

오염 물질 판별을 위한 관심영역 기반 히스토그램 유사성 비교 방법

손예진, 최승호*

송실대학교, *광운대학교

whitetong99@gmail.com, *jcn99250@naver.com

Region-of-Interest based Histogram Similarity Comparison for Contaminant Discrimination

Yea-Jin Shon, Seoung-Ho Choi*

Soongsil Univ., *KwangWoon Univ

요약

실생활에서 음식물을 먹다보면 옷에 음식물이 튀어 오염이 발생하는 경우가 발생한다. 오염 물질을 제대로 닦지 않으면 나중에 옷에 얼룩이 지는 문제가 발생한다. 위 문제를 경감하고자 오염 물질 판별 방법을 제안한다. 첫 번째로 오염물질의 Bounding box와 오염 물질의 Class를 추정하는 단계, 두 번째로 예측한 Bounding box의 좌표와, 정답 Bounding box의 좌표에 해당하는 영역을 Crop하는 단계, 세 번째로 Crop한 이미지를 Histogram으로 변환하는 단계, 네 번째로 예측 Histogram과 정답 Histogram의 유사성을 비교한다. 이를 검증하기 위해서 물체 검출 알고리즘과 유사성 평가지표 4가지를 이용한다. 검증한 결과 제안한 오염 물질 판별 방법이 실제로 오염 물질을 검출하고 유사성을 통해 분류함을 통해 오염물질을 판별함을 확인했다.

I. 서론

실생활에서 음식물이 옷에 묻어 얼룩이 생기는 경우가 종종 발생한다. 옷에 묻은 오염 물질을 인지하지 못하고 방치하게 되면 추후에 세탁을 해도 깨끗이 지워지지 않는다. 이러한 문제를 경감하고자 옷이 오염 되었을 때 즉시 사용자에게 알람을 제공하는 시스템을 제안한다. 따라서 본 논문의 공헌도는 다음과 같다.

오염 물질을 감지해 알람을 주는 방법을 처음으로 제안한다. 제안 방법은 얼룩 검출, 얼룩 잘라내기, 히스토그램 변환, 그리고 유사도 비교단계를 통해 얼룩을 검출한다. 제안 방법의 예측 결과를 4가지 성능 평가지표와 5가지 유사성 평가지표를 통해 분석한다.

II. 배경 지식

히스토그램은 도수분포표를 시각화한 것이다. 가로축은 계급, 세로축은 도수를 의미한다. 히스토그램의 유사도를 측정하는 기준은 대표적으로 Cross correlation distance, Chi-Square distance, Intersection distance, Bhattacharyya distance, 그리고 Cosine distance가 있다. Cosine distance는 두 벡터의 내적을 두 벡터의 크기의 곱으로 나눈다. $[-1, 1]$ 구간의 값을 가지며, 두 벡터가 유사할수록 1에 가까운 값을 갖는다. Cross correlation distance는 Sliding inner product를 통해 두 히스토그램의 분포가 얼마나 유사한지 측정한다. 완전히 일치하면 1, 완전 불일치하면 -1, 연관 관계가 없으면 0의 값을 가진다. Chi-Square distance는 기존 히스토그램에 대해 새로운 히스토그램의 분포가 나타날 가능성을 측정하는 방법이다. $[0, \infty)$ 구간의 값을 가지며 두 분포가 유사할수록 0에 가까운 값을 가진다. Intersection distance는 두 히스토그램의 겹치는 넓이를 구하는 방법이다. 보통 히스토그램을 0과 1사이로 정규화 하여 결과값도 $[0, 1]$ 구간의 값을 갖도록 한다. 데이터가 유사할수록 1에 가까운 값을 갖는다. Bhattacharyya distance는 $[0, 1]$ 구간의 값을 가지며 데이터가 유사할수록 0에 가까운 값을 갖는다.

오염 물질을 검출하기 위해 YOLOv5s 모델을 사용한다. YOLOv5s 모델은 Backbone network로 Cross Stage Partial Network (CSPNet)을 사용한다.

