

Central Lancashire Online Knowledge (CLOK)

Title	Deep learning for detection and segmentation of artefact and disease instances in gastrointestinal endoscopy
Type	Article
URL	https://clock.uclan.ac.uk/id/eprint/37165/
DOI	https://doi.org/10.1016/j.media.2021.102002
Date	2021
Citation	Ali, Sharib, Dmitrieva, Mariia, Ghatwary, Noha, Bano, Sophia, Polat, Gorkem, Temizel, Alptekin, Krenzer, Adrian, Hekalo, Amar, Guo, Yun Bo et al (2021) Deep learning for detection and segmentation of artefact and disease instances in gastrointestinal endoscopy. Medical Image Analysis, 70. p. 102002. ISSN 1361-8415
Creators	Ali, Sharib, Dmitrieva, Mariia, Ghatwary, Noha, Bano, Sophia, Polat, Gorkem, Temizel, Alptekin, Krenzer, Adrian, Hekalo, Amar, Guo, Yun Bo, Matuszewski, Bogdan, Gridach, Mourad, Voiculescu, Irina, Yoganand, Vishnusa, Chavan, Arnav, Raj, Aryan, Nguyen, Nhan T, Tran, Dat Q, Huynh, Le Duy, Boutry, Nicolas, Rezvy, Shahadate, Chen, Haijian, Choi, Yoon Ho, Subramanian, Anand, Balasubramanian, Velmurugan, Gao, Xiaohong W, Hu, Hongyu, Liao, Yusheng, Stoyanov, Danail, Daul, Christian, Realdon, Stefano, Cannizzaro, Renato, Lamarque, Dominique, Tran-Nguyen, Terry, Bailey, Adam, Braden, Barbara, East, James E and Rittscher, Jens

It is advisable to refer to the publisher's version if you intend to cite from the work.
<https://doi.org/10.1016/j.media.2021.102002>

For information about Research at UCLan please go to <http://www.uclan.ac.uk/research/>

All outputs in CLOK are protected by Intellectual Property Rights law, including Copyright law. Copyright, IPR and Moral Rights for the works on this site are retained by the individual authors and/or other copyright owners. Terms and conditions for use of this material are defined in the <http://clock.uclan.ac.uk/policies/>



Challenge Report

Deep learning for detection and segmentation of artefact and disease instances in gastrointestinal endoscopy[☆]

Sharib Ali^{a,c,*}, Mariia Dmitrieva^a, Noha Ghatwary^g, Sophia Bano^h, Gorkem Polat^j, Alptekin Temizel^j, Adrian Krenzer^k, Amar Hekalo^k, Yun Bo Guo^l, Bogdan Matuszewski^l, Mourad Gridach^x, Irina Voiculescu^m, Vishnusai Yoganandⁿ, Arnav Chavan^o, Aryan Raj^o, Nhan T. Nguyen^p, Dat Q. Tran^p, Le Duy Huynh^q, Nicolas Boutry^q, Shahadate Rezvy^r, Haijian Chen^s, Yoon Ho Choi^t, Anand Subramanian^u, Velmurugan Balasubramanian^v, Xiaohong W. Gao^r, Hongyu Hu^w, Yusheng Liao^w, Danail Stoyanov^h, Christian Daulⁱ, Stefano Realdon^e, Renato Cannizzaro^f, Dominique Lamarque^d, Terry Tran-Nguyen^b, Adam Bailey^{b,c}, Barbara Braden^{b,c}, James E. East^{b,c}, Jens Rittscher^a

^a Institute of Biomedical Engineering and Big Data Institute, Old Road Campus, University of Oxford, Oxford, UK

^b Translational Gastroenterology Unit, Experimental Medicine Div., John Radcliffe Hospital, University of Oxford, Oxford, UK

^c Oxford NIHR Biomedical Research Centre, Oxford, UK

^d Université de Versailles St-Quentin en Yvelines, Hôpital Ambroise Paré, France

^e Instituto Oncologico Veneto, IOV-IRCCS, Padova, Italy

^f CRO Centro Riferimento Oncologico IRCCS, Aviano, Italy

^g Computer Engineering Department, Arab Academy for Science and Technology, Alexandria, Egypt

^h Wellcome/EPSCRC Centre for Interventional and Surgical Sciences (WEISS) and Department of Computer Science, University College London, London, UK

ⁱ CRAN UMR 7039, University of Lorraine, CNRS, Nancy, France

^j Graduate School of Informatics, Middle East Technical University, Ankara, Turkey

^k Department of Artificial Intelligence and Knowledge Systems, University of Würzburg, Germany

^l School of Engineering, University of Central Lancashire, UK

^m Department of Computer Science, University of Oxford, UK

ⁿ Mimyk Medical Simulations Pvt Ltd, Indian Institute of Science, Bengaluru, India

^o Indian Institute of Technology (ISM), Dhanbad, India

^p Medical Imaging Department, Vingroup Big Data Institute (VinBDI), Hanoi, Vietnam

^q EPITA Research and Development Laboratory (LRDE), F-94270 Le Kremlin-Bicêtre, France

^r School of Science and Technology, Middlesex University London, UK

^s Department of Computer Science, School of Informatics, Xiamen University, China

^t Dept. of Health Sciences & Tech., Samsung Advanced Institute for Health Sciences & Tech. (SAIHST), Sungkyunkwan University, Seoul, Republic of Korea

^u Claritrics India Pvt Ltd, Chennai, India

^v School of Medical Science and Technology, Indian Institute of Technology, Kharagpur, West Bengal, India

^w Shanghai Jiaotong University, Shanghai, China

^x Ibn Zohr University, Computer Science HIT, Agadir, Morocco

ARTICLE INFO

Article history:

Received 21 October 2020

Revised 4 February 2021

Accepted 11 February 2021

Available online 17 February 2021

ABSTRACT

The Endoscopy Computer Vision Challenge (EndoCV) is a crowd-sourcing initiative to address eminent problems in developing reliable computer aided detection and diagnosis endoscopy systems and suggest a pathway for clinical translation of technologies. Whilst endoscopy is a widely used diagnostic and treatment tool for hollow-organs, there are several core challenges often faced by endoscopists, mainly: 1) presence of multi-class artefacts that hinder their visual interpretation, and 2) difficulty in identifying subtle precancerous precursors and cancer abnormalities. Artefacts often affect the robustness of deep learning methods applied to the gastrointestinal tract organs as they can be confused with tissue of interest. EndoCV2020 challenges are designed to address research questions in these remits. In this paper,

[☆] Endoscopy Computer Vision Challenge (EndoCV2020).

* Corresponding author.

E-mail address: sharib.ali@eng.ox.ac.uk (S. Ali).

Keywords:

Endoscopy
Challenge
Artefact
Disease
Detection
Segmentation
Gastroenterology
Deep learning

we present a summary of methods developed by the top 17 teams and provide an objective comparison of state-of-the-art methods and methods designed by the participants for two sub-challenges: i) artefact detection and segmentation (EAD2020), and ii) disease detection and segmentation (EDD2020). Multi-center, multi-organ, multi-class, and multi-modal clinical endoscopy datasets were compiled for both EAD2020 and EDD2020 sub-challenges. The out-of-sample generalization ability of detection algorithms was also evaluated. Whilst most teams focused on accuracy improvements, only a few methods hold credibility for clinical usability. The best performing teams provided solutions to tackle class imbalance, and variabilities in size, origin, modality and occurrences by exploring data augmentation, data fusion, and optimal class thresholding techniques.

© 2021 The Authors. Published by Elsevier B.V.

This is an open access article under the CC BY-NC-ND license (<http://creativecommons.org/licenses/by-nc-nd/4.0/>)

1. Introduction

Endoscopy is a widely used imaging technique for both diagnosis and treatment of patients with complications in hollow organs such as esophagus, stomach, colon, bladder, kidney and nasopharynx. During the endoscopic procedure, an endoscope, a long thin tube with a light source and a camera at its tip, is inserted into the organ cavity. The imaging procedure is usually displayed on a monitor on-the-fly and is often recorded for post analysis. Each organ imposes very specific constraints to the use of endoscopes, but the most common obstructions in all endoscopic surveillance consists of artefacts caused by motion, specularities, low contrast, bubbles, debris, bodily fluid and blood. These artefacts hinder the visual interpretation of clinical endoscopists (Ali et al., 2020c). Missed detection rates of precancerous and cancerous lesions are another limitation. Gastrointestinal (GI) cancer (especially colorectal cancer) has high mortality rates and 5-year relative survival rates for stage IIB is around 65% (Rawla et al., 2019). In general, the missed detection rates in endoscopic surveillance is considerably high, at over 15% (Lee et al., 2017). Therefore, the requirement for technology that can be effectively used in clinical settings during endoscopy imaging is necessary.

While a dedicated endoscopic procedure is followed for each specific organ, often these procedures are very similar, in particular for the GI tract organs like the esophagus, stomach, small intestine, colon and rectum. Notably, some precancerous abnormalities such as inflammation or dysplasia and even cancer lesions in these GI organs naturally look very similar. Often automated methods are only trained for a specific abnormality, organ and imaging modality (Zhang et al., 2019), whereas multiple different types of abnormalities can be present in different organs and several imaging protocols are used during endoscopy. Also, methods that are built for colonoscopy cannot be used during a gastroscopy (in the esophagus, stomach and small intestine), despite the nature and occurrence of many abnormalities being similar in these organs. Artefacts are prevalent in all endoscopy surveillance and are usually confused with lesions, which can lead to unreliable outcomes.

A pathway to develop and reliably deploy methods in clinical settings is by benchmarking methods on a curated multi-center, multi-modal, multi-organ and multi-disease dataset and through a thorough evaluation of built methods using standard imaging metrics and metrics that can test their clinical applicability, for example ranking based on accuracy, robustness and computational efficiency (Ali et al., 2020c). Most publicly available datasets are specific to a particular organ, modality or a single abnormality class, e.g., polyp detection and segmentation challenges (Bernal et al., 2017; Jorge and Aymeric, 2017). While dedicated organ specific challenges help to identify one particular disease type, they do not resemble the clinical workflow where the endoscopists are interested in biopsy and treatment of such abnormalities when

of potential threat. For polyp class, it is required to identify different stages of polyp such as benign, dysplastic or cancer. Recently, it was shown that polyps and artefacts can be confused mostly due to specularities (Soberanis-Mukul et al., 2020). Artefacts are the fundamental and inevitable issue in endoscopy that often add confusion in detecting tissue abnormalities in these organs. It is therefore vital to accelerate research in identifying these classes and restore frames where possible (Ali et al., 2021) or reduce the false detections by adding uncertainties for such confusions (Soberanis-Mukul et al., 2020). Other ways to address artefact problems in the endoscopy data is by using synthetically generated frames (Mahmood et al., 2018; Formosa et al., 2020; Incetan et al., 2020). Mahmood et al. (2018) used self-regularized transformer network that allowed to transform the real images into synthetic-like images with preserved clinically-relevant features. This allowed the authors to estimate depth in colonoscopy data robustly without being affected by adverse artefact problems. Incetan et al. (2020) demonstrated the use of a virtual active capsule environment that can simulate wide range of normal and abnormal tissue conditions such as inflated, dry and wet; organ types and endoscopy camera designs in capsule endoscopy. This allowed to optimize the analysis software for varied real conditions.

The Endoscopy Computer Vision Challenge (EndoCV2020)¹ is another crowd-sourcing initiative to address fundamental problems in clinical endoscopy and consists of: 1) Endoscopy artefact detection and segmentation (EAD2020), and 2) Endoscopy disease detection and segmentation (EDD2020). EndoCV2020 releases diverse datasets that include multi-center, multi-modal, multi-organ, multi-disease/abnormality, and multi-class artefacts. Among the two sub-challenges, EAD2020 is an extended sub-challenge of EAD2019 (Ali et al., 2019), however, unlike EAD2019 it includes both frame and sequence data with an addition of nearly 500 frames and a total of 41,832 annotations for detection task and 10,739 for segmentation task.

In this paper, we summarise and analyze the results of the top 17 (out of 43) teams participating in the EndoCV2020 challenge. Additionally, we benchmark these methods with the current state-of-the-art detection and segmentation methods. Each method is also evaluated for its efficacy to detect and segment multi-class instances. In addition to the standard computer vision metrics used to evaluate methods during the challenge, we perform a holistic analysis of individual methods to measure their clinical applicability.

2. Related work

With the advancements in deep learning for computer vision, object detection and segmentation algorithms have shown rapid

¹ <https://endocv.grand-challenge.org>.

development in recent years. This is due to the hidden feature representations provided by Convolutional Neural Networks (CNNs) that show significant improvement over hand-crafted features. CNN-based methods quickly gained the attention of the Medical Imaging community and are now widely used for automating the diagnosis and treatment for a range of imaging modalities, e.g. radiographs, CT, MRI, and endoscopy imaging. Below we present an overview of the recent deep learning-based object detection and segmentation techniques and discuss the related work in the context to medical image analysis with a particular focus on endoscopy imaging applications.

2.1. Detection and localization

Object detection and localization refers to determining the instances of an object (from a list of predefined object categories) that exist in an image. Object detection approaches can be broadly divided into three categories: single-stage, multi-stage and anchor-free detectors. A brief survey of these is presented below. *Single-stage detectors* Single-stage networks perform a single pass on the data and incorporate anchor boxes to tackle multiple object detection on the same image grid such as in YOLO-v2 (Redmon et al., 2016). Similarly, Liu et al. (2016) proposed the Single Shot MultiBox Detector (SSD) with additional layers to allow detection of multiple scales and aspect ratios. RetinaNet was introduced by Lin et al. (2017b) where the authors introduced focal loss that puts the focus on the sparse hard examples enabling a boost in performance and speed.

The domain of Gastroenterology has started to benefit from the success of single-stage object detectors. Wang et al. (Wang et al., 2018) proposed a model that is based on SegNet (Badrinarayanan et al., 2017) architecture to detect polyps during colonoscopy. Urban et al. (2018) used YOLO to detect polyps from colonoscopy images in real-time. Horie et al. (2019) used SSD to detect superficial and advanced esophageal cancer. RetinaNet was the most popular detector in the first EAD challenge held in 2019. RetinaNet detector with focal loss was used by some top performing teams (Kayser et al., 2019; Oksuz et al., 2019). Multi-stage detectors use a region proposal network to find regions of interest for objects and then a classifier to refine the search to get the final predictions. A two-stage architecture R-CNN using the classical region proposal method was proposed by Girshick et al. (2014) whose speed was improved later by integrating an end-to-end trainable region proposal network (RPN), widely known as Faster R-CNN (Ren et al., 2015). Due to the high precision of the Faster R-CNN, its architecture has become the base for many successful models in the object detection and segmentation domains, such as Cascade R-CNN (Cai and Vasconcelos, 2018) and Mask R-CNN (He et al., 2017). Although these two-stage networks have shown successful results on public datasets such as Pascal VOC (Everingham et al., 2012) and COCO (Lin et al., 2014), they are slow compared to the single-stage object detectors due to their region proposal mechanism.

In the field of Gastroenterology, Yamada et al. (2019) used Faster R-CNN with VGG16 as the backbone to detect challenging lesions which are generally missed by colonoscopy procedures. Their reported prediction speed was not suitable for real-time examination. Shin et al. (2018) detected Polyps using the Fast R-CNN architecture with a region proposal network and an inception ResNet backbone. The two-stage detectors tend to yield better results than their single-stage contemporaries and have performed better at medical image analysis challenges. In the EAD2019 challenge, the top performing team (Suhui Yang, 2019) used a Cascade R-CNN with a feature pyramid network (FPN) module and a ResNet backbone. Similarly, Pengyi and Xiaoqiong (2019) who used Mask aided

R-CNN with an ensemble of different ResNet backbones finished second.

Anchor-free detectors A newly emerging detector type are the anchor-free detectors. Single and multi-stage detectors rely on the presence of anchors. Anchor free architectures claim to detect objects while skipping the process of anchor definition. They rely on different geometrical characteristics like the center or corner points of objects (Law and Deng, 2018; Duan et al., 2019). Duan et al. (2019) utilized the upper left and lower right corner to mark an object. The authors used classical backbones to generate a heatmap from the feature map showing potential spots of the object corners. A corner pooling technique was then used to create the classic bounding box of object detection. Zhou et al. (2019) used a similar approach but instead they used a single point as the center of the bounding box.

Because of real-time dependencies in medical applications like the detection of polyps which have to be removed directly (Wang et al., 2019), anchor-free detectors are receiving more attention. Wang et al. (2019) designed an anchor-free automatic polyp detector which achieved the state-of-the-art results while maintaining real-time applicability. Liu et al. (2020) showed an anchor-free detector with state-of-the-art performance while maintaining real-time performance.

2.2. Semantic segmentation

Semantic segmentation involves pixel-level partitioning of an image into multiple segments where each segment represents a pre-defined object or scene category. Based on the success of deep learning approaches on medical imaging data for segmentation, we can divide these approaches broadly into the following groups: *Models based on fully convolutional networks* Fully Convolutional Network (FCN) architectures include only convolutional layers that enable them to take any arbitrary size input image to output a segmentation mask of the same size. These models are mostly based on the architecture developed by Long et al. (2015) for semantic image segmentation.

Sun et al. (2017) proposed a multi-channel FCN (MC-FCN) to segment liver tumors from multi-phase contrast-enhanced CT images. Kaul et al. (2019) proposed FocusNet for skin cancer and lung lesion segmentation. A benchmark study for polyp segmentation using FCNs was conducted by Gao et al. (2017). Similarly, Brandao et al. (2017) used FCN architecture with VGG backbone for a polyp segmentation task. The same group explored integration of depth information to improve segmentation accuracy in their FCN-based model (Brandao et al., 2018).

Models based on encoder-decoder architecture U-Net (Ronneberger et al., 2015), an encoder-decoder architecture, has become widely popular in medical image analysis community. U-Net based models have shown tremendous success, from cell segmentation (Falk et al., 2019) to liver tumor segmentation (Chlebus et al., 2017) and beyond (Sevastopolsky, 2017; Norman et al., 2018).

In endoscopy imaging, U-Net-based models were used for instrument segmentation on GI endoscopy data (Jha et al., 2020). Khan and Choo (2019) developed a model based on U-Net architecture for endoscopy artefact segmentation. Bano et al. (2020) directly used U-Net architecture for segmenting placental vessels from Fetoscopy imaging. Motion induced segmentation exploiting U-Net in the framework was used to segment kidney stones in the Uteroscopy data (Gupta et al., 2020). *Models based on pyramid-based architecture* In both detection and segmentation tasks, a crucial part is being able to identify objects and features of varying scales and sizes. One approach to this problem is to incorporate convolutional feature maps of varying resolutions during classification, which yields information about different scales of the image,

making it easier to detect both small and big objects. Such architectures are referred to as *pyramid networks*. PSPNet (Zhao et al., 2017) is one of such design that incorporates global context information for the task of scene parsing using a pyramid pooling module. A similar pyramid-based approach can be found in the task of object detection with Feature Pyramid Network (FPN) (Lin et al., 2017a). FPN extracts feature maps on a per-resolution-basis from the two bottom-up and top-down pathways of a pretrained architecture. The output maps can then be upsampled and concatenated to output a segmentation map (Seferbekov et al., 2018).

