하여 모델을 학습하고, 학습된 모델을 이용하여 선택되지 않은 데이터와 레이블 간의 손실(Loss)을 구한다. 이때, 손실이 작은 레이블들만 보존하여 다음 단계의 데이터로 사용하며 이 과정을 반복하여 정상 데이터 세트를 걸러내었다. [3]에서는 컴퓨터비전, 자연어처리 등 다양한 분야에 일반적으로 사용되는 10개 데이터 세트의 정답 레이블의 이상을 연구하였다. 해당 연구는 Confident Learning[4] 알고리즘을 이용하여 의심되는 정답 레이블을 추출하고, 이를 Amazon Mechanical Turk 크라우드소싱을 활용하여 검증하고 해당 레이블을 재할당하였다. 해당 방법을

ImageNet 데이터 세트에 적용한 결과로 Validation Set의 5.8%에 해당하는 2,916장의 레이블을 교정하였고, 동일한 VGG-19 모델에서 이전 Validation Set보다 더 높은 정확도를 달성하였다. 해당 논문은 검증 데이터 세트 레이블의 낮은 품질이 잘못된 모델 평가 결과를 낼 수 있음을 시사하였으나, 논문에서 제안한 방법을 객체 탐지 데이터 세트에 적용하기 어렵다는 한계점을 가지고 있다. 본 논문에서는 이를 해결하여 학습된 모델을 활용하여 객체 탐지 데이터 세트의 레이블을 교정하는 방안을 제안한다.

III. The Proposed Method

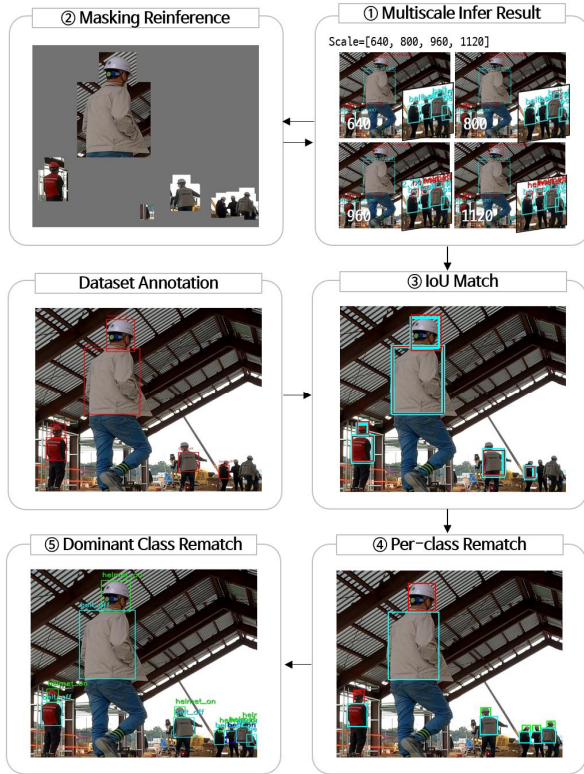


Fig. 2. Proposed Label Reassignment Strategy

본 논문에서 제안하는 방법은 모델로부터 다중 입력 크기로 객체를 추론하고 이를 이용하여 객체 재할당을 적용하는 과정으로 진행하며 Fig. 2에서 그 과정을 나타낸다.

3.1 Masking Re-Inference

제안 방법에서는 가장 먼저 데이터 세트로 학습된 모델에 학습 영상의 객체 추론을 진행한다. 동일한 입력 영상을 사전에 지정된 여러 입력 크기로 크기 조절하여 추론한

다. 이후 입력 영상에서 배경 영역을 제거하고 해당 영상을 다시 한번 모델로 추론한다. 본 논문에서 추론 결과의 모든 박스 영역을 제외한 나머지 영역을 배경 영역으로 가정하였다. 배경을 제거함으로써 이상 클래스로 인한 모델의 Bias를 완화하여 이상 클래스를 정상 클래스로 탐지하는 효과를 얻을 수 있다. 해당 과정에서 얻은 결과 박스는 원본 영상의 추론 결과 박스와 IoU(Intersection over Union)를 계산하고 동일한 위치에 해당하는 박스의 클래스 종류를 다시 할당한다.