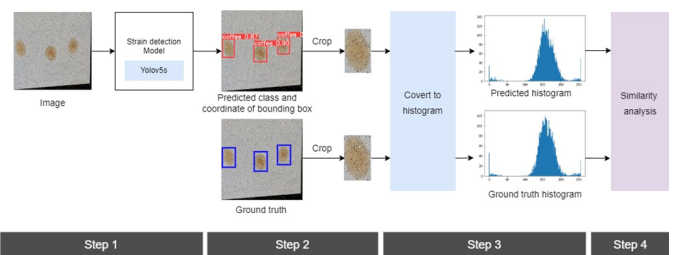


그림 1. 제안 방법

다.[1] CSPNet은 DenseNet을 기반으로 하며 Gradient의 중복을 제거해 더 적은 매개변수와 Floating point Operation Per Second (FLOPS)로 학습할 수 있다.[2]

III. 제안 방법

그림 1은 오염 물질 판별 방법의 단계를 나타낸다. 4단계로 구성된다. 첫 번째 단계에서는, YOLOv5s 모델을 통해 입력 이미지에서 오염 물질의 Bounding box를 검출하고 오염 물질의 Class를 추정한다. 두 번째 단계에서는 입력 이미지에서 예측한 Bounding box와 정답 Bounding box에 해당하는 영역을 각각 잘라낸다. 세 번째 단계에서는 잘라낸 이미지를 256 구간의 Histogram으로 변환한다. 네 번째 단계에서는 예측 Histogram과 정답 Histogram의 유사성을 Cross correlation distance, Chi-Square distance, Intersection distance, 그리고 Cosine distance 평가지표로 각각 분석한다. 그림 2는 3가지 종류의 수집한 오염 물질 데이터 샘플이다. (a)는 커피, (b)는 잉크, 그리고 (c)는 김치 얼룩이다. 수집한 데이터의 클래스 분포를 파악하기 위해 그림 3의 (i)에서 클래스 당 샘플 수를 확인했다. (ii)에는 클래스 당 이미지의 Red, Green, 그리고 Blue 채널의 픽셀 분포를 히스토그램으로 나타냈다. 세 가지 클래스 모두 170 근처에서 최빈값을 가지는 비슷한 분포임을 확인했다. 수집한 데이터 셋으로 YOLOv5s 모델을 전이학습시켜 오염 물질을 검출한다. 실험에 사용한 소스코드는 https://github.com/shonyeajin/Contaminant_Discrimination에 있다.

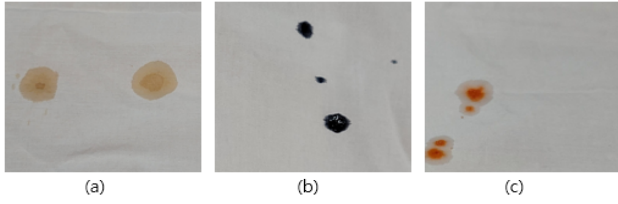


그림 2. 직접 획득한 오염 물질 이미지, (a) Coffee, (b) Ink, (c) Kimchi

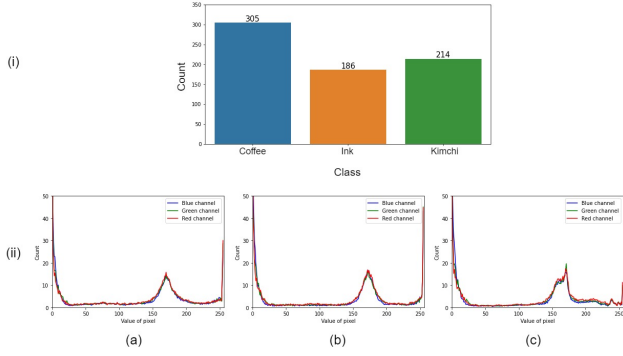


그림 3. 수집한 데이터의 분포, (i) 클래스 당 데이터 샘플 수, (ii) Red, Green, 그리고 Blue 채널의 히스토그램, (a) Coffee, (b) Ink, (c) Kimchi

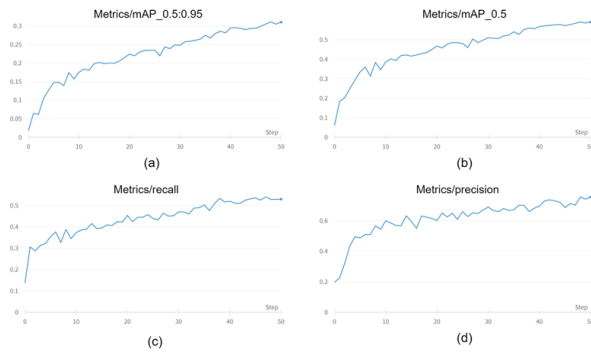


그림 4. 오염 물질 검출 모델의 훈련 중 성능 변화, (a) mAP@.5:.95, (b) mAP@.5, (c) Recall, (d) Precision