Guo et al. (2019) used PSPNet as part of an ensemble model including a U-Net and SegNet architecture for the task of automated polyp segmentation in colonoscopy images. Jia et al. (2020) trained a two-stage polyp detector named PLPNet which utilizes FPN for multiscale feature representation using both CVC-ColonDB (Bernal et al., 2012) and CVC-ClinicDB (Bernal et al., 2015). Their experimental results show that PLPNet outperforms other architectures in most regions on CVC-612 dataset (Bernal et al., 2015) and performs similarly on the ETIS dataset (Silva et al., 2014). Zhang and Xie (2019) utilized an FPN combined with a Cascade R-CNN for artefact detection in endoscopic images. *Models based on dilated convolution architecture* One of the challenges in the construction of semantic segmentation networks is to effectively control the size of the receptive field, providing adequate contextual information for pixel-level decisions while, at the same time, maintaining high spatial resolution and computational efficiency. The *dilated* or *atrous* convolution was proposed to address these challenges (Yu and Koltun, 2015). Chen et al. (2018) proposed a family of very effective semantic segmentation architectures, collectively named DeepLab (also an *encoder-decoder* network), all using the dilated convolution. DeepLabv3+ uses atrous kernels within the spatial pyramid pooling (ASPP) module and depth-wise separable convolution to improve the computational efficiency.

Guo et al. (2020a) proposed a fully convolutional network based on atrous kernels to segment polyps in endoscopy images, with their network winning the GIANA 2017 challenge (Jorge and Aymeric, 2017). Nguyen et al. (2020a) augmented DeepLabv3+ architecture, showing its favourable performance when compared with other state-of-the-art methods on the CVC-ClinicDB (Bernal et al., 2015) and ETIS-Larib (Silva et al., 2014) datasets. Ali et al. (2020a) used DeepLabv3+ with ResNet50 backbone to segment Barrett's area from esophageal endoscopy data. Yang and Cheng (2019) developed a model based on DeepLabv3+ for multi-class artefact segmentation used with different backbone architectures.

2.3. Endoscopy computer vision challenges

Biomedical challenges allow to set-up a benchmark for different computer vision methods. Several sub-challenge categories for the development of automated methods for wide-range of problems in endoscopy including surgical instrument segmentation (Ross et al., 2020), robotic scene segmentation (Allan et al., 2020), and computer aided detection and segmentation for polyps (Bernal et al., 2017; 2018) and Barrett's cancer detection² have been initiated under MICCAI EndoVis challenge³. Endoscopy artefact detection (EAD2019) is another challenge which was first initiated in 2019 and launched in conjunction with IEEE International Symposium on Biomedical Imaging (ISBI) 2019 (Ali et al., 2020c).

3. The endocv challenge: Dataset, evaluation and submission

In this section, we present the dataset compiled for the EndoCV2020 challenge, the protocol used to obtain the ground truth for this data, evaluation metrics that were defined to assess participants methods and a brief summary on the challenge setup and ranking procedure.

3.1. Dataset and challenge tasks

The EndoCV2020 challenge consists of two sub-challenges critical in clinical endoscopy. The EAD2020⁴ sub-challenge comprises of diverse endoscopy video frames collected from seven institutions worldwide, including three different modalities and five different human organs (see Fig. 2). Endoscopy video frames were annotated for detection and localization of eight different artefact class occurrences identified by clinical experts in the challenge team. These include specularly, saturation, misc. artefacts, blur, contrast, bubbles, instrument and blood. A total of 280 patient videos from multiple organs and institutions have been used for curating this dataset. Over 45,478 annotations were performed for this challenge on both single frame and sequence video data. Example annotations are shown in Fig. 1. Training data for the detection task consisted of total 2531 frames with 31,069 bounding boxes while 643 frames with 7511 binary masks were released for the segmentation task (except for blur, blood and contrast). The sequence data were sampled by manually observing the amount of changes in artefact categories in the selected sequence. Sequences were required to change from large areas of artefacts to small or no artefact frames and vice versa mimicking natural occurrence in endoscopic procedures. Sequence data for training included 5 sequences (232 frames) for detection and 2 sequences (70 frames) for semantic segmentation tasks sampled from 3 videos of 3 different patients. For the test set, two sequence (80 frames) for detection task were used from 2 independent patient videos. As observed in Fig. 2, due to the nature of occurrence of various artefact classes, the proportion of annotations for each class is different (Fig. 3). However, the proportion of training and test samples per-class were matched in the test data (also see Table 1).

Separately, EDD2020⁵ is a new disease detection and segmentation sub-challenge that consists of five disease categories (Ali et al., 2020b). The provided training set consisted of total 385 video frames comprising of 137 different patients used in this study with a total of 817 individual annotations. The annotations included non-dysplastic Barrett's esophagus (NDBE), suspicious, high-grade dysplasia (HGD), cancer, and polyp categories (also see Fig. 1). These disease classes were from three different endoscopic modalities (white light, narrow-band imaging, and chromoendoscopy) acquired from four different clinical centers, investigating four different GI organs. By including varied range of endoscopy data acquired from multiple organs like GI tract and liver in EAD sub-challenge and both upper and lower GI tract data for EDD sub-challenge, EndoCV2020 challenge aimed at developing more general methods that can potentially be applied in different endoscopy routine procedures independent to organ type. To our knowledge, this is the first comprehensive dataset for the multi-class detection and segmentation tasks. More details on the dataset are provided in Fig. 2. The detailed breakdown of training set and test set for each specific task is provided in Table 1.

EndoCV2020 posed three specific challenge tasks (see Fig. 4) that included: 1) detection and localization task, 2) semantic segmentation task and 3) out-of-sample generalization task. For detection and generalization tasks, participants were provided with

² <https://endovisub-barrett.grand-challenge.org>.

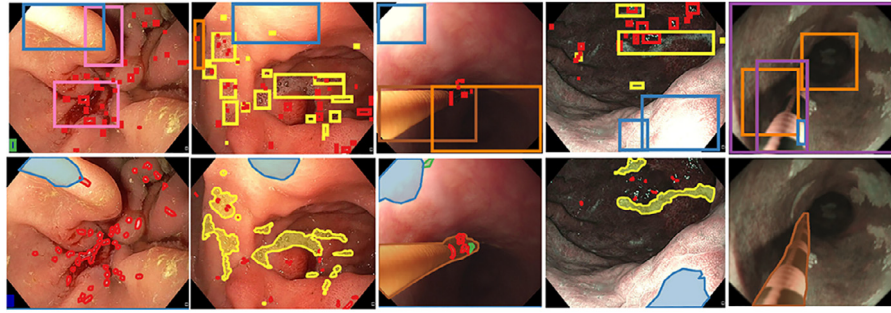
³ <https://endovis.grand-challenge.org>.

⁴ <https://ead2020.grand-challenge.org>.

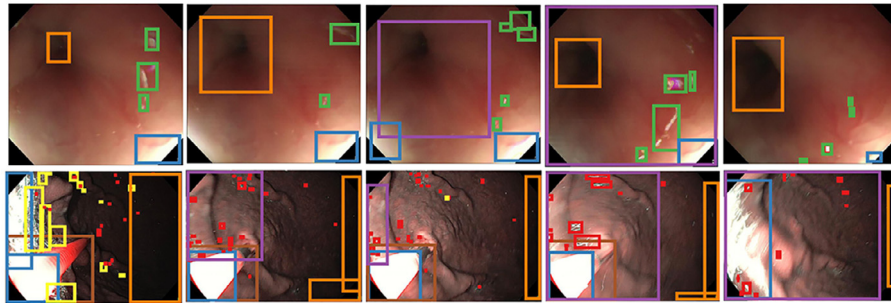
⁵ <https://edd2020.grand-challenge.org>.

a) Endoscopy artefact detection and segmentation

Single frame samples

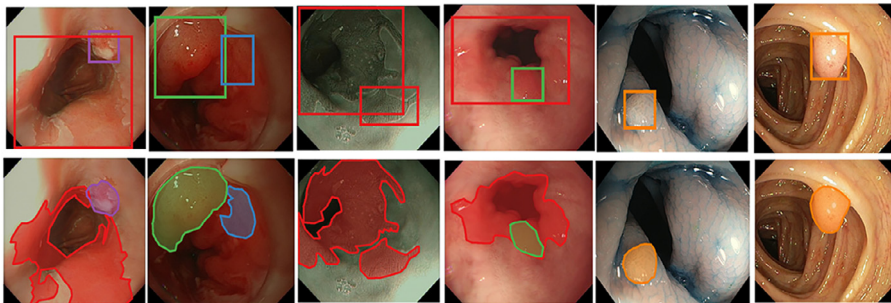


Sequence frame samples



● specularity ● saturation ● artifact ● blur ● contrast ● bubbles ● instrument ● blood

b) Endoscopy disease detection and segmentation



● BE ● suspicious ● High-grade dysplasia (HGD) ● cancer ● polyp

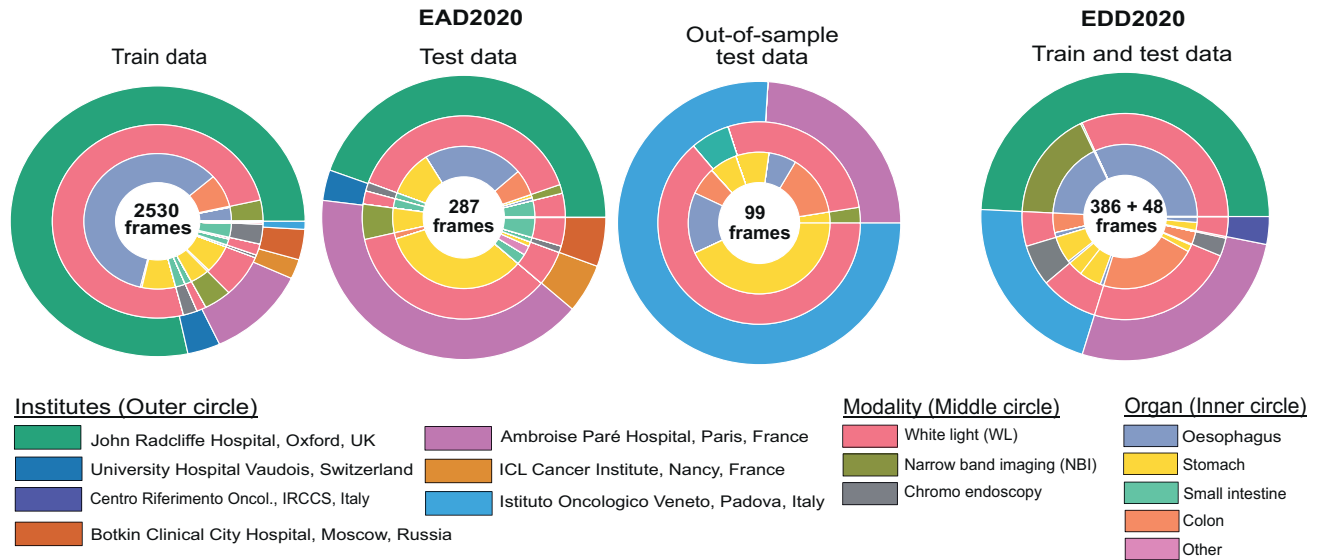
Fig. 1. EndoCV2020 train data samples. (a) Endoscopy artefact detection and segmentation sub-challenge (EAD2020) samples. Both single frame samples (top) and sequence frames (bottom) were released. While detection annotations involve 8 classes, segmentation classes were limited to 5 distinct class instances, mostly large indefinable shapes that include specularity, saturation, imaging artefact, bubbles and instrument. It can be observed that for sequence data most artefact instances follow upto few sequential frames so it is desirable to achieve such training datasets. 4th sample in the single frame data for segmentation shows that even though bounding boxes for detection are provided for all specular regions, some segmentation labels were missing. This shows the presence of annotator variability in the data. (b) Endoscopy disease detection and segmentation training samples for sub-challenge EDD2020. First four samples belong to esophageal endoscopy while the last two frames were acquired during colonoscopy. It can be observed that disease classes in esophagus confuse often, mostly the patient choice here is Barrett's where clearly suspected and high-grade dysplasia appear jointly. Similarly, for colonoscopy data protruded polyps can easily be confused with the surrounding ridge-like openings and specular areas.

Table 1

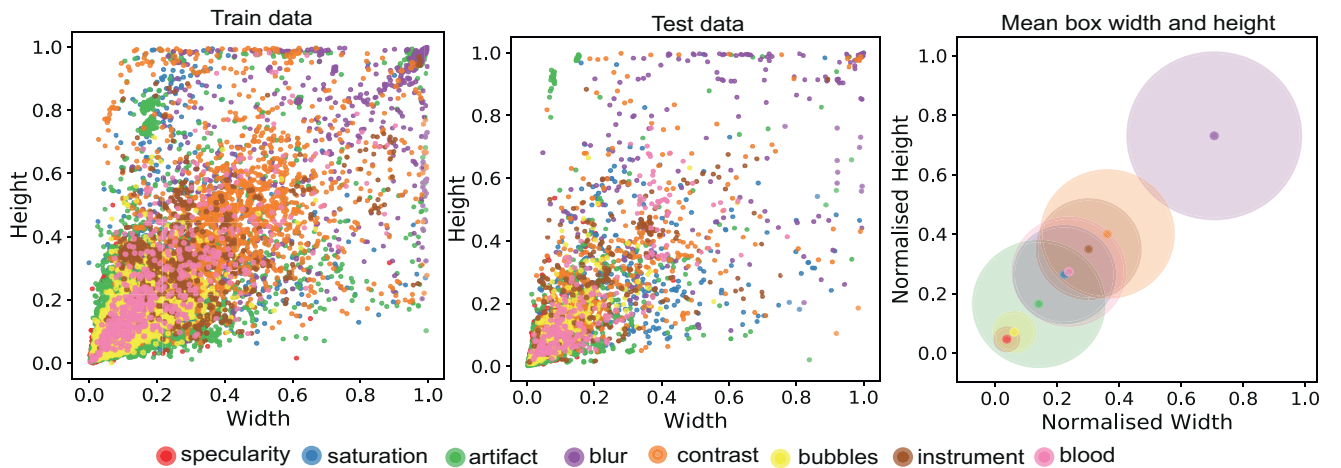
Breakdown of data: Number of samples and annotations released for EndoCV2020 challenge.

EndoCV	Tasks	# of classes	# of frames		# of annotations	
			Train	Test	Train	Test
EAD2020	Detection task	8	single: 2299 seq.: 232	single: 237 seq.: 80	31,069	7750
	Segmentation task	5	643	162	7511	3228
	Generalization task	8	na	99	na	3013
EDD2020	Detection task	5	386	43	749	68
	Segmentation task	5	386	43	749	68

a. EndoCV2020 multi-center data cohort: Train and test data for each sub-challenge



b. EAD2020 train and test sample with per class width and height for detection dataset



c. EDD2020 train and test sample with per class width and height for detection dataset

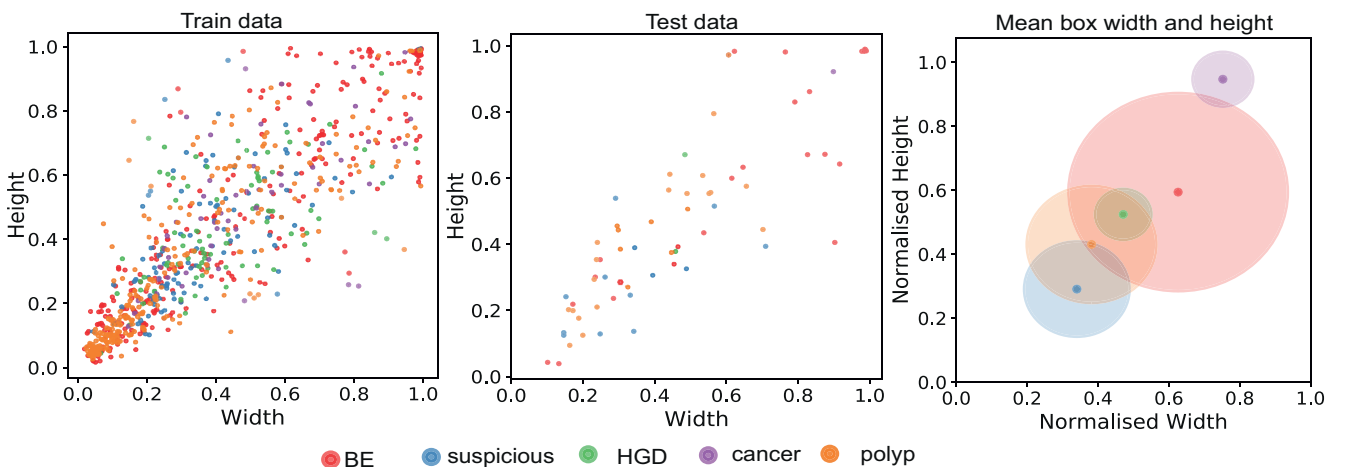


Fig. 2. Endoscopy computer vision EndoCV2020 challenge dataset details. (a) Multi-center, multi-modality and multi-organ dataset for EAD and EDD sub-challenges. For EAD2020, 2532 frames with 8 class bounding boxes for the detection task out-of which 573 included ground truth masks for segmentation task were provided. Participants were assessed on 317 frames for detection and 162 frames for segmentation tasks. An additional 99 frames were used to test out-of-sample generalization task for EAD sub-challenge. While EDD2020 consisted of 384 train samples and 43 test samples for 5 disease classes. (b-c) The distribution of 8 artefact classes of EAD and 5 disease classes of EDD w.r.t. their size compared to their height and width of image is provided. Each class size variability is also shown on right as blobs with mean at center and radius as standard deviation.

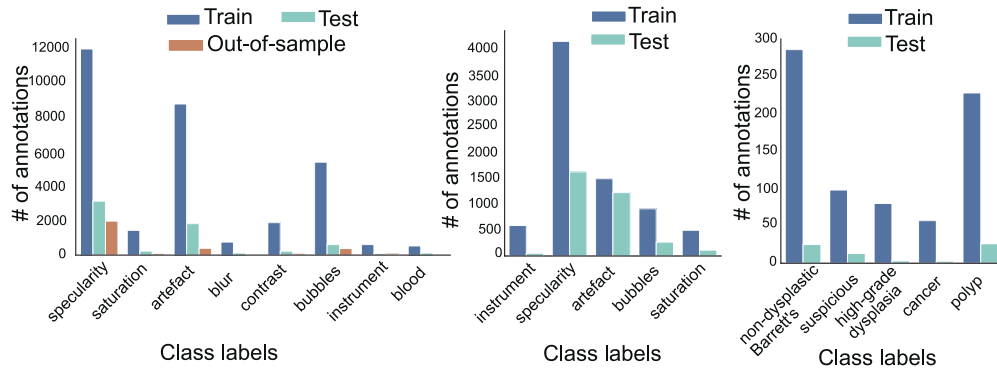


Fig. 3. EndoCV2020 train and test per-class sample proportion: Train and test annotations for sub-challenge on artefact (A,B) and disease (C) detection and segmentation for each class label.