3.2 Label Reassignment Strategy

레이블 재할당 방법은 마스킹 재추론, IoU 매칭, 클래스별 재매칭, 지배적 클래스 할당 과정으로 구성한다. 각 과정은 아래 세 가지의 결과 집합을 생성한다.

- (1) $B_{preserve}$: 매칭이 완료되어 보존할 박스 및 클래스 종류 목록 집합
- (2) $B_{inferonly}$: 추론된 결과에만 존재하는 박스 집합
- (3) $B_{labelonly}$: 레이블에만 존재하는 박스 집합

모든 재할당 과정을 거친 후, $B_{preserve}$ 집합의 박스들을 재할당 결과 데이터 세트 레이블로 저장한다. $B_{inferonly}$ 은 오탐으로 판단하여 사용하지 않고, $B_{labelonly}$ 는 잘못된 레이블로 판단하여 모두 사용하지 않는다. 각 과정에서는 두 박스의 위치 비교를 위하여 IoU를 계산하고 그 값이 임계값을 초과하는지 확인한다. 두 박스의 IoU가 임계값을 초과한 경우 같은 박스로 판단한다.

3.2.1 IoU Match

모든 입력 영상 크기에 대하여 원본 데이터 세트 레이블 박스와 마스킹 재추론 결과 박스간의 IoU를 계산하고 매칭된 박스와 추론된 클래스 종류들을 $B_{preserve}$ 에 저장한다. 이는 기존 데이터 세트 레이블에 존재하는 정상 객체를 보존하고 이상 객체를 여과하는 효과가 있다. 데이터 세트 레이블에 박스가 존재하나 추론되지 않은 객체들을 $B_{labelonly}$ 에 저장하고, 반대로 추론된 박스가 존재하나 데이터 세트 레이블에 매칭되는 박스가 없는 객체를 $B_{inferonly}$ 에 분류하여 저장한다.