IV. 실험 결과

그림 4는 Yolov5s 모델의 전이학습 과정에서 Epoch가 증가에 따른 mean Average Precision (mAP), Recall, 그리고 Precision의 변화를 나타낸다. Epoch가 증가함에 따라 성능이 개선됨을 확인했다. 이는 모델이 훈련 데이터에 Overfitting 되지 않고 패턴을 학습했기 때문이다. 표 1에는 테스트 데이터에 대한 오염 물질 검출 모델의 성능을 나타냈다. 그림 5는 Confusion matrix, F1 score curve, 그리고 Precision-Recall curve이다. 모든 성능 지표에 대해 3가지 오염 물질 중 Coffee 클래스의 예측 성능이 가장 높았다. 예측 히스토그램과 정답 히스토그램의 유사도를 Cross correlation distance, Chi-Square distance, Intersection distance, Bhattacharyya distance, 그리고 Cosine distance 측정하여 표 2에 나타냈다. Chi-Square distance는 $[0, \infty)$ 구간의 값을 가지기 때문에 측정값의 최솟값을 0으로, 최댓값을 1로 투영하여 구간을 $[0, 1]$ 로 변환했다. 전체 클래스에 대해 Cosine distance 적용 시 85.7%, Cross correlation distance 적용 시 85.1%, Intersection distance 적용 시 74.3%, Bhattacharyya distance 적용 시 72.9%, 그리고 Chi-Square 적용 시 54.1% 유사했다. 즉, Cosine distance, Cross correlation distance,

Intersection distance, Bhattacharyya distance, 그리고 Chi-Square distance 측정 지표 순으로 유사도가 높게 나타났다. 5가지 유사도 측정 지표에서 모두 Coffee 클래스의 성능이 가장 뛰어났다. 이는 전체 데이터셋에서 Coffee 클래스의 샘플 수가 가장 많기 때문에 모델이 커피 얼룩을 검출하는 특징을 더 풍부하게 학습한 것으로 추측한다.

V. 결론

오염 물질을 검출하고 유사도를 이용해 판별하는 방법을 제안했다. 수집한 데이터의 분포를 파악하기 위해 3가지 클래스의 샘플 개수를 확인하고 채널별로 픽셀 값의 빈도를 시각화했다. 테스트 데이터에 대한 모델의 성능을 4가지 유사성 지표로 측정했다. 그리고 예측 히스토그램과 정답 히스토그램의 유사도를 5가지 측정 지표를 적용해 비교했다. 이를 통해 제안한 오염 물질 알람시스템의 예측 성능을 다양한 관점으로 검증할 수 있음을 확인했다. 추후에는 Yolov5s 모델 외 다양한 객체 탐지 모델을 사용해 성능을 비교하고자 한다. 또한 제시한 유사도를 기반으로 오염물질을 효과적으로 세탁할 수 있는 물질을 추천하는 방법을 연구하고자 한다.

표 1. 테스트 데이터에 대한 오염 물질 검출 모델의 성능 평가

Class	Precision	Recall	mAP @.5	mAP @.5:.95
Coffee	0.800	0.623	0.700	0.421
Ink	0.761	0.544	0.587	0.300
Kimchi	0.705	0.42	0.485	0.21
All	0.756	0.529	0.591	0.31

표 2. 테스트 데이터에 대한 5가지 유사도 측정 지표 적용 결과, (a) Cross correlation distance, (b) Chi-Square distance, (c) Intersection distance, (d) Bhattacharyya distance, (e) Cosine distance

	(a)	(b)	(c)	(d)	(e)
Coffee	0.896	0	0.799	0.215	0.901
Ink	0.791	0.104	0.684	0.324	0.798
Kimchi	0.839	1	0.715	0.305	0.845
Whole class	0.851	0.459	0.743	0.271	0.857

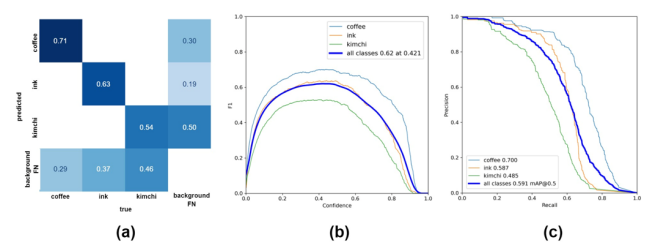


그림 5. 오염물질 검출 모델의 (a) Confusion matrix, (b) F1 score curve, (c) Precision-Recall curve

참 고 문 헌

- [1] Wang, C. Y., Liao, H. Y. M., Wu, Y. H., Chen, P. Y., Hsieh, J. W., and Yeh, I. H. "CSPNet: A new backbone that can enhance learning capability of CNN," in *CVPR*, 2020.
- [2] Huang, G., Liu, Z., Van Der Maaten, L., and Weinberger, K. Q. "Densely connected convolutional networks," in *CVPR*, 2017.