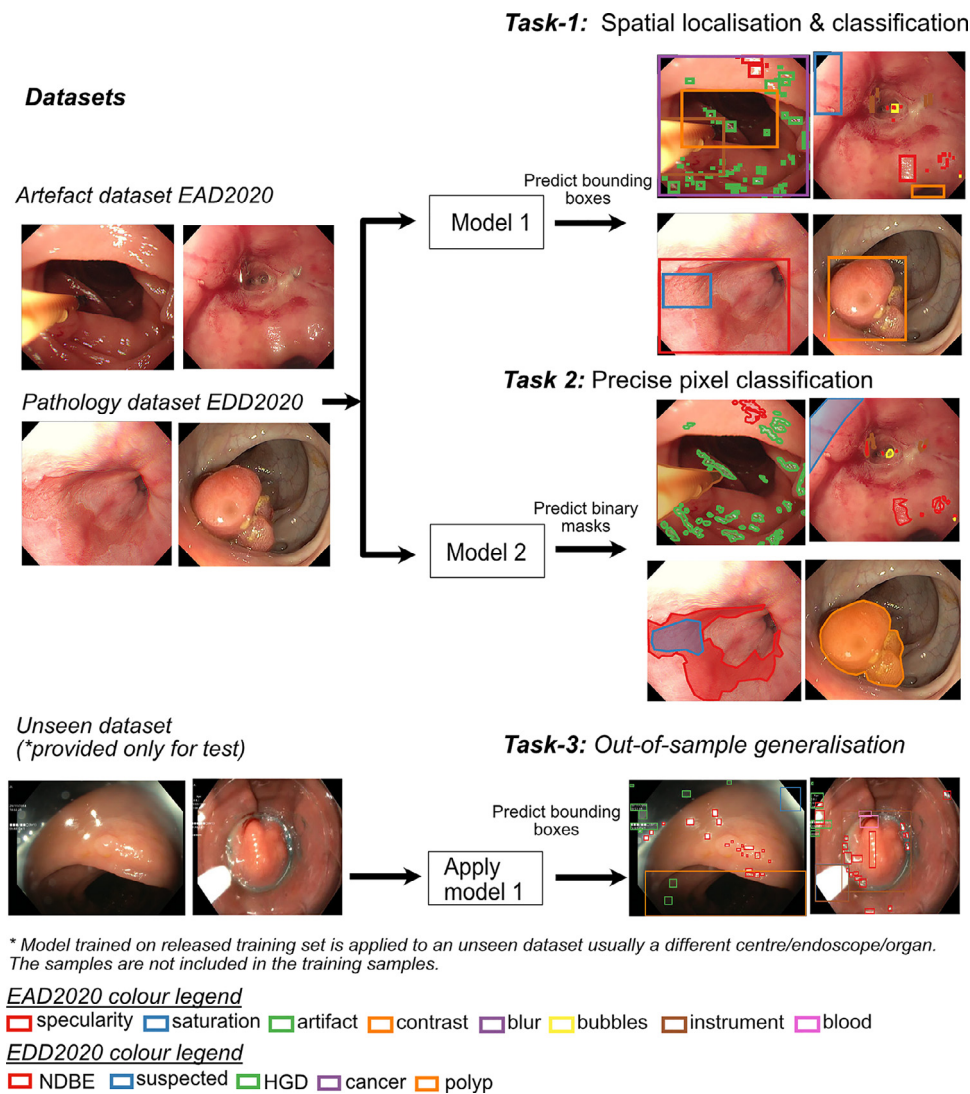


Fig. 4. EndoCV2020 challenge task descriptions for each sub-challenge. The three tasks of the EndoCV2020 challenge includes: (a) The “detection” task aimed at the coarse localization and classification. Given an input image (left) a detection model (middle) outputs the artefact/disease class and coordinates of the containing bounding box. (b) The “segmentation” task is aimed at precise delineation of artefact/disease object boundaries. The model predicts binary output images denoting the presence (‘1’) or absence (‘0’) of each class. (c) The “out-of-sample generalization” task is aimed at assessing the ability of a model trained on different dataset to generalize on an unseen dataset usually coming from a different center.

Table 2

Data collection information for each center: Data acquisition system and patient consenting information.

Centers	System info.	Ethical approval	Patient consenting type
John Radcliffe Hospital, Oxford, UK	Olympus GIF-H260Z, EVIS Lucera CV260	REC Ref: 16/YH/0247	Universal consent
Ambroise Paré Hospital, Paris, France	Olympus Exera 195	N° IDRCB: 2019-A01602-55	Endospectral study
Istituto Oncologico Veneto, Padova, Italy	Olympus endoscope H190	NA	Generic patients consent
Centro Riferimento Oncologico, IRCCS, Italy	Olympus VG-165, CV180, H185	NA	Generic patients consent
ICL, Cancer Institute, Nancy, France	Karl Storz 27005BA	NA	Generic patients consent
University Hospital Vaudois, Switzerland	NA (flexible cystoscopy)	NA	Generic patients consent
Botkin Clinical City Hospital, Moscow, Russia	BioSpec	NA	Generic patients consent

both frame label annotations for single and sequence images for the EAD2020 challenge while only single frames were released for EDD2020. The generalization task was only evaluated for the EAD2020 and only consisted of test data from an unseen institution that was not present in any training set. It is to be noted that test samples for all other tasks were taken from different patients as well even though they were collected from the same centers as that in the training set. EAD2020 attracted nearly 700 participants with 29 teams on the leaderboard and EDD2020 recorded nearly 550 participants with 14 teams on the leaderboard. Participation was permitted in either one or both sub-challenges. Both challenge datasets are publicly available for research and education. EAD2020 challenge data is available at Mendeley Data (<https://doi.org/10.17632/c7fjbxcgj9.3>) and EDD2020 dataset is available at IEEE dataPort (<https://doi.org/10.21227/f8xg-wb80>).

3.1.1. Ethical and privacy aspects of the data

Data for EAD2020 were collected from 7 different centers while for EDD2020 were from 4 centers. Each center was responsible for handling the ethical, legal and privacy of the relevant data sent to the challenge organizers. The data collection from each center included either two or all essential steps described below:

1. Patient consenting procedure at the home institution (required)
2. Review of the data collection plan by a local medical ethics committee or an institutional review board
3. Anonymization of the video or image frames (including demographic information) prior to sending to the organizers (required)

Table 2 illustrates the ethical and legal processes fulfilled by each center along with the endoscopy equipment used for the data collected for this challenge.

3.1.2. Annotation protocol

A team of two clinical experts and one post-doctoral researcher determined the class labels for the artefact detection challenge while for the disease detection challenge we consulted with four senior Gastroenterologists (over 20 years experience) regarding the class labels in the GI tract endoscopy. For each sub-challenge senior Gastroenterologists sampled the video frames from a small sub-set of video data collected from various institutions and multi-patient data cohort (see Fig. 2). These frames were then taken as reference to produce bounding box annotations for the remaining train-test dataset by four experienced postdoctoral fellows. Finally, further validation by three clinical endoscopists independently was carried out to assure the reference standard. The ground-truth labels were randomly sampled (1 per 20 frames) during this process. However, after the completion of this phase the entire annotation was discussed and reviewed together with the team of senior Gastroenterologists. Priority was given to indecisive frame annotations to have a collective opinion from experts. Following general annotation strategies were used by clinical experts and researchers:

- For the same region, multiple boxes (for detection/generalization) or pixel-wise delineation (for semantic seg-

mentation) were performed if the region belonged to more than 1 class

- The minimal box sizes were used to describe the class region and similarly possible small annotation areas for semantic segmentation were merged instead of having multiple small boxes/regions
- Each class type was determined to be distinctive and general across all datasets

For EAD dataset, defined class categories used included below descriptions (Ali et al., 2021). Related samples are presented in Fig. 1(a).

1. blur → fast camera motion
2. bubbles → a thin film of liquid with air that distorts tissue appearance
3. specularity → mirror-like reflection
4. saturation → overexposed bright pixel areas
5. contrast → low contrast areas from underexposure
6. misc. artefact → chromatic aberration, debris etc.
7. instrument → biopsy or any other instrument
8. blood → flow of red colored liquid due to biopsy or surgery

For EDD dataset, both upper-GI (gastroscopy) and lower-GI (colonoscopy) data were used with below defined class categories (please refer to the samples in Fig. 1(b)):

1. NDBE or BE → non-dysplastic Barrett's esophagus determined by a squamo-columnar junction above the gastric fold in the esophagus (Eluri and Shaheen, 2017)
2. HDG → high-grade dysplasia or early adenocarcinoma determined by irregular mucosal appearance (Wang et al., 2012)
3. suspected → aka low-grade dysplasia, an early sign of pathology (Eluri and Shaheen, 2017)
4. cancer → abnormal growth (Boland et al., 2005)
5. polyp → abnormal protrusion of the mucosa (Williams et al., 2013)

For the annotations of disease classes, pathology reports were also used to validate the class category for non-dysplastic Barrett's esophagus (BE), high-grade dysplasia (HGD), suspected (dysplasia or low-grade dysplasia), and cancer categories. That is, expert annotations (three senior gastroenterologists) were taken and supported with the pathology report for most disease categories including some indecisive cases. However, for the polyp class, both the protruded and flat polyps were marked by two experienced post-doctoral researchers and checked by a senior lower-GI specialist (no further categorization based on pathology report was done except for cancer cases).

3.2. Evaluation metrics

The challenge problems fall into three distinct categories. For each there already exist well-defined evaluation metrics used by the wider imaging community which we use for evaluation here.

Codes related to all evaluation metrics used in this challenge are also available online⁶.

3.2.1. Spatial localization and classification task

Metrics used for multi-class disease detection:

- IoU - intersection over union: This metric measures the overlap between two bounding boxes A and B , where A is segmented region and B is annotated GT. It is evaluated as the ratio between the overlapped area $A \cap B$ over the total area $A \cup B$ occupied by the two boxes:

$$\text{IoU} = \frac{A \cap B}{A \cup B} \quad (1)$$

where $\cap$, $\cup$ denote the intersection and union respectively. In terms of numbers of true positives (TP), false positives (FP) and false negatives (FN), IoU (aka Jaccard JC) can be defined as:

$$\text{IoU/JC} = \frac{TP}{TP + FP + FN} \quad (2)$$

- mAP - mean average precision: mAP of detected class instances is evaluated based on precision (p) defined as $p = \frac{TP}{TP + FP}$ and recall (r) as $r = \frac{TP}{TP + FN}$. This metric measures the ability of an object detector to accurately retrieve all instances of the ground truth bounding boxes. Average precision (AP) is computed as the Area Under Curve (AUC) of the precision-recall curve of detection sampled at all unique recall values ($r_1, r_2, \dots$) whenever the maximum precision value drops:

$$\text{AP} = \sum_n \{ (r_{n+1} - r_n) p_{\text{interp}}(r_{n+1}) \}, \quad (3)$$

with $p_{\text{interp}}(r_{n+1}) = \max_{\tilde{r} \geq r_{n+1}} p(\tilde{r})$. Here, $p(r_n)$ denotes the precision value at a given recall value. This definition ensures monotonically decreasing precision. The mAP is the mean of AP over all N classes given as

$$\text{mAP} = \frac{1}{N} \sum_{i=0}^N \text{AP}_i \quad (4)$$

This definition was popularised in the PASCAL VOC challenge (Everingham et al., 2012). The final mAP (mAP_d) was computed as an average mAPs for IoU from 0.25 to 0.75 with a step-size of 0.05 which means an average over 11 IoU levels is used for 5 categories in the competition ($\text{mAP} @ [0.25 : 0.75]$).

Participants were finally ranked on a final mean score (score_d), a weighted score of mAP and IoU represented as:

$$\text{score}_d = 0.6 \times \text{mAP}_d + 0.4 \times \text{IoU}_d \quad (5)$$

Standard deviation between the computed mAPs ($\pm \sigma_{\text{score}_d}$) are taken into account when the participants have the same score_d . Scores on both single frame data and sequence data were first separately computed and then averaged to get the final score_d of the detection task.

3.2.2. Segmentation task

Metrics widely used for multi-class semantic segmentation of disease classes have been used for scoring semantic segmentation. The final semantic score score_s comprises of an average score of F_1 -score (Dice Coefficient, DSC), F_2 -score, precision (PPV), recall (Rec) and accuracy (Acc).

Precision, recall, F_β -scores:

These measures evaluate the fraction of correctly predicted instances. Given a number of true instances #GT (ground-truth

bounding boxes or pixels in image segmentation) and number of predicted instances #Pred by a method, precision is the fraction of predicted instances that were correctly found, $\text{PPV} = \frac{\#TP}{\#Pred}$, where #TP denotes number of true positives and recall is the fraction of ground-truth instances that were correctly predicted, $\text{Rec} = \frac{\#TP}{\#GT}$. Ideally, the best methods should have jointly high precision and recall. F_β -scores gives a single score to capture this desirability through a weighted (β) harmonic means of precision and recall, $F_\beta = (1 + \beta^2) \cdot \frac{\text{PPV} \cdot \text{Rec}}{(\beta^2 \cdot \text{PPV}) + \text{Rec}}$.

Participants are ranked based on the value of their semantic performance score given by:

$$\text{score}_s = 0.25 \times (p + r + F_1 + F_2) \quad (6)$$

Standard deviation between each of the subscores are computed and averaged to obtain the final $\pm \sigma_{\text{score}_s}$ which is used during evaluation for participants with same final semantics score. We have also used provided accuracy of each semantic method in this paper for scientific completeness. Accuracy (Acc) can be defined as $\text{Acc} = \frac{TP + TN}{TP + TN + FP + FN}$.

3.2.3. Out-of-sample generalization task

Out-of-sample generalization of disease detection is defined as the ability of an algorithm to achieve similar performance when applied to a completely different institution data. To assess this, participants were challenged to apply their trained models on video frames that were neither included in the training nor in the test data of the other tasks. Assuming that participants applied the same trained weights, the out-of-sample generalization ability was estimated as the mean deviation between the mAP score of the detection and out-of-sample generalization test datasets of each class i for deviation greater than a tolerance of $\{0.1 \times \text{mAP}_d^i\}$.

$$\text{dev}_g = \frac{1}{N} \sum_i \text{dev}_g^i \quad (7)$$

$$\text{dev}_g^i = \begin{cases} 0, & \text{for } |\text{mAP}_d^i - \text{mAP}_g^i| / \text{mAP}_d^i \leq 0.1 \\ |\text{mAP}_d^i - \text{mAP}_g^i|, & \text{for } |\text{mAP}_d^i - \text{mAP}_g^i| / \text{mAP}_d^i > 0.1 \end{cases} \quad (8)$$

The best algorithm should have high mAP_g and low $\text{dev}_g (\rightarrow 0)$. Participants were finally ranked using a weighted ranking score for out-of-sample generalization as $R_{\text{gen}} = 1/3 \cdot \text{Rank}(\text{dev}_g) + 2/3 \cdot \text{Rank}(\text{mAP}_g)$ where $\text{Rank}(\text{mAP}_g)$ is the rank of a participant when sorted by mAP_g in ascending order.

3.3. Challenge setup, and ranking procedure

The challenge proposal was submitted to the IEEE ISBI challenge organisers and was peer-reviewed by two reviewers. Upon the acceptance, the challenge website⁷ was launched on 1st November 2019. Training datasets for each sub-challenge (EAD and EDD) were first provided (via AWS amazon S3 for EAD data and IEEE data portal for EDD data⁸). The test data was released nearly 20 days before the leaderboard closing through a docker container set-up. A docker based online leaderboard was established separately for EAD2020⁹ and EDD2020¹⁰ where each participating team was allowed to submit a maximum of 2 submissions per day on the final test data. A wiki-page¹¹ was set-up for the submission guidelines

⁷ <https://endocv.grand-challenge.org>.

⁸ <https://ieee-dataport.org/competitions/endoscopy-disease-detection-and-segmentation-edd2020>.

⁹ <https://ead2020.grand-challenge.org/evaluation/leaderboard/>.

¹⁰ <https://edd2020.grand-challenge.org/evaluation/leaderboard/>.

¹¹ <https://github.com/sharibox/EndoCV2020/wiki>.

⁶ <https://github.com/sharibox/EndoCV2020>.

and a code repository with evaluation metrics used in the challenge was also provided¹².

For the ranking of different task categories, we used the metrics described in Section 3.2. The participants were able to see only the final score in the leaderboard and all other sub-scores were hidden for the final test data. This was done to avoid any class specific refinement on the released test set. Notably, the detection task was bounded by two IoU thresholds (mAP @ IoU thresholds [.25 : .05 : .75]) and the overall IoU scores itself. For the detection task, participants were ranked on a final weighted score of mAP and IoU (see Eq. (5)), while for the segmentation task, participants were ranked based on a final weighted average of DSC or F1-score, F2-score, precision and recall (see Eq. (6)). For the generalization task, both the mAP score gap dev_g and mAP on generalization data mAP_g were taken into account.

4. Method summary of the participants

In this Section, we present summary of top participating teams for both EAD2020 and EDD2020 sub-challenges. Each of these teams has participated in either detection task or segmentation task or both.

4.1. EAD2020 Participating teams

- **Team polatgorkem** (Polat et al., 2020) The team used an ensemble of three object detectors: Faster R-CNN (ResNet50 with FPN), Cascade R-CNN (ResNet50 with FPN), RetinaNet (ResNet101 with FPN). Class-agnostic NMS operation, where the model predictions were passed through the NMS procedure together for all classes, was applied to the output of each individual model. During ensemble, only the bounding boxes for which majority of the models agree were kept. False-positive elimination was applied as a post-processing step to eliminate same-type predicted boxes located close to each other. For each class, an IoU threshold was determined.
 - **Team CVML** (Guo et al., 2020b) CVML team's model was inspired by DeepLabV3+. The team experimented with several changes including the backbone, the global pooling, the dilated kernels and the convolution kernels with dilation rates. Moreover, the squeeze-and-excitation module is added behind the balanced ASPP module to introduce attention gating at the output of the original encoder to better utilize the information available in the computed feature maps. In addition, the original multi-class classifier is replaced with 5 binary classifiers to enable segmentation of the overlapping objects. At test time, they used some post-processing techniques such as rotation, holes filling and removal of objects from the image boundary.
 - **Team mouradai_ox** (Gridach and Voiculescu, 2020) The team proposed a novel neural network called OxEndoNet to tackle the segmentation challenge. The network uses the pyramid dilated module (PDM) consisting of multiple dilated convolutions stacked in parallel. For each input image, pre-trained ResNet50 (on ImageNet) was used as the backbone to extract the feature map followed by multiple PDM layers to form an end-to-end trainable network. In the final architecture, they used four PDM layers; each layer used four parallel dilated convolutions with a filter size of 3×3 and dilation rates of 1, 2, 3, and 4. They fed the final PDM layer to a convolution layer followed by a bilinear interpolation to up-scale the feature map to the original image size.
 - **Team mimykgcp** (Y et al., 2020) The team re-trained the ResNeXt101 backbone with the cardinality parameter set to 64.
- To enable detection of artefacts at different scales, an FPN was integrated into the object detectors. Data-Augmentation techniques based on RandAugment (Cubuk et al., 2019) were incorporated to improve the generalization capability. For the segmentation task, a U-Net with an ImageNet pre-trained ResNext50 backbone was used.
- **Team DuyHUYNH** (Huynh and Boutry, 2020) For segmentation, the team exploited a model based on U-Net++ using pre-trained EfficientNet on ImageNet as the backbone. The model was trained to minimize F2-loss using the Adam optimizer. At the test-time the team used five transformations: horizontal, vertical flipping, and three rotations. For detection, the team used the bounding boxes deduced from the results of their segmentation model on the EDD dataset, while for EAD, they used YOLOv3 pre-trained on COCO.
 - **Team mathew666** (Hu and Guo, 2020) The team used Cascade RCNN architecture with the ResNeXt backbone in a FPN based feature extraction paradigm. Data augmentation with probability of 0.5 for horizontal flip was applied. The team also utilised multi-scale detection to tackle with variable sized object detection.
 - **Team arnavchavan04** (Jadhav et al., 2020) For the object detection task, the team used an ensemble of three models: Faster R-CNN (ResNet101 + FPN), RetinaNet (ResNet101 + FPN) and Faster R-CNN (ResNet101 + DC5). For the segmentation task, an ensemble of multiple depth EfficientNet models with FPN trained on multiple optimization plateaus (DSC, BCE, IoU) was designed. Data augmentation techniques like horizontal and vertical flip, cutout (random holes), random contrast, gamma, brightness, rotation along with CutMix (Yun et al., 2019) strategy for the segmentation task were incorporated to improve generalization capability.
 - **Team anand_subu** (Subramanian and Srivatsan, 2020) The team used RetinaNet with ResNet101 backbone. For the segmentation task, the team used an ensemble network with U-Net with a ResNet50 backbone and DeepLabV3. However, the team reported U-Net with ResNet101 as their best architecture of choice. All the backbones were pre-trained on the ImageNet. Real-time augmentation techniques like rotation, shear, random-image-flip, image contrast, brightness, saturation, and hue variations were incorporated while training to improve the generalization capability of the network.
 - **Team higersky** (Chen et al., 2020) The team implemented Hyper Task Cascade and Cascade R-CNN with ResNeXt101 backbone as a feature extractor and FPN module for multi-scale feature representation for the object detection task. They applied Soft-NMS (Bodla et al., 2017) to avoid mistakenly discarded bounding-boxes. For the semantic segmentation task, the team incorporated DeepLabV3+ with ResNet101 backbone and trained with BCE and DICE losses. The backbones for both tasks were pre-trained on ImageNet.
 - **Team MXY** (Yu and Guo, 2020) The team used a Cascade R-CNN with an ImageNet pre-trained ResNet101 backbone and a FPN module. Post-detection, soft-NMS was added to remove false predictions. The dataset was augmented by random resizing technique to improve the final output scores. The team used more weight for the losses of specularly, artefact, and bubbles classes to overcome classification difficulties between those classes.
 - **Team StarStarG** The team used Cascade-RCNN as network architecture and adopted COCO2017 pre-trained ResNeXt as backbone with FPN and multi-stage RCNN framework. The authors also integrated Deformable Convolutional Networks in backbone to improve the model performance.
 - **Tesam xiaohong1** (Gao and Braden, 2020) The team built their detection and segmentation method upon Yolact-based instance

¹² <https://github.com/sharibox/EndoCV2020>.