3.2.2 Per-class Rematch

이 단계에서는 클래스별로 $B_{inferonly}$ 에 존재하는 모든 박스에 대하여 IoU를 계산하고 같은 위치에 존재하는 박

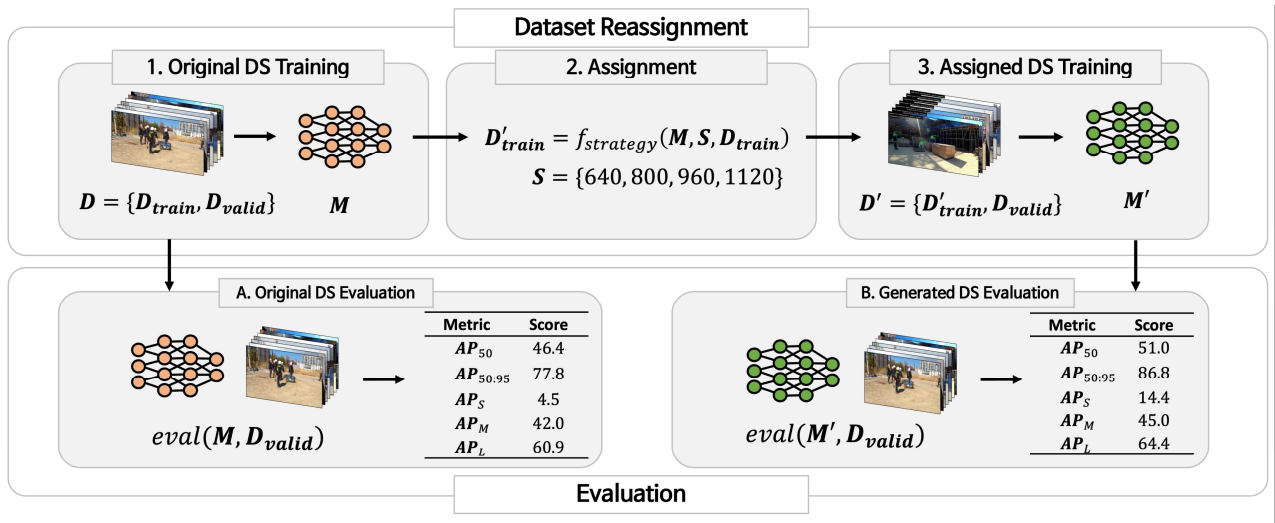


Fig. 3. Proposed Reassignment And Evaluation Pipeline

스들과 빈도를 찾는다. 해당 과정에서 동일 위치에 존재하는 박스들은 발생 빈도가 2 이상이면 해당 박스들 중 대표 박스 하나와 클래스 종류들을 $B_{preserve}$ 에 보존한다. 이 과정은 데이터 세트 레이블에 없는 새로운 박스를 보존하고 발생 빈도에 따라 오답인 박스를 제거한다.

3.2.3 Dominant Class Rematch

클래스별 재매칭 단계에서는 클래스가 다르지만 동일한 위치에 존재하는 박스들을 매칭하지 않기 때문에 지배적 클래스 재매칭 단계에서는 박스로만 매칭을 수행한다. 본 단계에서는 내부 비교와 외부 비교를 진행하고 비교 결과에 따라 지배적인 클래스 할당을 진행한다.

내부 비교는 이미 매칭된 결과인 $B_{preserve}$ 내에서 모든 박스들이 나머지 박스들과 서로를 비교한다. 비교 과정에서 두 박스가 동일한 위치에 존재하는 경우, 각 박스의 클래스 종류들을 합하여 가장 빈도가 높은 클래스 종류를 해당 박스의 클래스로 재할당한 뒤 최다 빈도 클래스 종류가 없는 박스를 $B_{preserve}$ 에서 제거한다. 위 과정은 두 박스에 대하여 지배적인 클래스 할당을 진행함으로써 데이터 세트 레이블에서 비롯된 이상 클래스를 다시 할당한다.

외부 비교는 $B_{preserve}$ 와 $B_{labelonly}$ 에 내부 비교와 동일한 과정을 적용한다. 재할당 단계를 완료한 레이블은 COCO나 VOC 포맷으로 저장하여 추후 모델 학습에 사용한다.

IV. Experiments

실험 과정은 Fig. 3에서 나타낸다. 우선으로 원본 데이터 세트로 모델을 학습하고 학습 결과를 평가하여 기본 결과를 생성해 낸다. 다음으로, 데이터 세트 레이블에 재할당 방법을 적용한 후에 나온 결과 레이블을 데이터 세트로 모델에 학습하고 원본 데이터 세트로 학습한 모델과 비교한다.

4.1 Dataset

실험에는 한국지능정보사회진흥원 AIHub에서 제공한 공사현장 안전장비 인식 영상 공개 데이터 세트를 활용하였다. 전체 데이터 세트 중 헬멧과 안전띠 두 가지 클래스를 선별하였고, 클래스별 착용/미착용 상태를 구분하여 총 4개의 클래스로 구성하였다. 실험 데이터 세트는 학습 데이터 823,021장, 검증 데이터 7,774장 규모의 데이터 세트로 구성된다.

AIHub의 공사현장 안전장비 인식 영상 공개 데이터 세트의 경우 박스가 잘못된 위치에 존재하거나 작은 객체에 대한 박스가 존재하지 않는 경우가 다수 존재한다. 본 논문에서는 검증된 레이블을 확보하기 위하여 검증 데이터 세트 재확인 작업을 진행하였다. 재확인 작업에서는 실제 객체에 대한 레이블이 존재하는지와 실제 객체와 레이블에 기록된 클래스 종류가 같은지에 대하여 확인하며, 비정상 레이블을 찾아내고 수정한다. 재확인 작업을 위하여 검증 세트 내 레이블을 1차, 2차로 검증을 진행하였고 그 결과, helmet_off 객체 365개에 대하여 검증하여 최종적으로 1,056장의 테스트 데이터 세트를 확보하였다.