Table 3
Endoscopy artefact detection and segmentation (EAD2020) method summary for top 13 teams (out-of 33 valid submissions).

Team EAD2020	Algorithm	Preprocessing	Nature	Basis-of-choice	Backbone	Data aug.	Pretrained	Computation		code
Detection								GPU	Test time	
polatgorkem (METU_DLCV)	Faster RCNN + CascadeRCNN + Retinanet	Resize Normalise	Ensemble	Accuracy+	ResNet50, ResNet101	Yes (R, F) ^a	COCO	RTX 2080	0.76	GorkemP/EAD
qzheng5 (CVML)	Faster RCNN	Resize Normalise	Context	Accuracy+	ResNet101	Yes (R, T, LD) ^a	COCO	GTX1060	0.20	CVML/EAD2020
xiaohong1	YOLOACT + NMS-within-class	None	Context	Accuracy+, speed+	ResNet101	None	ImageNet	Tesla K80	0.14	yolact
mathew666	Faster RCNN + NMS	None	Context	Accuracy+	ResNet101	Yes	NA	RTX 2080	NA	NA
VinBDI	EfficientDet D0	Resize (512x512)	Multiscale scalable	Speed+	EfficientNet B0	Yes (S, Sc, R, N, MU) ^a	COCO	RTX 2080Ti	NA	endocv2020-seg
higersky	Cascade R-CNN	None	Cascading	Accuracy+	ResNeXt101	Yes	NA	GTX1080 Ti	NA	NA
StarStarG	Cascade R-CNN	Resize Normalise	Cascading	Accuracy+	ResNeXt101	Yes (F, S) ^a	NA	RTX 2080	NA	NA
anand_subu	RetinaNet	Resize Normalise	Context	Accuracy+, speed+	ResNet101	Yes (R, Sh, F, C, B, St, H) ^a	ImageNet	GTX1050Ti	0.36	anand-subu/EAD2020
arnavchavan04	RetinaNet + FasterRCNN (FPN + DC5)	Resize (512x512)	Ensemble	Accuracy+	ResNet50; ResNeXt101	Yes (F, C, R) ^a	ImageNet	Tesla T4	NA	ubamba98/EAD2020
MXV	Cascase RCNN + FPN	Resize Normalise	Cascading	Accuracy+	ResNet101	Yes (F) ^a	ImageNet	RTX 2080 Ti	0.80	Carboxy/EAD2020
mimykqcp	Faster RCNN + RetinaNet	Resize Normalise	Ensemble	Accuracy+, speed+	ResNeXt101	Yes (RA) ^a	COCO	GTX 1080Ti	0.58	NA
DuyHUYNH (LRDE)	YOLOv3	Normalise	Multiscale	Accuracy+, speed++	Darknet53	Yes (RA) ^a	COCO	GTX1080 Ti	0.07	dhuynh/endocv2020
Segmentation										
qzheng5 (CVML)	DeepLabv3+	Resize (513x513) Normalise	Encoder-decoder, mutiscale	Accuracy+	SE-ResNeXt50	(R, T, LD + TTA) ^a	ImageNet	GTX1080Ti	0.50; 5 (+TTA)	CVML/EAD2020
mouradai_ox	Pyramid dilated module	Resize (512x512) Normalise	Multiscale	Accuracy+, speed+	ResNet50	Yes (T, R, LD) ^a	ImageNet	Colab	0.37	NA
arnavchavan04	FPN + EfficientNet	Resize (512x512)	Ensemble	Accuracy+	EfficientNet	Yes (F, C, R) ^a	ImageNet	Tesla T4	NA	ubamba98/EAD2020
VinBDI	U-Net + BiFPN	Resize (512x512)	Ensemble, Endcoder-decoder	Accuracy+, speed+	EfficientNet B4; ResNet50	Yes (S, Sc, R, F) ^a	COCO ImageNet	RTX 2080Ti	NA	endocv2020-seg
higersky	DeepLabv3+	None	Encoder-decoder, mutiscale	Accuracy+	ResNet101	Yes (F;S;Sc;BI) ^a	ImageNet	GTX1080 Ti	NA	NA
anand_subu	U-Net	Resize (512x512)	Encoder-decoder	Accuracy+	ResNet50	Yes (S, F, R, N, Cr, BI, H, St, C, Sp) ^a	ImageNet	GTX1050Ti	0.17	anand-subu/EAD2020
DuyHUYNH (LRDE)	U-Net+	Normalise	Encoder-decoder	Accuracy+, speed+	EfficientNet B1	Yes (R, S, F, Sc, LD, TTA) ^a	ImageNet	GTX1080 Ti	0.97	dhuynh/endocv2020
mimykqcp	U-Net	Resize Normalise	Encoder-decoder	Accuracy+, speed+	ResNeXt50	Yes (RA) ^a	ImageNet	RTX 2070	0.25	NA

^a B: brightness, C: contrast, F: Flip, H: hue, LD: Local deformation, N: noise, R: Rotation, RA: RandAugment, S: Shift, Sc: scaling Sh: shear, St: saturation, Mu: mixup, T: Translation, TTA: test-time augmentation

Table 4
Endoscopy disease detection and segmentation (EDD2020) method summary for top 7 teams (out-of 14 submission).

Team EDD2020	Algorithm	Preprocessing	Nature	Basis-of-choice	Backbone	Data aug.	Pretrained	Computation		code
Detection								GPU	Test time	
Adrian	YOLOv3+ Faster R-CNN	Resize	Ensemble	Accuracy+, speed+	Darnet53 ResNet101	Yes (F, D) ^a	COCO public polyp dataset	Tesla P100	0.41	Adrian398/EDD
shahadate	Mask R-CNN	Resize Normalise	Multiscale	Accuracy, speed+	ResNet101	Yes (Sc, R, F, Cr, S, N) ^a	COCO	RTX2060	NA	EDD-Mask-rcnn
VinBDI	EfficientDet D0	Resize (512x512)	Ensemble	Speed+	EfficientNet B0	Yes (S, Sc, R, N, MU) ^a	COCO	RTX 2080Ti	NA	endocv2020-seg
YH_Choi	CenterNet	NA	Context	Accuracy+	ResNet50	Yes(Du, R, F, C, B) ^a	PASCAL VOC2012	RTX 2080	2	NA
DuyHUYNH (LRDE)	U-Net+	Normalise	Encoder-decoder	Speed	EfficientNet B1	Yes (R, S, F, Sc, LD, TTA) ^a	ImageNet	GTX1080 Ti	1.53	dhuyinh/endocv2020
mimykcp (vishnusai)	Faster RCNN + RetinaNet	Resize (256x256) normalise	Ensemble	Accuracy+, speed+	ResNeXt101	Yes (RA) ^a	COCO	GTX1080Ti	0.58	NA
Segmentation										
Adrian	YOLOv3 + Faster R-CNN + Cascade RCNN	Resize	Ensemble	Accuracy+	Darnet53 ResNet101	Yes (F, D) ^a	COCO public polyp dataset	Tesla P100		Adrian398/EDD2020
shahadate	MaskRCNN	Resize Normalise	Multiscale	Accuracy, speed+	ResNet101	Yes (Sc, R, F, Cr, S, N) ^a	COCO	RTX2060		EDD-Mask-rcnn
VinBDI	U-Net + BiFPN	Resized (512x512)	Ensemble Endcoder-decoder	Accuracy+, speed+	EfficientNet B4 ResNet50	Yes (S, Sc, R, F) ^a	COCO ImageNet	RTX 2080 Ti	NA	endocv2020-seg
YH_Choi	U-Net	NA	Encoder-decoder	Accuracy+	ResNet50	Yes(Du, R, F, C, B) ^a	PASCAL VOC2012	RTX 2080	7	NA
DuyHUYNH (LRDE)	U-Net+	Normalise	Encoder-decoder	Accuracy+, speed+	EfficientNet B1	Yes (R, S, F, Sc, LD, TTA) ^a	ImageNet	GTX1080 Ti	1.53	endocv2020
drvelmuruganb	SUMNet	NA	Encoder-decoder	Accuracy+, speed++	VGG11	Yes(R, A, Sc, P, and Cr) ^a	ImageNet	GTX1080 Ti	0.16	drvelmuruganb/EDD2020
mimykcp	U-Net	Resize Normalise	Encoder-decoder	Accuracy+	ResNeXt50	Yes (RA) ^a	ImageNet	RTX2070	1.25	NA

^a A: affine, B: brightness, C: contrast, Cr: cropping, D: distortion, Du: duplication, F: flip, H: hue, LD: local deformation, Mu: mixup, N: noise, P: perspective transformation, R: rotation, RA: RandAugment library, S: shift, Sc: scaling, Sh: shear, St: saturation, T: translation, TTA: test-time augmentation

segmentation system. Yolact (Bolya et al., 2019) adds a segmentation component to the RetinaNet to ensure the tasks of detection, classification and delineation which are performed simultaneously. The network uses ResNet101 as an imageNet pre-trained backbone.

4.2. EDD2020 Participating teams

- **Team Adrian** (Krenzer et al., 2020) The team compared two different models: YOLOv3 with darknet-53 backbone and Faster R-CNN with ResNet-101 backbone. For post-processing, both algorithms in the final architecture were combined. For the second task, the team leveraged the state-of-the-art Cascade Mask R-CNN with ResNeXt-151 as a backbone. The team trained YOLOv3 using categorical cross-entropy for classification and default localization loss, while for Cascade Mask-RCNN, they used binary cross entropy for classification and mask, and L1 smooth for boundary box regression.
- **Team Shahadate** (Rezvy et al., 2020) The team implemented a modified benchmark Mask R-CNN infrastructure model on the EDD2020 dataset. They used COCO trained weights and biases with the ResNet101 backbone as an initial feature extractor. The network head of the backbone model was replaced with new untrained layers that consisted of a fully-connected classifier with five classes and an additional background class. Non-maximum suppression was used to reduce overlapped detection. Finally, the team merged multiple bounding boxes for the same class label as one bounding box to match with the mask annotation.
- **Team VinBDI** (Nguyen et al., 2020b) For the object detection task, the team designed an ensemble of six EfficientDet models (with BiFPN modules) trained on six different EfficientNet backbones. A total of eleven augmentation techniques were incorporated to increase the output prediction scores of the model. For the segmentation task, an ensemble of U-Net and EfficientNet-B4 and BiFPN with the ResNet50 backbone was devised. The same team also participated in the EAD2020 sub-challenge.
- **Team YH_Choi** (Choi et al., 2020) The team implemented a CenterNet-based model with the PASCAL VOC pretrained ResNet50 backbone for the object detection task. A similar backbone with U-Net was devised for the segmentation task. The dataset was randomly duplicated to tackle class-imbalance. To improve generalization performance, each image was augmented 86 times by randomly choosing augmentation techniques from the pool of rotation, flipping, contrast enhancement and brightness adjustment.
- **Team drvelmuruganb** (Balasubramanian et al., 2020) For the segmentation of disease classes the team used an encoder-decoder based SUMNet architecture with the ImageNet pretrained VGG11 backbone. The authors also applied several augmentation strategies including variable brightness and HSV values, multiple crops and geometric transformations such as rotation, affine, scaling and projective were also applied to improve the accuracy.

5. Results

For the EAD2020 sub-challenge, we present the results of 12 participating teams for multi-class artefact detection task and 8 teams for segmentation task. Similarly, for EDD2020 sub-challenge, we have included top 6 teams for detection and 7 teams for segmentation of multi-class diseases. In this section we present the quantitative and qualitative results for each team based on the evaluation metrics discussed in Section 3.2. For the EAD2020 sub-challenge, 3 different test dataset were released: 1) single-frame data for detection and segmentation, 2) sequence dataset for

detection only and 3) out-of-sample data for generalization task only. For the detection task, the average of the aggregated sum of the detection scores for the single frame data and the sequence data were considered for final scoring. While, for the EDD2020 challenge only single frame detection and segmentation data were released. Below we present the result for each sub-challenges separately.

5.1. Quantitative results

5.1.1. EAD2020 Sub-challenge

In this section, the results of the participant teams in the EAD2020 challenge to detect and segment artefacts are presented.

Detection task for EAD2020

Table 5 and Table 6 present the mAP values computed at different IoU thresholds (i.e., 25%, 50%, and 75%), overall mAP, overall IoU, and the final score for the detection of the artefacts from single frame and sequence data, respectively. Additionally, we also provide results of baseline methods that include YOLOv3 and RetinaNet with darknet53 and ResNet101 backbones, respectively. In Table 5 (i.e., single frame detection), it can be observed that the team *polatgorkem* that implemented ensemble technique with Cascaded RCNN, Faster-RCNN and RetinaNet surpassed the other teams by achieving the highest final score on the leaderboard (score_d , Eq. 5) of 25.123 ± 7.124 with the best overall mIoU of 36.579 providing a high overlap ratio between the generated bounding box with ground truth per frame. The method proposed by the team *arnavchavan04* comes in the second place with score_d of 24.079 ± 9.342 with 9% more mAP than the winning team but large sacrifice in the mean IoU. Similarly, for sequence data in Table 6, team *polatgorkem* maintained the first position with a final score of 25.529 ± 10.326 . While the second scorer team *VinBDI* suggested a method that obtained a better balanced between mAP and mIoU scores.

Furthermore, Table 7 shows the overall ranking for the teams in terms of Score (R_{score_d}), mAP (R_{mAP}), and generalizability performance (R_g) in addition to, mAP_d , mAP_{seq} , score_d , mAP_g and dev_g . The baseline RetinaNet recorded the least deviation but also the least mAPs. On considering the mAP_g and dev_g together for the final ranking of the generalization task, teams *VinBDI* and *StarStarG* secured the first place. On observing at the class-wise performance in Fig. 5 (a) (i.e., single frame), it can be seen that there was a high detection score (score_d) and AP for larger artefact instances such as saturation and contrast. Similarly, most of the teams had a high IoU with the ground truth when detecting the instrument class. On the other hand, the detection and localization of smaller artefact instances such as bubble and saturation showed the degraded performances by all the participating teams and by the baseline methods.

Segmentation task for EAD2020

Table 8 presents the JC, DSC, F2, PPV, recall, and accuracy obtained by each team and baseline methods. As shown, the method proposed by team *arnavchavan04* and team *VinBDI* had the best performance in terms of JC ($> 62\%$), DSC ($> 67\%$), F2 ($> 67\%$) and PPV ($> 80\%$) proving the ability to segment less false positive regions. However, the method suggested by team *qzheng5* and team *DuyHUYNH* segmented more true positive regions compared to other teams obtaining top recall values of 0.8352 and 0.828. The baseline methods showed a low performance in terms of final score compared to the methods proposed by the participants. Furthermore, Fig. 6 (a) shows class-wise scores for DSC, PPV and Recall. Similar to detection, segmenting larger instances like the saturation and the instrument obtained the high scores. Specularity, bubble and the artefact classes were among least performing classes for many teams and baseline methods.

Table 5

EAD2020 results for the detection task on the single frame dataset. mAP at IoU thresholds 25%, 50% and 75% are provided along with overall mAP and overall IoU computations. Overall scores are computed at 11 IoU thresholds and averaged. Weighted detection score $score_d$ is computed between overall mAP and IoU scores only. Three best scores for each metric criteria are in bold.

Team names	mAP ₂₅	mAP ₅₀	mAP ₇₅	overall mAP _d	overall mIoU _d	mAP _δ	score _d ± δ
polatgorkem	26.886	17.883	5.608	17.486	36.579	7.124	25.123 ± 7.124
qzheng5	33.134	20.084	5.570	19.720	27.185	8.820	22.706 ± 8.820
xiahong1	30.627	19.384	4.935	18.512	26.388	8.428	21.663 ± 8.428
mathew666	20.360	19.440	7.783	18.091	32.692	5.617	23.931 ± 5.617
VinBDI	38.429	25.426	7.053	24.069	12.644	10.291	19.499 ± 10.291
higersky	36.920	25.770	9.452	24.771	17.298	8.707	21.781 ± 8.707
StarStarG	41.800	29.984	10.733	28.380	16.250	10.042	23.528 ± 10.042
anand_subu	29.755	19.893	5.271	18.886	24.029	7.619	20.943 ± 7.619
arnavchavan04	38.752	27.247	9.858	26.021	21.165	9.342	24.079 ± 9.342
MXV	25.373	18.967	7.171	17.82	28.056	5.754	21.914 ± 5.754
mimykcp	39.897	26.296	6.839	25.082	10.209	10.765	19.133 ± 10.765
DuyHUYNH	20.512	12.234	2.978	11.894	27.063	5.671	17.962 ± 5.671
baselines							
YOLOv3	22.798	13.736	2.804	13.249	24.883	6.525	17.903 ± 6.525
RetinaNet\$ResNet101)	15.270	8.927	2.061	8.754	23.202	4.275	14.533 ± 4.275

Table 6

EAD2020 results for the sequence dataset. mAP at IoU thresholds 25%, 50% and 75% are provided along with overall mAP and overall IoU computations. Overall scores are averaged with 11 IoU thresholds. Weighted detection score $score_d$ is computed between overall mAP and IoU scores only. Three best scores for each metric criteria are in bold.

Team names	mAP ₂₅	mAP ₅₀	mAP ₇₅	overall mAP _{seq}	overall mIoU _{seq}	mAP _δ	score _d ± δ
polatgorkem	38.464	24.803	4.138	23.137	29.117	10.326	25.529 ± 10.326
qzheng5	48.210	25.717	3.997	25.665	20.949	14.222	23.779 ± 14.222
xiahong1	46.087	25.813	2.684	25.136	18.398	15.128	22.441 ± 15.128
mathew666	31.599	21.878	3.053	19.623	20.858	9.718	20.117 ± 9.718
VinBDI	45.295	26.723	4.396	25.285	23.426	13.972	24.542 ± 13.972
higersky	47.716	29.841	4.473	28.334	12.865	14.579	22.147 ± 14.579
StarStarG	46.965	30.202	5.432	28.107	8.371	13.367	20.213 ± 13.367
anand_subu	38.352	25.535	3.843	23.014	20.703	10.859	22.089 ± 10.859
arnavchavan04	34.511	21.524	4.886	20.700	11.827	9.839	17.151 ± 9.839
MXV	31.391	19.838	3.620	18.601	21.504	8.688	19.762 ± 8.688
mimykcp	44.972	26.780	4.400	25.937	6.892	13.697	18.319 ± 13.697
DuyHUYNH	28.632	15.524	0.815	15.468	16.968	9.381	16.068 ± 9.381
baselines							
YOLOv3	32.199	18.473	1.137	17.176	16.351	10.596	16.846 ± 10.596
RetinaNet\$ResNet101)	17.646	6.447	0.767	8.079	10.000	5.151	9.252 ± 5.151

Table 7

EAD2020 team ranking based on different metric criteria for detection and generalization task. Overall mAPs (mAP_d and mAP_{seq}) computed on single frame and sequence data are averaged. Final score_d is then computed as the weighted value between the final IoU_d and the averaged mAP. Rankings for each metric are also provided based on ascending order of the scores except for deviation score for out-of-sample data. Three best scores for each metric criteria are in bold.

Team Names	mAP _d	mAP _{seq}	final IoU	final score _d	mAP _g	dev _g	R _{score_d}	R _{mAP}	R _{gen}
polatgorkem	17.486	23.137	32.848	25.326	21.008	9.359	1	9	6
qzheng5	19.720	24.174	23.751	22.668	23.749	8.522	2	6	5
xiahong1	18.512	25.136	22.393	22.051	24.579	8.169	3	7	3
mathew666	18.091	19.651	26.783	22.035	16.714	5.674	4	10	4
VinBDI	24.069	25.282	18.033	22.018	24.140	5.607	5	4	1
higersky	24.771	28.252	15.061	21.931	24.850	7.686	6	2	2
StarStarG	28.380	28.107	12.311	21.870	25.340	7.537	7	1	1
anand_subu	18.886	23.004	22.359	21.510	20.203	7.896	8	8	5
arnavchavan04	26.021	20.700	16.496	20.614	21.138	6.968	10	5	3
MXV	17.820	18.597	24.779	20.836	17.294	6.077	9	11	4
mimykcp	25.082	25.843	8.536	18.691	23.929	7.999	11	3	4
DuyHUYNH	11.894	15.468	22.016	17.015	11.304	4.807	13	13	4
baselines									
YOLOv3	13.249	17.176	20.617	17.374	15.456	4.397	12	12	3
RetinaNet (ResNet101)	8.754	8.079	16.601	11.690	7.763	1.985	14	14	3

5.1.2. EDD2020 Sub-challenge

In this section, we report the performance of the participating teams in the EDD2020 challenge for the detection and segmentation.

Detection task for EDD2020

In Table 9, the team *adrian* achieved the highest score among other participants and the baseline methods with a final score_d

of 33.602 ± 8.523 with the highest overall mAP (37.594) and the second highest overall mIoU (27.614). The best localization score was obtained by the team *sahadate* but with nearly 5% lower mAP than the top scorer team. Furthermore, the baseline method *RetinaNet* with the ResNet101 backbone performed better than most of the participating teams. From Table 10, it is evident that most teams and baselines failed to detect suspicious class instance while

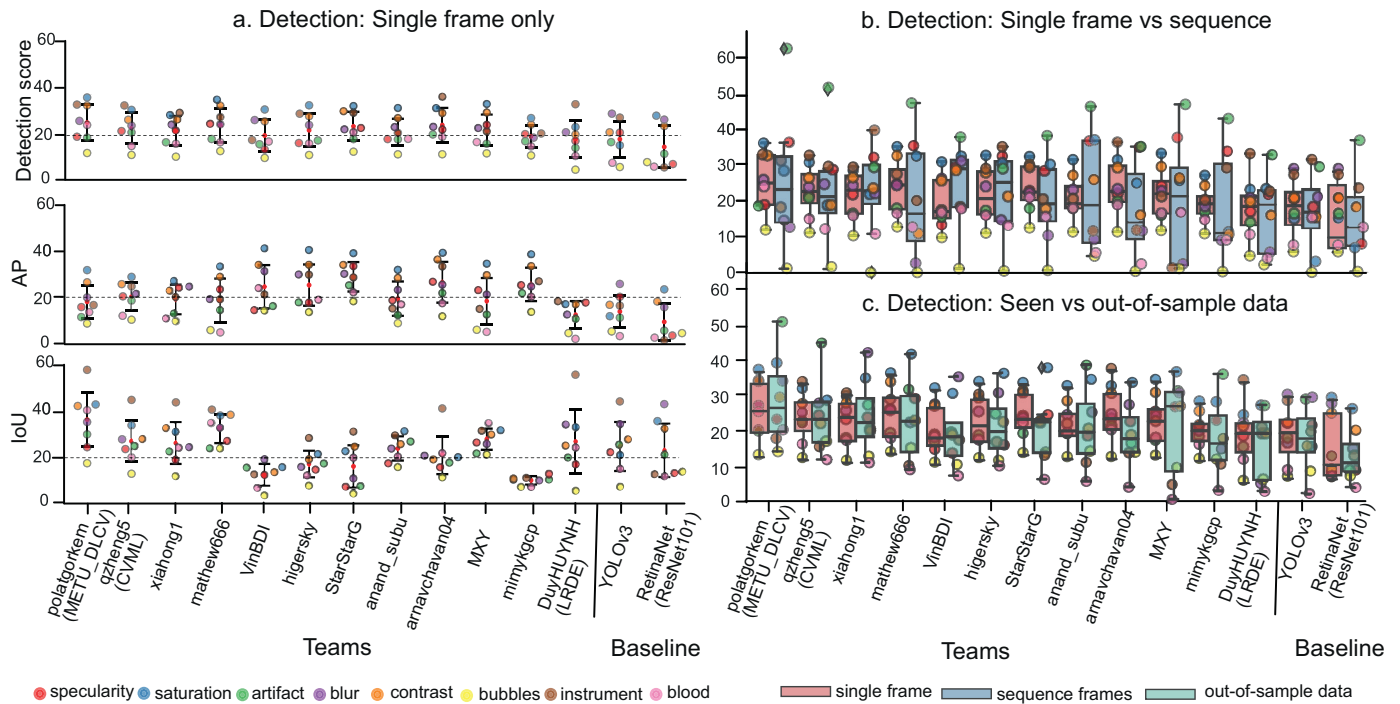


Fig. 5. Detection and out-of-sample generalization tasks for EAD2020 sub-challenge. a) Error bars and swarm plots for the intersection over union (IoU, top), average precision (AP, middle) and challenge detection score (mAP_d, bottom) for each team is presented on 237 single frame test data. b-c) Comparison of mAP_d w.r.t. mAP_{seq} (mAP on sequence test data with 80 frames) and mAP_{oos} (mAP on out-of-sample data 99 frames) are provided. a-c) On the right, results from baseline detection methods: YOLOv3 and RetinaNet (with ResNet101 backbone) are also presented. Teams are arranged by decreasing overall detection ranking R_{score_d} (see Table 7).

Table 8

Evaluation of the artefact segmentation task. Top three best scores for each metric criteria are in bold.

Team Names	JC	DSC	F2	PPV	Rec	Acc	Score _s	R _{score_s}
qzheng5	0.477	0.532	0.561	0.556	0.835	0.973	0.621	8
VinBDI	0.628	0.673	0.670	0.837	0.738	0.978	0.730	2
higersky	0.529	0.579	0.587	0.675	0.758	0.975	0.650	5
anand_subu	0.304	0.354	0.361	0.430	0.747	0.975	0.473	14
arnavchavan04	0.622	0.673	0.683	0.800	0.767	0.977	0.731	1
DuyHUYNH	0.502	0.557	0.583	0.593	0.829	0.974	0.640	6
mimykgcp	0.531	0.576	0.579	0.723	0.726	0.977	0.651	4
mouradai_ox	0.581	0.632	0.647	0.711	0.800	0.974	0.697	3
baselines								
FCN8	0.500	0.548	0.550	0.670	0.708	0.976	0.619	9
UNet-ResNet34	0.310	0.364	0.373	0.419	0.766	0.974	0.481	13
PSPNet	0.497	0.541	0.534	0.698	0.680	0.975	0.613	10
DeepLabv3 (ResNet50)	0.448	0.495	0.492	0.599	0.704	0.974	0.572	12
DeepLabv3+ (ResNet50)	0.485	0.533	0.535	0.646	0.726	0.976	0.610	11
DeepLabv3+ (ResNet101)	0.501	0.547	0.546	0.683	0.718	0.973	0.624	7

Table 9

EDD2020 results for the detection task on the single frame dataset. mAP at IoU thresholds 25%, 50% and 75% are provided along with overall mAP and overall IoU computations. Overall scores are computed at 11 IoU thresholds and averaged. Weighted detection score $score_d$ is computed between overall mAP and IoU scores only. Three best scores for each metric criteria are in bold.

Team names	mAP ₂₅	mAP ₅₀	mAP ₇₅	overall mAP _d	overall mIoU _d	mAP _δ	score _d ± δ
adrian	48.402	33.562	27.098	37.594	27.614	8.523	33.602 ± 8.523
sahadate	37.612	23.284	15.837	26.834	32.420	8.325	29.068 ± 8.325
VinBDI	43.202	26.981	17.001	30.219	17.773	9.478	25.241 ± 9.478
YHChoi	23.183	11.082	8.800	15.783	24.623	6.216	19.319 ± 6.216
DuyHUYNH	23.959	9.587	5.659	12.479	13.829	6.284	13.019 ± 6.284
mimykgcp	34.884	20.982	4.463	20.742	2.270	9.359	13.353 ± 9.359
drvelmuruganb	31.018	18.421	11.768	21.790	7.322	7.424	16.002 ± 7.424
baselines							
YOLOv3	34.305	21.227	14.650	22.980	24.351	6.456	23.528 ± 6.456
RetinaNet (ResNet50)	26.833	14.441	9.907	17.552	25.580	6.464	20.763 ± 6.464
RetinaNet (ResNet101)	42.579	27.000	11.194	27.974	26.434	11.949	27.358 ± 11.949

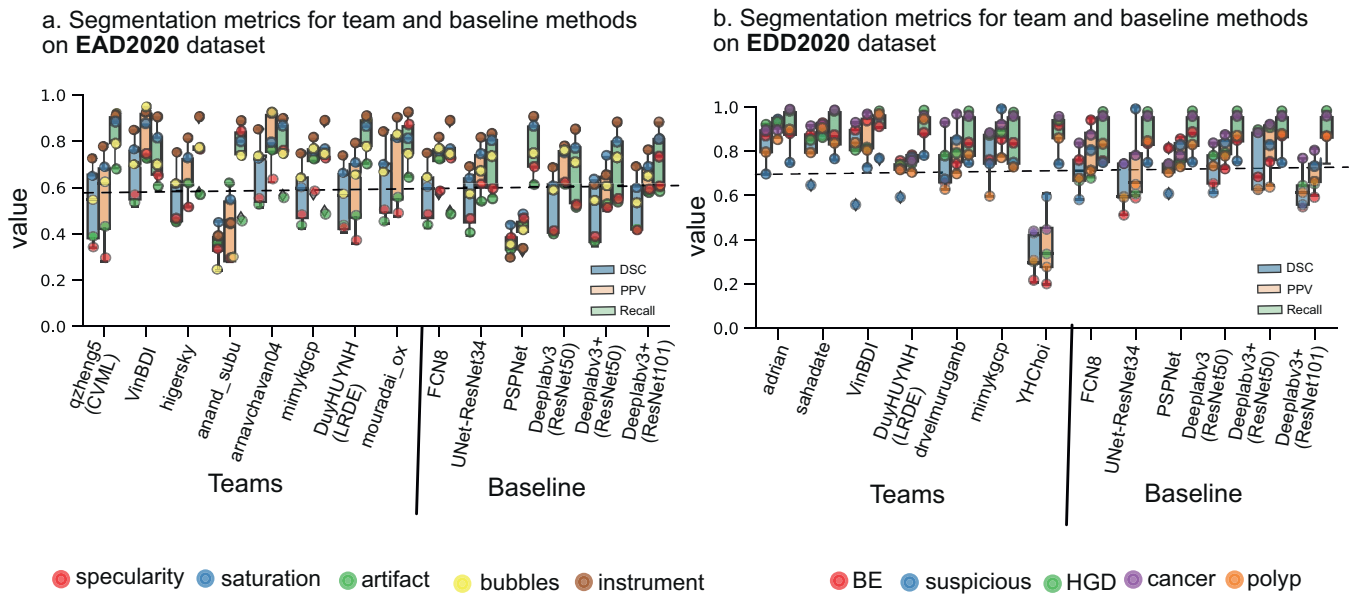


Fig. 6. Semantic segmentation for EAD and EDD sub-challenges: Error bars with overlaid swarm plots for dice similarity coefficient (DSC), positive predictive value (PPV) or precision and recall are presented for each team and baseline methods for the EAD2020 (a) and EDD2020 (b) challenges. 6 different baseline methods are also provided for comparison.

Table 10

Per class evaluation results for the detection task of the EDD2020 sub-challenge.

Teams EDD2020	NDBE	suspicious	HGD	cancer	polyp	δ
adrian	28.911	1.776	32.727	64.286	60.269	22.841
sahadate	46.193	1.099	22.727	10.000	54.152	20.414
VinBDI	48.489	3.497	25.852	10.000	63.260	22.660
YHChoi	26.900	0.000	22.727	0.000	29.289	13.057
DuyHUYNH	20.281	1.499	11.364	0.000	29.254	11.134
mimykqcp	50.089	4.592	23.064	5.852	20.112	16.429
drvelmuruganb	34.775	0.000	22.727	0.000	51.446	19.993
baselines						
YOLOv3 (darknet53)	38.839	0.000	6.970	16.667	52.426	19.712
RetinaNet (ResNet50)	23.636	0.000	18.182	0.000	45.943	17.086
RetinaNet (ResNet101)	29.483	0.000	22.727	31.818	55.840	17.909

Table 11

Evaluation of the disease segmentation methods proposed by the participating teams and the baseline methods. Top three evaluation criteria are highlighted in bold.

Team Names	JC	DSC	F2	PPV	Rec	Acc	Score _s	R _{score_s}
adrian	0.820	0.836	0.842	0.921	0.894	0.955	0.873	1
sahadate	0.797	0.816	0.819	0.906	0.883	0.955	0.856	2
VinBDI	0.788	0.805	0.812	0.859	0.912	0.952	0.847	3
DuyHUYNH	0.6843	0.7058	0.718	0.762	0.905	0.931	0.773	9
drvelmuruganb	0.7166	0.7349	0.734	0.819	0.857	0.959	0.786	6
mimykqcp	0.7561	0.7721	0.770	0.893	0.845	0.957	0.820	4
YHChoi	0.314	0.340	0.356	0.385	0.896	0.892	0.494	13
baselines								
FCN8	0.687	0.705	0.709	0.811	0.850	0.953	0.769	10
UNet-ResNet34	0.617	0.637	0.638	0.732	0.868	0.958	0.719	11
pspnet	0.698	0.721	0.723	0.797	0.876	0.959	0.779	8
DeepLabv3 (RetinaNet50)	0.704	0.724	0.724	0.810	0.878	0.962	0.784	7
DeepLabv3+ (RetinaNet50)	0.725	0.744	0.749	0.818	0.882	0.960	0.798	5
DeepLabv3+ (RetinaNet101)	0.608	0.627	0.629	0.698	0.880	0.962	0.709	12

most teams performed comparatively better on polyp and NDBE classes. Only the winning team *adrian* and RetinaNet (ResNet101) provided a descent score for cancer class with most teams recording mAP below 10. For HGD class category, top performing teams were *adrian* and *VinBDI* with mAP over 25.

Segmentation task for EDD2020

From Table 11, it can be observed that the three teams (*Adrian*, *sahadate* and *nhanthanhnghuyen94*) achieved a DSC over 0.80. More-

over, they maintained the high performance for other metrics as well that include JC (>0.78), F2 (>0.81), and PPV (>0.85) securing first, second and third ranks, respectively. Teams *VinBDI* and *DuyHUYNH* were able to segment more true positive regions reaching the top recall values. Fig. 6(b) represents per-class metric values. It can be observed that unlike detection task, most teams reported high performance for cancer class. Also, most teams showed higher DSC, PPV and recall for BE class instance as well (> 0.8 for top

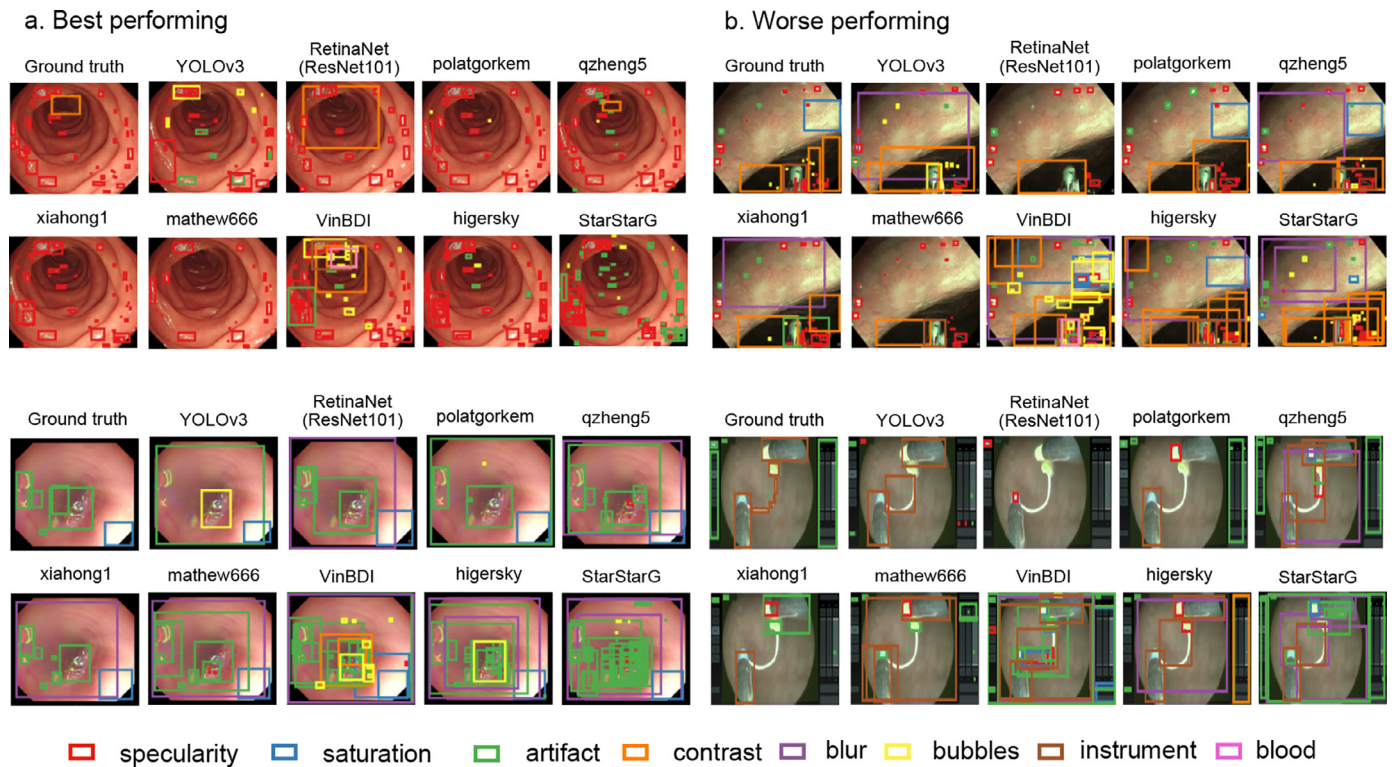


Fig. 7. EAD2020 best and worse performing samples for the detection task. a) Best performing samples for 6 top ranked team results. b) Worse performing samples for the same teams in (a). Results with baseline methods are also included together with ground truth sample.

three teams). However, similar to the detection task, most team and baseline methods reported least values for the suspicious class.