4.2 Model

본 논문에서는 실시간 객체 탐지를 위하여 YOLOX 모델을 사용하였다. YOLOX는 YOLO 계열[13][6][14]에서 파생된 실시간 객체 탐지 모델이다. 이는 앵커 프리 구조로 객체 탐지에 앵커 박스를 사용하지 않고 Head 레이어에서 객체의 위치를 직접 추론한다. 이러한 특성으로 YOLOX는 다른 모델 대비 비교적 간단하게 객체 탐지 파이프라인의 구성이 가능하다. YOLOX는 Decoupled Head 구조를 적용하여 객체 위치 회귀와 분류에 각각 다른 Head 레이어를 사용함으로써 모델이 데이터 세트에 더 빠르게 수렴하는 특성을 가진다.

Table 1. Applied Hyperparameters

Training Parameter	Value
Pretrained model	COCO
Input size	640
Multiscale train range	640~1120
Augmentation	Mosaic
Learning rate	0.001
Learning rate scheduler	yoloxwarmcos
Epoch count	10
Inference Parameter	Value
Inference scale	640~1120
NMS IoU Threshold	0.3
Confidence Threshold	0.23

원본 데이터 세트 학습과 재할당된 데이터 세트 학습 단계에서는 YOLOX-L 모델과 COCO 데이터 세트로 사전 학습된 가중치에 전이 학습을 수행하였으며, 두 학습 과정에 Table 1의 설정값을 동일하게 적용하였다.

모든 학습에는 데이터 증강을 위하여 Mosaic[11] 데이터 증강 기법을 적용하여 학습 데이터 세트의 영상 중 4장을 무작위로 선별하여 한 장의 영상을 만들어 내었다. 이에 더하여, Multi-Scale Training[6]을 적용하여 입력 영상의 크기를 10개 배치 진행에 한 번씩 {640, 800, 960, 1120} 중에서 무작위로 선택하였다.

4.3 Environments

실험에 활용한 학습 서버 하드웨어는 Ubuntu 20.04 Server OS, i9-10890XE CPU, 256GB RAM, RTX 3090 GPU 4개로 구성하였으며, 소프트웨어는 Kubernetes 학습 클러스터 환경과 PyTorch 1.11.0a0 도커 이미지로 구성하였다. 학습과 추론 과정에는 AMP(Automatic Mixed Precision[12])를 적용하여 연산 과정을 가속한다. AMP는 합성곱 연산을 반정밀도 부동소수점(FP16) 형식으로 연산하여 GPU에 내장된 Tensor 코어를 활용하여 연산 과정을 가속하는 기술이다. FP16의 정확도 손실을 최소화하기

위하여 동적 범위가 큰 합성곱 이외의 연산은 일반적으로 사용되는 단정밀도 부동소수점(FP32) 형식으로 연산을 진행하였다.

제안한 방법의 마스킹 재추론 결과 박스와 원본 영상 추론 결과 박스간의 IoU 비교 과정에서 임계값을 0.8로 사용하였으며, 레이블 재할당 과정에서 임계값을 모델 추론에 사용하는 IoU 임계값과 같은 값을 사용하였다.

실험을 위하여 사전에 데이터 세트 레이블을 JSON 형태의 COCO 어노테이션 포맷으로 변환하였다. 또한, 다중 GPU 환경을 활용하기 위하여 PyTorch 분산 병렬처리 라이브러리(DDP)를 이용하였고, 레이블 재할당 과정에서는 전체 데이터 세트 영상을 균등히 나누어 각 GPU에 할당된 뒤 4개의 작업을 병렬로 진행하고 모든 결과를 한 파일로 병합하였다.

V. Results

실험 결과 평가는 원본 데이터 세트 D 의 학습 결과와 재할당된 데이터 세트 D' 의 학습 결과를 평가하고 비교한다. 각 결과의 평가는 데이터 세트 확인 과정에서 재확인된 1,056장의 데이터를 사용하였다. 평가 결과의 입력 크기별 AP에서 사용하는 영상 크기는 각각 아래와 같이 정의한다.

- (1) $AP^{small} : area(img) < 32^2$
- (2) $AP^{medium} : 32^2 \leq area(img) < 96^2$
- (3) $AP^{large} : area(img) \geq 96^2$

재할당 알고리즘을 적용한 데이터로 학습을 진행한 결과 성능은 Table 2와 같다. 제안한 재할당 알고리즘을 적용한 결과 $AP_{50:95}$, AP_{50} 모두 기존 대비 우수한 성능을 보여주고 있다. 또한, AP^{small} 의 경우 4.71%에서 14.53%로 3.1배 성능이 증가하였다. 각 클래스별 AP_{50} 의 경우 helmet_on 클래스에서 약 1.2배, helmet_off 클래스에서 약 1.4배로 제안한 방법이 기존보다 성능의 큰 향상을 확인할 수 있다.

5.1 Confidence Threshold Search

모델로부터 최적의 결과 박스를 추론하기 위하여, 원본 데이터 세트를 학습한 모델 M 의 평가 과정에서 사용할 최적의 객체 신뢰도 임계값 검색을 수행하였다. 임계값 검

Table 2. Model Evaluation Result

Model	$AP_{50:95}$ (%)	AP_{50} (%)	AP^{small} (%)	AP^{medium} (%)	AP^{large} (%)	$AP_{50}(class)$ (%)			
						helmet_on	helmet_off	belt_on	belt_off
Baseline M	47.15	77.59	4.71	42.88	61.34	73.74	58.39	89.11	89.13
Proposed M'	51.07	86.28	14.53	45.05	64.47	88.31	80.81	87.39	88.61

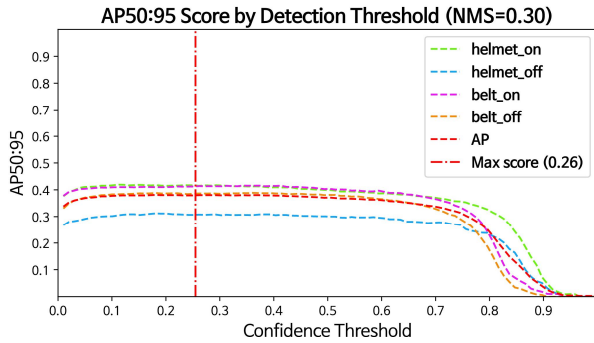


Fig. 4. mAP Score By Threshold Search Space

색에는 Mean Average Precision 점수를 성능 척도로 사용하고, 신뢰도 임계값의 검색 공간으로 0.01부터 0.99까지 0.01 단위의 총 99개 값을 정의하였다. Fig. 4는 그 결과를 나타낸다. 임계값 검색을 수행한 결과, 신뢰도 임계값 0.23이 AP_{50} 가 39.23%, $AP_{50:95}$ 가 71.04%로 가장 높은 성능이 나타나는 것을 확인하였다. 본 논문에서는 실험 결과에 따라 최적 신뢰도 임계값 0.23을 추론 단계에 NMS 신뢰도 임계값으로 적용하였다.

5.2 Inference



[영상1] 640 결과 (7.8x)

[영상1] 1120 결과 (7.8x)



[영상2] 640 결과 (6.8x)

[영상2] 1120 결과 (6.8x)

Fig. 5. Inference Result Of Two Input Sizes

데이터 세트 레이블 재할당을 위하여 학습 과정과 동일한 스케일 {640, 800, 960, 1120} 4가지의 크기에 대하여 각각 입력 영상을 추론하여 마스킹 재추론 방법의 신뢰성을 검증한다. Fig. 5는 두 영상에 대하여 입력 크기 640, 1120에서 추론 결과를 각각 확대하여 나타내고 있다. 실험 결과를 통해 작은 영상에서 미탐으로 처리된 객체들이 영상이 커짐에 따라 검출됨을 확인하였다. 이는 제안한 마스킹 재추론 과정에서 다양한 크기의 영상으로 추론하여 기존에 탐지하지 못한 객체들을 검출함에 따라 제안한 방법이 미탐 객체들에 대한 레이블 생성이 가능하다는 것을 검증하였다.