5.2. Qualitative results

Detection task

Fig. 7 shows the best (panel a) and the worse (panel b) performing frames from single frame dataset for EAD2020. It can be observed that specularity and artefacts are detected and well localized by top teams (see Fig. 7 a). Similarly, in the bottom example, saturation is also detected by all the participants. Even though, blur is not present for this sample, most methods also detected it. While for the worse performing frame (see Fig. 7 b), instrument class is confused with contrast or artefact on the top sample, while in the bottom sample instrument is detected by some teams but often either detected only partially or overlapped by different classes such as saturation or artefact.

For out-of-sample generalization task, it can be seen in Fig. 8 (a) that besides YOLOv3 baseline method, all the baselines and teams detected saturation class. While some teams (*mathew666*, *VinBDI*, *higersky*) detected multiple bounding boxes for the same class, they also detected blur class for this frame. While for worse performing frame (see Fig. 8(b)), instrument class (at the center of the image) is well localized only by the team *xiahong1* while most teams either partially detected the instrument (e.g., team *qzheng5*) or could not detect the instrument class at all (e.g., team *polatgorkem*). In both cases, the three teams *VinBDI*, *higersky* and *StarStarG* produced multiple overlapping and different size bounding boxes.

Qualitative results for the EDD2020 challenge is shown in Fig. 9. The best performing samples in Fig. 9(a) shows polyp class (at the top); non-dysplastic Barrett's esophagus (NDBE) and suspicious classes on the bottom. It can be observed that polyp class is detected and well localized by all the teams and baseline methods. However, for bottom row NDBE is detected by most of the meth-

ods while confusion is observed across the suspicious class with high-grade dysplasia (HGD) class. Team *mimykgep* produced numerous bounding boxes failing to optimally localize adherent disease classes. For the worse performing frames (Fig. 9(b)), cancer class (top) in the ground truth is confused with the polyp class instance for most of the teams and the baseline methods. While, for the NDBE class in the bottom of Fig. 9(b), teams were either not able to detect the NDBE class (except team *adrian*, team *YH-Choi* and YOLOv3) at all or partially detected the NDBE areas (e.g., teams *VinBDI* and *drvvelmuruganb*). Again, for the presented case, team *mimykgep* detected numerous bounding boxes.

Segmentation task

Endoscopic artefact segmentation samples representing best and worse performing teams is provided in Fig. 10. For the sample with only the instrument class (see Fig. 10 a, top panel) it can be observed that almost all the baseline and teams were able to predict precise delineation of the instrument class. Similarly, in the bottom panel of Fig. 10(a), specularity, saturation and artefact classes were segmented well by most of the teams and baseline methods. Even though, a single instrument class is present in the sample image in Fig. 10(b), none of the methods were able to segment the instrument. Also, for the bottom panel in the Fig. 10(b), specularity areas were segmented well by the teams *mouradaiox* and *mimykgep*. However, saturation area was under segmented by most of the teams and baseline methods. Fig. 11 (a) represents the polyp class (at the top); NDBE and suspicious classes (at the bottom). It can be observed that polyp is segmented well by all the baselines and most teams (except team *drvvelmuruganb* who misclassified the pixels to suspicious class). While, most teams and baselines were able to precisely delineate NDBE class for the frame in the bottom panel but missed suspicious area. In the worse performing sample (see Fig. 11(b)), most teams were able to segment NDBE area but large HGD area was missed by all the teams. Also, some teams confused HGD area with suspicious class. For the bot-

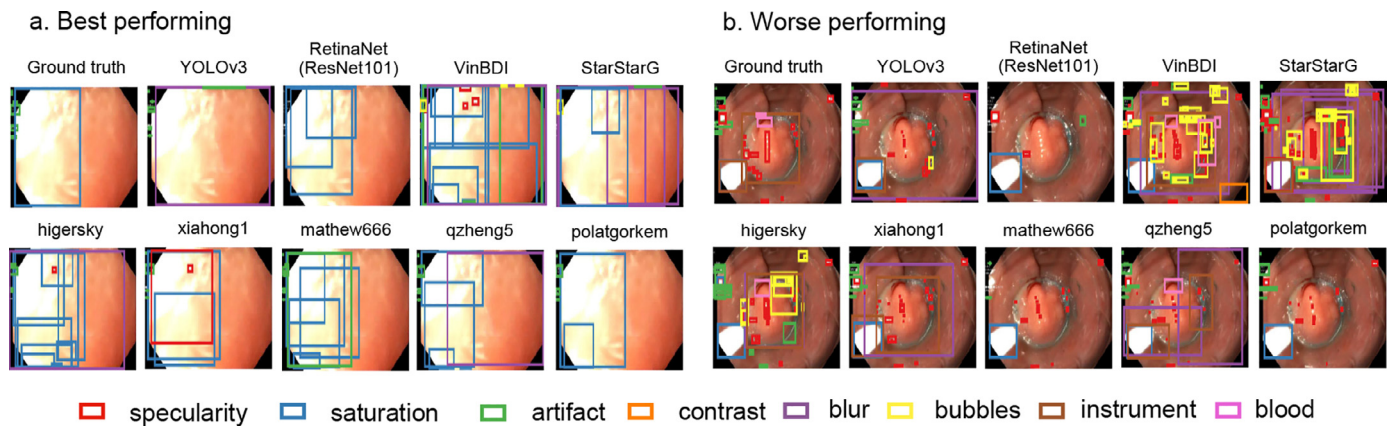


Fig. 8. EAD2020 best and worse performing samples for the generalization task. a) Best performing samples for 7 top ranked team results. b) Worse performing samples for the same teams in (a). Results with baseline methods are also included together with ground truth sample.

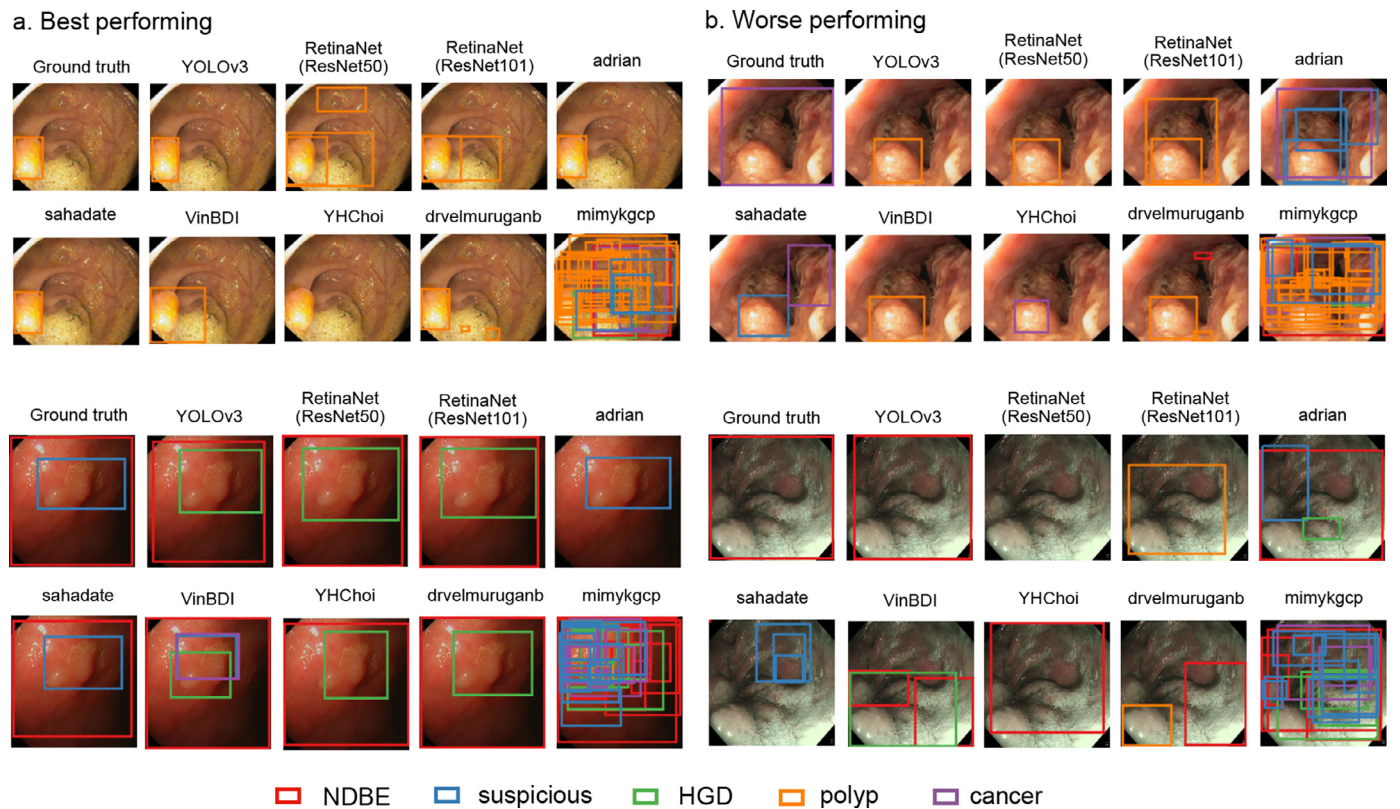


Fig. 9. EDD2020 best and worse performing samples for the detection task. a) Best performing samples for 6 top ranked team results. b) Worse performing samples for the same teams in (a). Results with baseline methods are also included together with ground truth sample.

tom panel in Fig. 11(b), instead of suspicious class present in the ground truth, almost all the teams detected this as polyp or cancer. However, the region delineation was close to the ground truth for most teams.

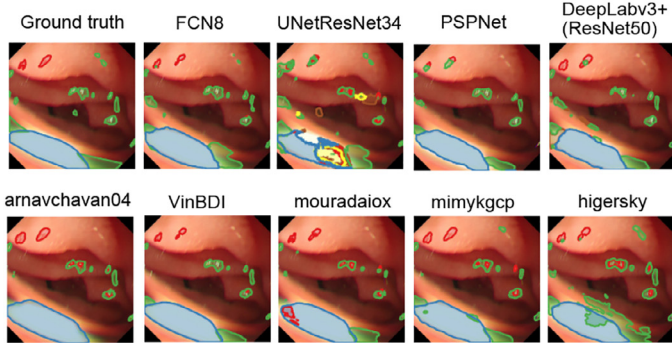
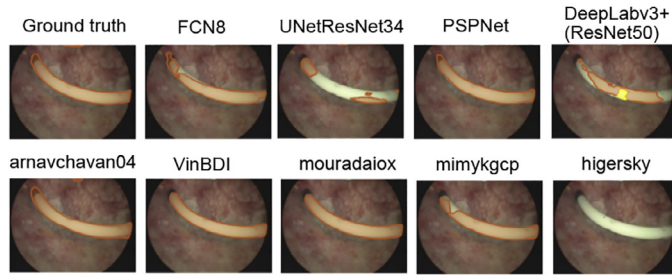
6. Discussion

Deep learning methods are rapidly being translated for the use of computer aided detection (CADE) and diagnosis (CADx) of diseases in complex clinical settings including endoscopy. However, the amount of data variability particularly in endoscopy is significantly higher than in natural scenes which possess a significant challenge in the process. It is therefore vital to determine an effective translational pathway in endoscopy. Majority of challenges in endoscopy are due to its complex surveillance that lead to se-

vere artefacts that may confuse with disease. Similarly, a system designed for a particular organ may not generalize to be used in the other.

Most deep learning methods that were used in the EndoCV2020 challenge can be categorised into multiscale, symbiotic, ensemble, encoder-decoder and cascading nature, or a combination of these (see Table 3 and Table 4). Fig. 12 presents the overview of the used methods for the detection (a) and segmentation (b) challenge tasks based on the architecture usage. It can be observed that the majority of detection methods used two-stage Faster-RCNN with 4/7 teams combining it with one-stage RetinaNet or YOLOv3 or a combination of all. Cascade R-CNN which is built upon Faster R-CNN cascaded architecture was exploited by 4 teams. Similarly, U-Net-based architectures were utilised by most teams for semantic

a. Best performing



□ specularity □ saturation □ artifact □ contrast □ blur □ bubbles □ instrument □ blood

b. Worse performing

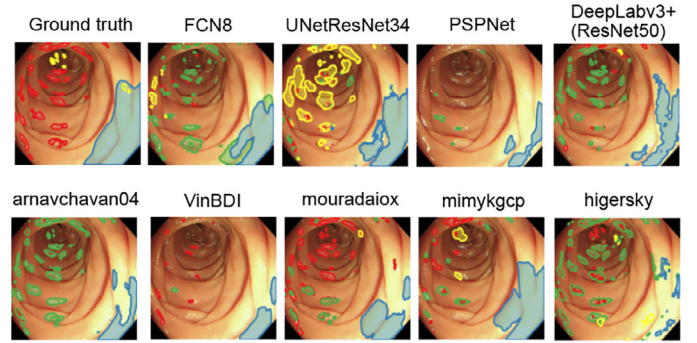
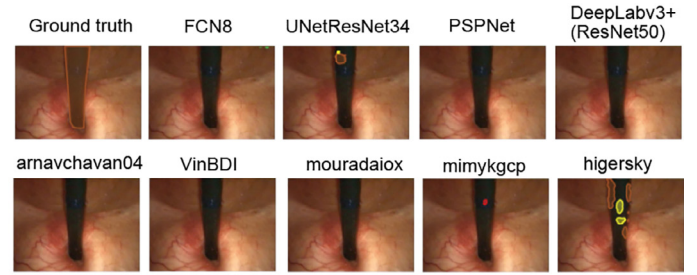
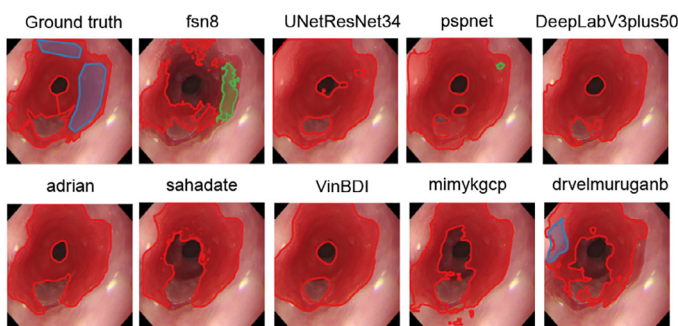
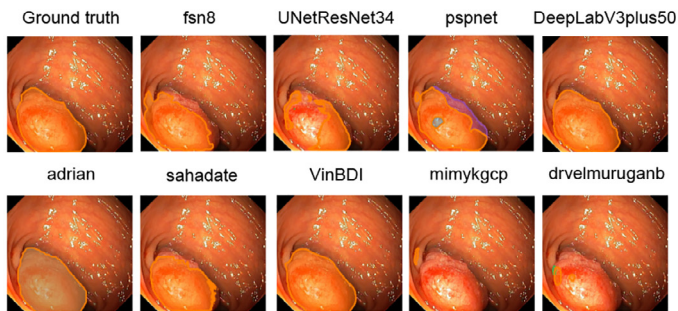


Fig. 10. EAD2020 best and worse performing samples. a) Best performing samples for 5 top ranked team results. b) Worse performing samples for the same teams in (a). Results with baseline methods are also included together with ground truth sample (top). Single class samples are chosen at the top and multi-class samples are at the bottom in each category.

a. Best performing



□ NDBE □ suspicious □ HGD □ polyp □ cancer

b. Worse performing

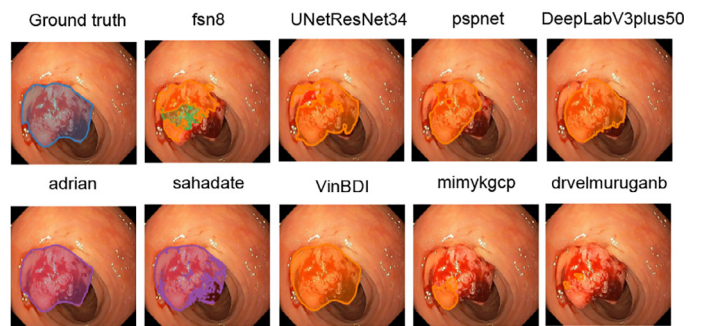
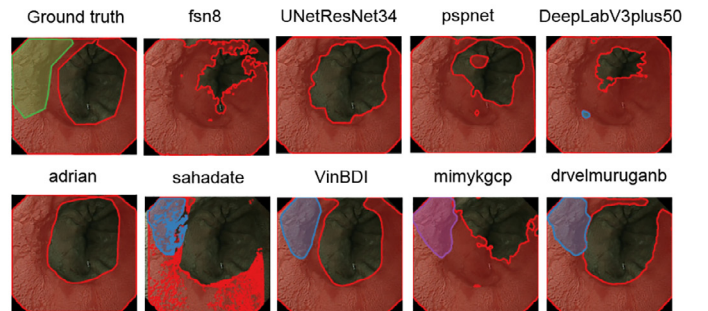


Fig. 11. EDD2020 best and worse performing samples. a) Best performing samples for 5 top team results. b) Worse performing samples for the same teams in (a). Results with baseline methods are also included together with ground truth sample (top).

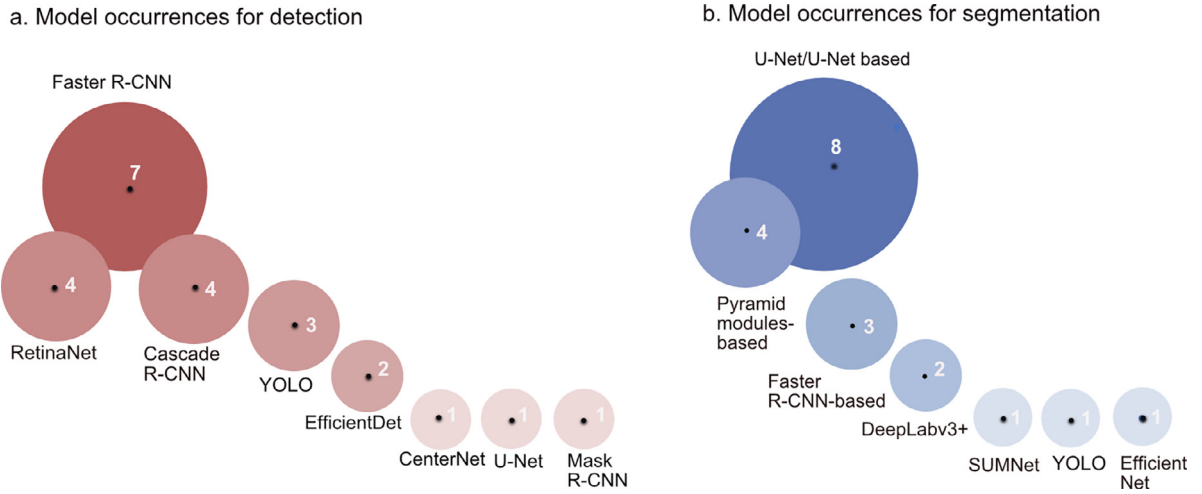


Fig. 12. EndoCV2020 method categories in blob-representation. Model occurrences are presented for detection (a) and segmentation (b) tasks for both EAD2020 and EDD2020 sub-challenges. The number of occurrences is provided inside each blob.

segmentation task with 4 teams exploring pyramid module-based architectures and 2 teams used Deeplabv3+ architecture. Faster RCNN-based model was also explored with additional thresholding (e.g., team *adrian*) or per pixel prediction heads (e.g., team *sahadate*). Even though similar techniques were used in EAD2019 challenge (Ali et al., 2020c), a direct comparison is not possible. This is due to the inclusion of more data for EAD2020 in both train and test sets. Also, EAD2020 includes sequence data which was not provided in EAD2019 challenge.

For the detection task, the top performing teams on the challenge metric in both EAD (team *polatgorkem*) and EDD (team *adrian*) were those using ensemble networks, i.e., maneuvering outputs from multiple architectures. However, these networks sacrifice the speed of detection which can be observed from the computational time which were significantly higher than teams that used a single architecture (see Table 7 and Table 9). Other teams that used such an approach included team *arnavchavan04* and *mimykqcp* who combined Faster R-CNN with RetinaNet but both teams were respectively on 10th and 11th ranking. Just using Faster R-CNN alone with ResNet101 backbone, teams *qzhang5* and *matthew666* were able to detect both small and large size bounding boxes with sub-optimal accuracy that put them at 2nd and 4th positions, respectively. Similarly, team *sahadate* claimed 2nd position on EDD detection task using Mask R-CNN which is based on the Faster R-CNN architecture. For EAD2019 challenge (Ali et al., 2020c), team *yangsuhui* also used an ensemble network with Cascade RCNN and FPN approach for the detection task similar to the EAD2020 top scorer team *polatgorkem*.

An intelligent choice for improved speed and accuracy using a scalable network was presented by the teams *xiahong1* (used YOLACT) and *VinBDI* (used EfficientDet D0) which were placed 3rd and 5th, respectively, on the final detection score of the EAD2020. On the sequence data, team *VinBDI* was the 2nd best method demonstrating the reliability of the used EfficientNet and FPN architectures. However, for almost all team methods the standard deviation was higher than for single frame data. No team exploited the sequence data provided for training. Team *VinBDI* was also ranked 3rd on the EDD detection task. Teams *higerssky*, *StarStarG* and *MXy* that used cascaded R-CNN were ranked respectively on 6th, 7th and 9th positions. Additionally, the team *StarStarG* was ranked 1st and team *higerssky* was ranked 2nd on the overall mAP. However, it is to be noted that taking only mAP scores into account for detection could lead to over detection of the bounding boxes that increases the chance of finding a particular class but at

the same time weakens the localization capability of the algorithm (see Fig. 7). Similar observations were found for the EDD dataset where the team *mimykqcp* obtained an overall mAP of 20.742 but only 2.270 for the overall IoU (see Table 9). As a result, over detection of the bounding boxes can be seen in Fig. 9. In order to deal with the over detection of the bounding boxes, YOLACT architecture used by *xiahong1* suppressed the duplicate detections using already-removed detections in parallel (*fast NMS*). Similarly, teams such as *polatgorkem* from the EAD and *adrian* from the EDD were able to eliminate the duplicate detections using ensemble network and a class agnostic NMS.

Hypothesis I: In the presence of multiple class objects, object detection methods may fail to precisely regress the bounding boxes. Methods need better penalization on the bounding box regression or a technique to perform effective non-maximal suppression.

The choice of networks from each team depended on their ambition of either obtaining very high accuracy without focusing on speed or a trade-off between the speed and the accuracy or focusing on both and thinking out-of-the-box to use more recent developed methods which beats faster networks (such as YOLOv3) that included EfficientDet D0 architecture used by the team *VinBDI* (see Table 3). Due to the efficiency of the EfficientDet D0 network that used biFPN and efficientNet backbone, team *VinBDI* achieved second least deviation in mAP (i.e., $dev_g = 5.607$) with competitive mAP_g ($= 24.140$) and won the generalization task together with the team *StarStarG* who had slightly higher mAP_g ($= 25.340$) but larger mAP deviation between detection and generalization datasets. Most methods for the detection task on both the EAD and EDD dataset performed better than the baseline one-stage methods (YOLOv3 and RetinaNet). However, it was found that even though team *polatgorkem* won the detection task, the method failed on generalization data where the team was ranked only last. The main reason behind this could be because the generalization gap mAP_g was estimated between two mAP's (mAP_d and mAP_g) and not IoU. Also, the final ranking was done taking into account the rank of dev_g and mAP_g only. It can be observed in Fig. 8 that the bounding box localization of team *polatgorkem* is precise in (a) while it misses instrument area at the center in (b). However, the winning teams *VinBDI* and *StarStarG* both over detect the boxes. The generalization ability of the methods were not explored for EDD dataset.

Hypothesis II: Metrics are critical but using a single metric does not always gives the right answer. Weighted metrics are desired in object detection task to establish a good trade-off between detection and precise localization.

A major problem in the detection of EDD dataset was class confusion mostly for suspicious, HGD and cancer classes. This could be because of smaller number of samples for each of these classes compared to NDBE and polyp (see Fig. 3). While most methods were able to detect and localize NDBE and polyp class in general (3/7 teams with an overall mAP > 45 and 4/7 teams with > 50), all teams failed in suspicious class (overall mAP < 5.0) and most teams for cancer class (overall mAP < 15.0) (see Table 10). Fig. 9 shows that polyp is detected and localized very well by most teams (a, top). Similarly, NDBE is localized by most methods, however, in this case suspicious class is confused mostly with the HGD. Also, in Fig. 9(b, top), it can be observed that the cancer class instance is confused with mostly polyp class.

Hypothesis III: Detection bounding boxes confuse with classes that have similar morphology and smaller number of samples failing to learn the contextual features. To improve detection, such samples need to be identified and more data demonstrating such attributes need to be injected (both positive and negative samples).

Similar to the detection task, teams that used ensemble techniques were among the best performing teams for the segmentation task. Teams *arnavchavan04* and *VinBDI* secured first ($score_s = 0.731$) and second ($score_s = 0.730$) positions, respectively, on the EAD2020 segmentation task (see Table 8) and the team *adrian* won the EDD2020 segmentation task challenge with $score_s$ of 0.873 (see Table 11). The team *arnavchavan04* used multiple augmentation techniques including cutmix and a feature pyramid network with a combination of EfficientNet backbones from B3 to B5. Similarly, team *VinBDI* ensembled a U-Net architecture with EfficientNet B4 and BiFPN network with ResNet50 backbone. Compared to EAD2019 where the winning team *yangsuhui* used DeepLabV3+ model with two different backbones, both of the top scorer teams of 2020 revealed the strength of recent EfficientNet and FPN-based segmentation approaches.

In the EDD2020 segmentation task, the team *adrian* combined predictions from three object detection architectures where the YOLOv3 and Faster R-CNN class predictions were used to correct the instance segmentation masks from Cascade R-CNN. A direct instance segmentation approach used by the team *sahadate* secured second position ($score_s = 0.856$) on the same while ensemble network of the team *VinBDI* secured the third position ($score_s = 0.847$). Direct usage of a single existing state-of-the-art methods utilising different augmentation techniques (e.g., *DuyHUYNH*) or different backbones (e.g., *mimykqcp*, *qzheng5*) resulted in improved results compared to the original baseline methods, however, much lower than the top performing methods (see Table 8 and Table 11).

Hypothesis IV: The choice of combinatorial networks that well synthesizes width, depth and resolution to capture optimal receptive field, and a domain agnostic knowledge transfer mechanism are critical to tackle heterogeneous (multi-center and variable size) multi-class object segmentation task.

From Fig. 6 it can be observed that the top three performing teams of the EAD2020 segmentation task (*arnavchavan04*, *VinBDI*, *mouradai_ox*) has high DSC value (0.538, 0.548 and 0.492 respectively) compared to most methods for the specularity class instance. It is to be noted that the specularities are often confused with either artefact or bubbles which makes them hard to differentiate. For the instrument, saturation and bubbles class instances (see Fig. 10 a.), most methods obtained high performance compared to other classes (e.g., the top three teams obtained 0.853, 0.844, 0.848 for the instrument; 0.722, 0.758, 0.703 for the saturation; and 0.738, 0.693, 0.693 for the bubbles class instance, respectively), artefact (DSC < 0.520) was among the worst class for most teams and for the baseline methods. This is mostly due to the variable size of artefacts; and the bubbles class instance is predominantly confused with either artefact or the specularity class (see Fig. 10 b.). Additionally, due to small sized and sparsely scattered

specularity or bubble regions in some cases (for e.g., 4th image from left in Fig. 3(a)), the annotator variability for these samples can have affected method performances for these classes. While checking for such biases is beyond the conducted study, we refer to the work by Rolnick et al. (2017). The authors suggested that in general deep learning models are capable of generalizing from training data where the correct labels are outnumbered by the incorrect ones. However, the authors also acknowledged that a decrease in performance is inevitable and necessary steps such as using larger batch size and downscaling learning rate can help mitigate these issues.

Unlike the EAD2020, the EDD2020 segmentation task comprised of larger shaped regions and only a few classes confused (see 1 b.). Most methods scored comparably high DSC values with over 75% for most of the disease classes except for suspicious class by most of the team. However, Fig. 11(b) (top) shows that while majority of teams were able to segment NDBE class area, the teams either missed the HGD area or miss classified HGD as suspicious class instance. It is to be noted that there is a very subtle difference between the HGD and the suspicious region even for the expert endoscopists. Similar observation can be found for the segmentation of protruded structures (Fig. 11(b), bottom) where most methods confused the class with the polyp class and the top two teams (*adrian*, *sahadate*) classified it as cancer class. Looking up into our expert consensus notes we found that these samples had hard to reach agreement cases (i.e., suspicious and HGD classes; and cancer and polyp region).

Hypothesis V: Instead of hard scoring of predicted mask classes that penalizes the method performance heavily in presence of marginal visual difference between classes and variability due to existing expert consensus in the dataset, probability maps can be used to mitigate such problem. Additionally, teams should be encouraged to report results for different batch size and learning rates for obtaining better insight regarding performance especially when datasets are prone to have some incorrect labels.

7. Conclusion

We provided a comprehensive analysis of the deep learning methods built to tackle two distinct challenges in the gastrointestinal endoscopy: a) artefact detection and segmentation and b) disease detection and segmentation. It has been possible by the crowd-sourcing initiative of the EndoCV2020 challenges. We have laid out the summary of the methods developed by the top 17 participating teams and compared their methods with the state-of-the-art detection and segmentation methods. Additionally, we dissected-different paradigms used by the teams and present a detailed analysis and discussion of the outcomes. We also suggested pathways to improve the methods for building reliable and clinically transferable methods. In future, we aim towards more holistic comparison of the built methods for clinical deployability by testing for hardware and software reliability in clinical setting.

Author contributions

S. Ali conceptualized the work, led the challenge and workshop, prepared the dataset, software and performed all analyses. M. Dmitrieva, N. Ghatwary, and S. Bano served as organising committee and participated in annotations. A. Bailey, B. Braden, J.E. East, R. Cannizzaro, D. Lamarque, S. Realdon were involved in the validation and quality checks of the annotations used in this challenge. G. Plolat, A. Temizel, A. Krenzer, A. Hekalo, YB. Guo, B. Matuszewski, M. Gridach, V. Yoganand assisted in compiling the related work and method section of the manuscript. S. Ali wrote most of the manuscript with inputs from M. Dmitrieva, N. Ghat-

wary, S. Bano and all co-authors. All authors participated in the revision of this manuscript and provided substantial input.

Declaration of Competing Interest

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

Acknowledgments

The research was supported by the National Institute for Health Research (NIHR) Oxford Biomedical Research Centre (BRC). The views expressed are those of the authors and not necessarily those of the NHS, the NIHR or the Department of Health. S. Ali, B. Braden, A. Bailey and J. E. East is supported by NIHR BRC and J. Rittscher by Ludwig Institute for Cancer Research and EPSRC Seebibyte Programme Grant (EP/M013774/1). We want to acknowledge Karl Storz for co-sponsoring the challenge workshop. We would also like to acknowledge the annotators and our EndoCV2020 challenge workshop proceedings reviewers and IEEE International Symposium on Biomedical Imaging 2020 challenge committee Michal Kozubek and Hans Johnson.

References

- Ali, S., Bailey, A., East, J.E., Leedham, S.J., Haghighat, M., Investigators, T., Lu, X., Rittscher, J., Braden, B., 2020. Artificial intelligence-driven real-time 3D surface quantification of Barrett's oesophagus for risk stratification and therapeutic response monitoring. *medRxiv* doi:10.1101/2020.10.04.20206482.
- Ali, S., Ghatwary, N.M., Braden, B., Lamarque, D., Bailey, A., Realdon, S., Cannizzaro, R., Rittscher, J., Daul, C., East, J., 2020. Endoscopy disease detection challenge 2020. *CoRR abs/2003.03376*.
- Ali, S., Zhou, F., Bailey, A., Braden, B., East, J.E., Lu, X., Rittscher, J., 2021. A deep learning framework for quality assessment and restoration in video endoscopy. *Med. Image Anal.* 68, 101900. doi:10.1016/j.media.2020.101900.
- Ali, S., Zhou, F., Braden, B., Bailey, A., Yang, S., Cheng, G., Zhang, P., Li, X., Kayser, M., Soberanis-Mukul, R.D., et al., 2020. An objective comparison of detection and segmentation algorithms for artefacts in clinical endoscopy. *Sci. Rep.* 10 (1), 1–15.
- Ali, S., Zhou, F., Daul, C., Braden, B., Bailey, A., Realdon, S., East, J., Wagnières, G., Loschenov, V., Grisan, E., et al., 2019. Endoscopy artefact detection (EAD 2019) challenge dataset. *arXiv preprint arXiv:1905.03209*.
- Allan, M., Kondo, S., Bodenstedt, S., Leger, S., Kadkhodamohammadi, R., Luengo, I., Fuentes, F., Flouty, E., Mohammed, A., Pedersen, M., Kori, A., Alex, V., Krishnamurthi, G., Rauber, D., Mendel, B., Palm, C., Bano, S., Saibro, G., Shih, C.-S., Chiang, H.-A., Zhuang, J., Yang, J., Iglovikov, V., Dobrenkii, A., Reddiboina, M., Reddy, A., Liu, X., Gao, C., Unberath, M., Kim, M., Kim, C., Kim, C., Kim, H., Lee, G., Ullah, I., Luna, M., Park, S. H., Azizian, M., Stoyanov, D., Maier-Hein, L., Speidel, S., 2020. 2018 robotic scene segmentation challenge. 2001.11190.
- Badrinarayanan, V., Kendall, A., Cipolla, R., 2017. Segnet: a deep convolutional encoder-decoder architecture for image segmentation. *IEEE Trans Pattern Anal. Mach. Intell.* 39 (12), 2481–2495.
- Balasubramanian, V., Kumar, R., Kamireddi, S.J., Sathish, R., Sheet, D., 2020. Semantic segmentation, detection AND localisation of mucosal lesions from gastrointestinal endoscopic images using SUMNET. In: *Proceedings of the 2nd International Workshop and Challenge on Computer Vision in Endoscopy, EndoCV@ISBI 2020*, Iowa City, Iowa, USA, 3rd April 2020. In: *CEUR Workshop Proceedings*, 2595, pp. 82–83.
- Bano, S., Vasconcelos, F., Shepherd, L.M., Poorten, E.V., Vercouteren, T., Ourselin, S., David, A.L., Deprest, J., Stoyanov, D., 2020. Deep placental vessel segmentation for fetoscopic mosaicking. *arXiv preprint arXiv:2007.04349*.
- Bernal, J., Sánchez, F.J., Fernández-Esparrach, G., Gil, D., Rodríguez, C., Vilariño, F., 2015. Wm-dova maps for accurate polyp highlighting in colonoscopy: validation vs. saliency maps from physicians. *Computer. Med. Imag. and Graph.* 43, 99–111.
- Bernal, J., Sánchez, J., Vilarino, F., 2012. Towards automatic polyp detection with a polyp appearance model. *Patt. Recognit.* 45 (9), 3166–3182.
- Bernal, J., et al., 2017. Comparative validation of polyp detection methods in video colonoscopy: results from the miccai 2015 endoscopic vision challenge. *IEEE Trans. Med. Imag.* 36 (6), 1231–1249.
- Bernal, J., et al., 2018. Polyp detection benchmark in colonoscopy videos using gtcreator: A novel fully configurable tool for easy and fast annotation of image databases. In: *Proc. Comput. Assist. Radiol. Surg. (CARS)*.
- Bodla, N., Singh, B., Chellappa, R., Davis, L.S., 2017. Soft-nms—improving object detection with one line of code. In: *Proceedings of the IEEE international conference on computer vision*, pp. 5561–5569.
- Boland, C., Luciani, M., Gasche, C., A., G., 2005. Infection, inflammation, and gastrointestinal cancer. *Gut* 54 (9), 1321–1331. doi:10.1136/gut.2004.060079.
- Bolya, D., Zhou, C., Xiao, F., Lee, Y.J., 2019. YOLACT: Real-time instance segmentation. *IEEE international conference on computer vision*.
- Brandao, P., Mazomenos, E., Ciuti, G., Cali, R., Bianchi, F., Mencias, A., Dario, P., Koulaouzidis, A., Arezzo, A., Stoyanov, D., 2017. Fully convolutional neural networks for polyp segmentation in colonoscopy. In: *Medical Imaging 2017: Computer-Aided Diagnosis*. International Society for Optics and Photonics. SPIE, pp. 101–107. doi:10.1117/12.2254361.
- Brandao, P., Zisimopoulos, O., Mazomenos, E., Ciuti, G., Bernal, J., Visentini-Scarzanella, M., Mencias, A., Dario, P., Koulaouzidis, A., Arezzo, A., Hawkes, D.J., Stoyanov, D., 2018. Towards a computed-aided diagnosis system in colonoscopy: automatic polyp segmentation using convolution neural networks. *Journal of Medical Robotics Research* 03 (02), 1840002. doi:10.1142/S2424905X18400020.
- Cai, Z., Vasconcelos, N., 2018. Cascade r-cnn: Delving into high quality object detection. In: *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pp. 6154–6162.
- Chen, H., Lian, C., Wang, L., 2020. Endoscopy artefact detection and segmentation using deep convolutional neural network. In: *Proceedings of the 2nd International Workshop and Challenge on Computer Vision in Endoscopy, EndoCV@ISBI 2020*, Iowa City, Iowa, USA, 3rd April 2020. *CEUR-WS.org*, pp. 37–41.
- Chen, L.-C., Zhu, Y., Papandreou, G., Schroff, F., Adam, H., 2018. Encoder-decoder with atrous separable convolution for semantic image segmentation. In: *Proceedings of the European conference on computer vision (ECCV)*, pp. 801–818.
- Chlebus, G., Meine, H., Moltz, J.H., Schenk, A., 2017. Neural network-based automatic liver tumor segmentation with random forest-based candidate filtering. *arXiv preprint arXiv:1706.00842*.
- Choi, Y.H., Lee, Y.C., Hong, S., Kim, J., Won, H., Kim, T., 2020. Centernet-based detection model and u-net-based multi-class segmentation model for gastrointestinal diseases. In: *Proceedings of the 2nd International Workshop and Challenge on Computer Vision in Endoscopy, EndoCV@ISBI 2020*, Iowa City, Iowa, USA, 3rd April 2020. *CEUR-WS.org*, pp. 73–75.
- Cubuk, E.D., Zoph, B., Shlens, J., Le, Q.V., 2019. RandAugment: practical data augmentation with no separate search. *arXiv preprint arXiv:1909.13719* 2 (4), 7.
- Duan, K., Bai, S., Xie, L., Qi, H., Huang, Q., Tian, Q., 2019. Centernet: Keypoint triplets for object detection. In: *Proceedings of the IEEE International Conference on Computer Vision*, pp. 6569–6578.
- Eluri, S., Shaheen, N., 2017. Barrett's esophagus: diagnosis and management. *Gastrointest. Endosc.* 85 (4), 889–903. doi:10.1016/j.gie.2017.01.007.
- Everingham, M., Van Gool, L., Williams, C.K.I., Winn, J., Zisserman, A., 2012. The PASCAL Visual Object Classes Challenge 2012 (VOC2012) Results. <http://www.pascal-network.org/challenges/VOC/voc2012/workshop/index.html>.
- Falk, T., Mai, D., Bensch, R., Çiçek, Ö., Abdulkadir, A., Marrakchi, Y., Böhm, A., Deubner, J., Jäckel, D., Seiwald, K., Dovzhenko, A., Tietz, O., Dal Bosco, C., Walsh, S., Saltukoglu, D., Tay, T.L., Prinz, M., Palme, K., Simons, M., Diester, I., Brox, T., Ronneberger, O., 2019. U-Net: Deep learning for cell counting, detection, and morphometry. *Nat. Methods* 16 (1), 67–70. doi:10.1038/s41592-018-0261-2.
- Formosa, G.A., Micah Prendergast, J., Sean Humbert, J., Rentschler, M.E., 2020. Non-linear dynamic modeling of a robotic endoscopy platform on synthetic tissue substrates. *J. Dyn. Syst. Meas. Control* 143 (1). doi:10.1115/1.4048190.
- Gao, J., Vázquez, D., Bernal, J., Sánchez, F.J., Fernández-Esparrach, G., López, A.M., Romero, A., Drozdal, M., Courville, A., 2017. A benchmark for endoluminal scene segmentation of colonoscopy images. *J. Healthc. Eng.* 4037190.
- Gao, X.W., Braden, B., 2020. Artefact detection and segmentation based on a deep learning system. In: *Proceedings of the 2nd International Workshop and Challenge on Computer Vision in Endoscopy, EndoCV@ISBI 2020*, Iowa City, Iowa, USA, 3rd April 2020. In: *CEUR Workshop Proceedings*, 2595, pp. 80–81.
- Grishick, R., Donahue, J., Darrell, T., Malik, J., 2014. Rich feature hierarchies for accurate object detection and semantic segmentation. In: *Proceedings of the IEEE conference on computer vision and pattern recognition*, pp. 580–587.
- Gridach, M., Voiculescu, I., 2020. OXENDONET: A dilated convolutional neural network for endoscopic artefact segmentation. In: *Proceedings of the 2nd International Workshop and Challenge on Computer Vision in Endoscopy, EndoCV@ISBI 2020*, Iowa City, Iowa, USA, 3rd April 2020. *CEUR-WS.org*, pp. 26–29.
- Guo, X., Zhang, N., Guo, J., Zhang, H., Hao, Y., Hang, J., 2019. Automated polyp segmentation for colonoscopy images: a method based on convolutional neural networks and ensemble learning. *Med. Phys.* 46 (12), 5666–5676. doi:10.1002/mp.13865.
- Guo, Y., Bernal, J., J. Matuszewski, B., 2020. Polyp segmentation with fully convolutional deep neural networks—extended evaluation study. *Journal of Imaging* 6 (7), 69.
- Guo, Y.B., Zheng, Q., Matuszewski, B.J., 2020. Deep encoder-decoder networks for artefacts segmentation in endoscopy images. In: *Proceedings of the 2nd International Workshop and Challenge on Computer Vision in Endoscopy, EndoCV@ISBI 2020*, Iowa City, Iowa, USA, 3rd April 2020. *CEUR-WS.org*, pp. 18–21.
- Gupta, S., Ali, S., Goldsmith, L., Turney, B., Rittscher, J., 2020. Mi-unet: Improved segmentation in ureteroscopy. In: *2020 IEEE 17th International Symposium on Biomedical Imaging (ISBI)*, pp. 212–216.
- He, K., Gkioxari, G., Dollár, P., Girshick, R., 2017. Mask r-cnn. In: *Proceedings of the IEEE international conference on computer vision*, pp. 2961–2969.
- Horie, Y., Yoshio, T., Aoyama, K., Yoshimizu, S., Horiuchi, Y., Ishiyama, A., Hirasawa, T., Tsuchida, T., Ozawa, T., Ishihara, S., et al., 2019. Diagnostic outcomes of esophageal cancer by artificial intelligence using convolutional neural networks. *Gastrointest. Endosc.* 89 (1), 25–32.
- Hu, H., Guo, Y., 2020. Endoscopic artefact detection in mmdetection. In: *Proceedings of the 2nd International Workshop and Challenge on Computer Vision in Endoscopy, EndoCV@ISBI 2020*, Iowa City, Iowa, USA, 3rd April, pp. 78–79. Vol. 2595, *CEUR Workshop Proceedings*.