5.3 Label Reassignment

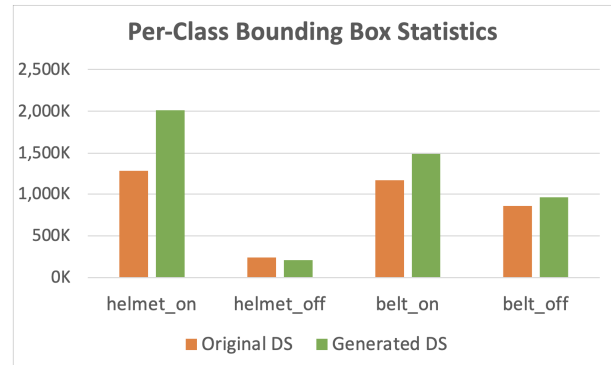


Fig. 6. Statistics Of Bounding Boxes Per Class

재할당 알고리즘을 적용한 결과 데이터 세트 전체의 박스 개수는 원본 데이터 세트 대비 31.4%를 더 할당하여 총 4,674,717개가 할당되었다. Fig. 6에서는 클래스 종류별 박스의 개수 차이를 나타내고 있다. 실험 결과에서 helmet_off 클래스의 기존 대비 개수가 감소하였으나, 이는 이상 레이블로 존재하던 helmet_off 클래스가 정상 레이블인 helmet_on 클래스로 재할당되는 과정에서 수가 감소한 현상이다.

VI. Conclusions

본 논문에서는 데이터 세트의 품질을 정제하는 새로운 방법을 제안하였다. 제안 기법은 원본 데이터 세트 레이블

에 새로운 객체와 클래스를 재할당하는 방법이다. 또한 제안한 방법을 검증하기 위하여 원본 데이터 세트의 모델 학습 결과와 재할당 데이터 세트의 학습 결과를 서로 비교하여 제안 방법을 적용한 결과 성능의 우수함을 확인하였다. 제안 방법은 현재 재난 알림 시스템에 활용하는 오픈 데이터 세트 중 하나인 AIHub의 공사현장 안전장비 인식 영상 공개 데이터 세트의 레이블을 교정함에 따라 해당 시스템의 성능 향상에 기여할 수 있음을 확인하였다.

ACKNOWLEDGEMENT

This work was supported by Basic Science Research Program through the National Research Foundation of Korea (NRF) funded by the Ministry of Education (NRF-2022R1F1A1076454), funded by "Regional Innovation Strategy (RIS)" through the National Research Foundation of Korea (NRF) funded by the Ministry of Education (MOE) (2022RIS-005)

REFERENCES

- [1] M. Everingham, L. Van Gool, C.K.I. Williams, et al., "The PASCAL Visual Object Classes (VOC) Challenge", *International Journal of Computer Vision*, Vol. 88, pp. 303-338, Sep. 2009. DOI: 10.1007/s11263-009-0275-4
- [2] T. Y. Lin, et al., "Microsoft COCO: Common Objects in Context", *Lecture Notes in Computer Science*, Vol. 8693, pp. 740-755, Sep. 2014, DOI: 10.1007/978-3-319-10602-1_48
- [3] C. Northcutt, A. Athalye, J. Mueller, "Pervasive Label Errors in Test Sets Destabilize Machine Learning Benchmarks", *Conference on Neural Information Processing Systems (NeurIPS), Dataset and Benchmark Track 1*, Virtual Conference, Dec. 2021
- [4] C. Northcutt, L. Jiang, I. Chuang, "Confident Learning: Estimating Uncertainty in Dataset Labels", *ACM Journal of Artificial Intelligence Research*, Vol. 70, pp. 1373-1411, May 2021, DOI: 10.1613/jair.1.12125
- [5] Z. Ge, S. Liu, F. Wang, Z. Li, J. Sun, "YOLOX: Exceeding YOLO Series in 2021", *arXiv preprint*, Aug. 2021, DOI: 10.48550/arXiv.2107.08430
- [6] J. Redmon, A. Farhadi, "YOLO9000: Better, Faster, Stronger", *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, pp. 7263-7271, Jul. 2017, DOI: 10.1109/CVPR.2017.690
- [7] Construction Site Safety Equipment Detection Image Dataset, National Information Society Agency, <https://aihub.or.kr/aihubdata/data/view.do?dataSetSn=163>
- [8] Ministry of Employment and Labor, Analysis of the current status of industrial accidents in 2020, https://www.moel.go.kr/info/public/publicDataView.do?bbs_seq=20211201900
- [9] J. Sukhwan, L. Joowon, K. Mingyun, H. Sungeun, B. Junil, K. Hwajong, "A Study on Helmet Wear Identification Using YOLOv5 Object Recognition Model", *The Korean Institute of Communications and Information Sciences (KICS)*, pp. 1293-1294, Pyeongchang, Korea, Feb. 2022
- [10] J. Jaegyung, G. Taehun, K. Gyeongmin, L. Jaemoon, K. Sungyoung, O. Byoung-woo, "Development of Personal Safety Equipment Wearing Confirmation System", *Proceedings of KIIT Conference*, pp. 664-667, Jeju, Korea, Jun. 2022
- [11] A. Bochkovskiy, C. Wang, H. M. Liao, "YOLOv4: Optimal Speed and Accuracy of Object Detection", *arXiv preprint*, arXiv:2004.10934, Apr. 2020
- [12] P. Micikevicius, S. Narang, et al., "Mixed Precision Training", *International Conference on Learning Representations (ICLR)*, Poster Session, Canada, May 2018
- [13] J. Redmon, S. Divvala, R. Girshick, A. Farhadi, "You Only Look Once: Unified, Real-Time Object Detection", *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, pp. 779-788, Jun. 2016
- [14] J. Redmon, A. Farhadi, "YOLOv3: An Incremental Improvement", *arXiv preprint*, Apr. 2018
- [15] Y. Shen, S. Sanghavi, "Learning with bad training data via iterative trimmed loss minimization", in *Proceedings of International Conference on Machine Learning (ICML)*, pp. 5739-5748, Jun. 2019
- [16] P. Chen, B. Liao, G. Chen, S. Zhang, "Understanding and utilizing deep neural networks trained with noisy labels", in *Proceedings of International Conference on Machine Learning (ICML)*, pp. 1062-1070, Jun. 2019

Authors



Jung-In An received B.S. degrees in Computer Science and Engineering from Kangwon National University, Korea, in 2020. He is currently a M.S. Student in the department of Computer Science and

Engineering at Kangwon National University. His research interests are machine learning, model optimization and computer vision.



Yoon Kim received a B.S. degree in 1993, and M.S. degree in 1995, and a Ph.D. degree in 2003, in electronic engineering with the Department of Electronic Engineering from Korea University. In 2004, he joined the

Department of Computer Science and Engineering, Kangwon National University, where he is currently a professor. From 1995 to 1999, he was with the LG-Philips LCD Co., where he was involved in research and development on digital image equipment. His research interests are in the areas of machine learning, multimedia communications, and computer vision.



Hyun-Soo Choi received the B.S. degree in computer and communication engineering for the first major and in brain and cognitive science for the second major from Korea University, in 2013.

In 2020, he also received the integrated M.S./Ph.D. degree in electrical and computer engineering from Seoul National University, South Korea. From 2020 to 2021, he was a senior researcher in Vision AI Labs of SK Telecom. From 2021 to 2021, he was in Kangwon National University as an assistant professor, South Korea. Since Oct 2021, he has joined Ziovision as chief technical officer. Since 2023, he is working in Seoul National University of Science and Technology as an assistant professor.