- Huynh, L.D., Boutry, N., 2020. A u-net++ with pre-trained efficientnet backbone for segmentation of diseases and artifacts in endoscopy images and videos. In: Proceedings of the 2nd International Workshop and Challenge on Computer Vision in Endoscopy, EndoCV@ISBI 2020, Iowa City, Iowa, USA, 3rd April 2020. CEUR-WS.org, pp. 13–17.
- Incetan, K., Celik, I.O., Obeid, A., Gokceler, G.I., Ozyoruk, K.B., Almalioğlu, Y., Chen, R.J., Mahmood, F., Gilbert, H., Durr, N.J., Turan, M., 2020. Vr-caps: A virtual environment for capsule endoscopy. 2008.12949.
- Jadhav, S., Bamba, U., Chavan, A., Tiwari, R., Raj, A., 2020. Multi-plateau ensemble for endoscopic artefact segmentation and detection. In: Proceedings of the 2nd International Workshop and Challenge on Computer Vision in Endoscopy, EndoCV@ISBI 2020, Iowa City, Iowa, USA, 3rd April 2020. CEUR-WS.org, pp. 22–25.
- Jha, D., Ali, S., Emanuelsen, K., Hicks, S., Thambawita, V., Garcia-Ceja, E., Riegler, M., de Lange, T., Schmidt, P.T., Johansen, H., Johansen, D., Halvorsen, P., 2020. Kvasir-instrument: Diagnostic and therapeutic tool segmentation dataset in gastrointestinal endoscopy. OSF Preprints.
- Jia, X., Mai, X., Cui, Y., Yuan, Y., Xing, X., Seo, H., Xing, L., Meng, M.Q., 2020. Automatic polyp recognition in colonoscopy images using deep learning and two-stage pyramidal feature prediction. IEEE Trans. Autom. Sci. Eng. 17 (3), 1570–1584.
- Jorge, B., Aymeric, H., 2017. GastrointestinalG imagel analysisANA challenge. <https://endovissub2017-giana.grand-challenge.org/>.
- Kaul, C., Manandhar, S., Pears, N., 2019. Focusnet: an attention-based fully convolutional network for medical image segmentation. In: 2019 IEEE 16th International Symposium on Biomedical Imaging (ISBI 2019). IEEE, pp. 455–458.
- Kayser, M., Soberanis-Mukul, R.D., Albarqouni, S., Navab, N., 2019. Focal loss for artefact detection in medical endoscopy. In: Proceedings of the 1st International Workshop and Challenge on Computer Vision in Endoscopy.
- Khan, M.A., Choo, J., 2019. Multi-class artefact detection in video endoscopy via convolution neural networks. In: Proceedings of the 1st International Workshop and Challenge on Computer Vision in Endoscopy.
- Krenzer, A., Hekalo, A., Puppe, F., 2020. Endoscopic detection and segmentation of gastroenterological diseases with deep convolutional neural networks. In: Proceedings of the 2nd International Workshop and Challenge on Computer Vision in Endoscopy, EndoCV@ISBI 2020, Iowa City, Iowa, USA, 3rd April 2020. CEUR-WS.org, pp. 58–63.
- Law, H., Deng, J., 2018. Cornernet: Detecting objects as paired keypoints. In: Proceedings of the European Conference on Computer Vision (ECCV), pp. 734–750.
- Lee, J., Park, S.W., Kim, Y.S., Lee, K.J., Sung, H., Song, P.H., Yoon, W.J., Moon, J.S., 2017. Risk factors of missed colorectal lesions after colonoscopy. Medicine 96 (27), e7468–e7468.
- Lin, T.-Y., Dollár, P., Girshick, R., He, K., Hariharan, B., Belongie, S., 2017. Feature pyramid networks for object detection. In: Proceedings of the IEEE conference on computer vision and pattern recognition, pp. 2117–2125.
- Lin, T.-Y., Goyal, P., Girshick, R., He, K., Dollár, P., 2017. Focal loss for dense object detection. In: Proceedings of the IEEE international conference on computer vision, pp. 2980–2988.
- Lin, T.-Y., Maire, M., Belongie, S., Hays, J., Perona, P., Ramanan, D., Dollár, P., Zitnick, C.L., 2014. Microsoft coco: Common objects in context. In: European conference on computer vision. Springer, pp. 740–755.
- Liu, W., Anguelov, D., Erhan, D., Szegedy, C., Reed, S., Fu, C.-Y., Berg, A.C., 2016. Ssd: Single shot multibox detector. In: European conference on computer vision. Springer, pp. 21–37.
- Liu, Y., Zhao, Z., Chang, F., Hu, S., 2020. An anchor-free convolutional neural network for real-time surgical tool detection in robot-assisted surgery. IEEE Access 8, 78193–78201.
- Long, J., Shelhamer, E., Darrell, T., 2015. Fully convolutional networks for semantic segmentation. In: Proceedings of the IEEE conference on computer vision and pattern recognition, pp. 3431–3440.
- Mahmood, F., Chen, R., Durr, N.J., 2018. Unsupervised reverse domain adaptation for synthetic medical images via adversarial training. IEEE Trans. Med. Imaging 37 (12), 2572–2581. doi:10.1109/TMI.2018.2842767.
- Nguyen, N.-Q., Vo, D.M., Lee, S.-W., 2020. Contour-aware polyp segmentation in colonoscopy images using detailed upsampling encoder-decoder networks. IEEE Access 8, 99495–99508.
- Nguyen, N.T., Tran, D.Q., Nguyen, D.B., 2020. Detection and segmentation of endoscopic artefacts and diseases using deep architectures. In: Proceedings of the 2nd International Workshop and Challenge on Computer Vision in Endoscopy, EndoCV@ISBI 2020, Iowa City, Iowa, USA, 3rd April 2020. CEUR-WS.org, pp. 64–67.
- Norman, B., Pedoia, V., Majumdar, S., 2018. Use of 2d u-net convolutional neural networks for automated cartilage and meniscus segmentation of knee mr imaging data to determine relaxometry and morphometry. Radiology 288 (1), 177–185.
- Oksuz, I., Clough, J.R., King, A.P., Schnabel, J.A., 2019. Artefact detection in video endoscopy using retinanet and false loss function. In: Proceedings of the 1st International Workshop and Challenge on Computer Vision in Endoscopy, EndoCV@ISBI 2019.
- Pengyi, Z., Xiaoqiong, L.Y.Z., 2019. Ensemble mask-aided r-cnn. In: Ali, S., Zhou, F. (Eds.), Proceedings of the 1st International Workshop and Challenge on Computer Vision in Endoscopy, EndoCV@ISBI 2019.
- Polat, G., Sen, D., Inci, A., Temizel, A., 2020. Endoscopic artefact detection with ensemble of deep neural networks and false positive elimination. In: Proceedings of the 2nd International Workshop and Challenge on Computer Vision in Endoscopy, EndoCV@ISBI 2020, Iowa City, Iowa, USA, 3rd April 2020. CEUR-WS.org, pp. 8–12.
- Rawla, P., Sunkara, T., Barsouk, A., 2019. Epidemiology of colorectal cancer: incidence, mortality, survival, and risk factors. Prz. Gastroenterol. 14 (2), 89–103.
- Redmon, J., Divvala, S., Girshick, R., Farhadi, A., 2016. You only look once: Unified, real-time object detection. In: Proceedings of the IEEE conference on computer vision and pattern recognition, pp. 779–788.
- Ren, S., He, K., Girshick, R., Sun, J., 2015. Faster r-cnn: Towards real-time object detection with region proposal networks. In: Advances in neural information processing systems, pp. 91–99.
- Rezvy, S., Zebin, T., Braden, B., Pang, W., Taylor, S., Gao, X.W., 2020. Transfer learning for endoscopy disease detection & segmentation with mask-rcnn benchmark architecture. In: Proceedings of the 2nd International Workshop and Challenge on Computer Vision in Endoscopy, EndoCV@ISBI 2020, Iowa City, Iowa, USA, 3rd April 2020. CEUR-WS.org, pp. 68–72.
- Rolnick, D., Veit, A., Belongie, S.J., Shavit, N., 2017. Deep learning is robust to massive label noise. CoRR abs/1705.10694.
- Ronneberger, O., Fischer, P., Brox, T., 2015. U-net: Convolutional networks for biomedical image segmentation. In: International Conference on Medical image computing and computer-assisted intervention. Springer, pp. 234–241.
- Ross, T., Reinke, A., Full, P.M., Wagner, M., Kenngott, H., Apitz, M., Hempe, H., Mindroc Filimon, D., Scholz, P., Tran, T.N., Bruno, P., Arbelez, P., Bian, G.-B., Bodenstedt, S., Bolmgren, J.L., Bravo-Sánchez, L., Chen, H.-B., González, C., Guo, D., Halvorsen, P., Heng, P.-A., Hosgor, E., Hou, Z.-G., Isensee, F., Jha, D., Jiang, T., Jin, Y., Kirtac, K., Kletz, S., Leger, S., Li, Z., Maier-Hein, K.H., Ni, Z.-L., Riegler, M.A., Schoeffmann, K., Shi, R., Speidel, S., Stenzel, M., Twick, I., Wang, G., Wang, J., Wang, L., Zhang, Y., Zhou, Y.-J., Zhu, L., Wiesenfarth, M., Kopp-Schneider, A., Müller-Stich, B.P., Maier-Hein, L., 2020. Comparative validation of multi-instance instrument segmentation in endoscopy: results of the robust-mis 2019 challenge. Med. Image Anal. 101920. doi:10.1016/j.media.2020.101920.
- Seferbekov, S.S., Iglovikov, V.I., Buslaev, A.V., Shvets, A.A., 2018. Feature pyramid network for multi-class land segmentation. CoRR abs/1806.03510.
- Sevastopolsky, A., 2017. Optic disc and cup segmentation methods for glaucoma detection with modification of u-net convolutional neural network. Pattern Recognit. Image Anal. 27 (3), 618–624.
- Shin, Y., Qadir, H.A., Aabakken, L., Bergsland, J., Balasingham, I., 2018. Automatic colon polyp detection using region based deep cnn and post learning approaches. IEEE Access 6, 40950–40962.
- Silva, J., Histic, A., Romain, O., Dray, X., Granado, B., 2014. Toward embedded detection of polyps in wce images for early diagnosis of colorectal cancer. Int. J. Comput. Assist. Radiol. and Surg. 9 (2), 283–293.
- Soberanis-Mukul, R.D., Kayser, M., Zvereva, A., Klare, P., Navab, N., Albarqouni, S., 2020. A learning without forgetting approach to incorporate artifact knowledge in polyp localization tasks. CoRR abs/2002.02883.
- Subramanian, A., Srivatsan, K., 2020. Exploring deep learning based approaches for endoscopic artefact detection and segmentation. In: Proceedings of the 2nd International Workshop and Challenge on Computer Vision in Endoscopy, EndoCV@ISBI 2020, Iowa City, Iowa, USA, 3rd April 2020. CEUR-WS.org, pp. 51–56.
- Suhui Yang, G.C., 2019. Endoscopic artefact detection and segmentation with deep convolutional neural network. In: Ali, S., Zhou, F. (Eds.), Proceedings of the 1st International Workshop and Challenge on Computer Vision in Endoscopy, EndoCV@ISBI 2019.
- Sun, C., Guo, S., Zhang, H., Li, J., Chen, M., Ma, S., Jin, L., Liu, X., Li, X., Qian, X., 2017. Automatic segmentation of liver tumors from multiphase contrast-enhanced ct images based on fcns. Artif. Intell. Med. 83, 58–66.
- Urban, G., Tripathi, P., Alkayali, T., Mittal, M., Jalali, F., Karnes, W., Baldi, P., 2018. Deep learning localizes and identifies polyps in real time with 96% accuracy in screening colonoscopy. Gastroenterology 155 (4), 1069–1078.
- Vishnusi, V., Prakash, P., Shivashankar, N., 2020. A submission note on EAD 2020: Deep learning based approach for detecting artefacts in endoscopy. In: Proceedings of the 2nd International Workshop and Challenge on Computer Vision in Endoscopy, EndoCV@ISBI 2020, Iowa City, Iowa, USA, 3rd April 2020. CEUR-WS.org, pp. 30–36.
- Wang, D., Zhang, N., Sun, X., Zhang, P., Zhang, C., Cao, Y., Liu, B., 2019. Afp-net: Real-time anchor-free polyp detection in colonoscopy. In: 2019 IEEE 31st International Conference on Tools with Artificial Intelligence (ICTAI). IEEE, pp. 636–643.
- Wang, K.K., Tian, J.M., Gorospe, E., Penfield, J., Prasad, G., Goddard, T., Wong-keesong, M., Buttar, N., Lutzke, L., Krishnadath, S., 2012. Diseases of the esophagus : official journal of the international society for diseases of the esophagus. Dis Esophagus 25 (4), 349–355. doi:10.1111/j.1442-2050.2012.01342.x.
- Wang, P., Xiao, X., Brown, J.R.G., Berzin, T.M., Tu, M., Xiong, F., Hu, X., Liu, P., Song, Y., Zhang, D., et al., 2018. Development and validation of a deep-learning algorithm for the detection of polyps during colonoscopy. Nat. Biomed. Eng. 2 (10), 741–748.
- Williams, J.G., Pullan, R.D., Hill, J., Horgan, P.G., Salmo, E., Buchanan, G.N., Rasheed, S., McGee, S.G., Haboubi, N., 2013. Management of the malignant colorectal polyp: acpgbi position statement. Colorectal Disease 15 (s2), 1–38. doi:10.1111/codi.12262.
- Yamada, M., Saito, Y., Imaoka, H., Saiko, M., Yamada, S., Kondo, H., Takamaru, H., Sakamoto, T., Sese, J., Kuchiba, A., et al., 2019. Development of a real-time endoscopic image diagnosis support system using deep learning technology in colonoscopy. Sci. Rep. 9 (1), 1–9.
- Yang, S., Cheng, G., 2019. Endoscopic artefact detection and segmentation with deep convolutional neural network. In: Proceedings of the 1st International Workshop and Challenge on Computer Vision in Endoscopy.
- Yu, F., Koltun, V., 2015. Multi-scale context aggregation by dilated convolutions. arXiv preprint arXiv:1511.07122.

- Yu, Z., Guo, Y., 2020. Endoscopic artefact detection using cascade R-CNN based model. In: Proceedings of the 2nd International Workshop and Challenge on Computer Vision in Endoscopy, EndoCV@ISBI 2020, Iowa City, Iowa, USA, 3rd April 2020. CEUR-WS.org, pp. 42–46.
- Yun, S., Han, D., Oh, S.J., Chun, S., Choe, J., Yoo, Y., 2019. Cutmix: Regularization strategy to train strong classifiers with localizable features. In: Proceedings of the IEEE International Conference on Computer Vision, pp. 6023–6032.
- Zhang, X., Chen, F., Yu, T., An, J., Huang, Z., Liu, J., Hu, W., Wang, L., Duan, H., Si, J., 2019. Real-time gastric polyp detection using convolutional neural networks. PLoS ONE 14 (3), 1–16. doi:[10.1371/journal.pone.0214133](https://doi.org/10.1371/journal.pone.0214133).
- Zhang, Y.-y., Xie, D., 2019. Detection and segmentation of multi-class artifacts in endoscopy. Journal of Zhejiang University-SCIENCE B 20 (12), 1014–1020.
- Zhao, H., Shi, J., Qi, X., Wang, X., Jia, J., 2017. Pyramid scene parsing network. In: Proceedings of the IEEE conference on computer vision and pattern recognition, pp. 2881–2890.
- Zhou, X., Zhuo, J., Krahenbuhl, P., 2019. Bottom-up object detection by grouping extreme and center points. In: Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, pp. 850–859.