

AxIoU: An Axiomatically Justified Measure for Video Moment Retrieval

Riku Togashi
Cyberagent, Inc., Waseda University

Mayu Otani
Cyberagent, Inc.

Yuta Nakashima
Osaka University

Esa Rahtu
Tampere University

Janne Heikkilä
University of Oulu

Tetsuya Sakai
Waseda University

Abstract

Evaluation measures have a crucial impact on the direction of research. Therefore, it is of utmost importance to develop appropriate and reliable evaluation measures for new applications where conventional measures are not well suited. Video Moment Retrieval (VMR) is one such application, and the current practice is to use $R@K, \theta$ for evaluating VMR systems. However, this measure has two disadvantages. First, it is rank-insensitive: It ignores the rank positions of successfully localised moments in the top- K ranked list by treating the list as a set. Second, it binarizes the Intersection over Union (IoU) of each retrieved video moment using the threshold θ and thereby ignoring fine-grained localisation quality of ranked moments.

We propose an alternative measure for evaluating VMR, called Average Max IoU (AxIoU), which is free from the above two problems. We show that AxIoU satisfies two important axioms for VMR evaluation, namely, **Invariance against Redundant Moments** and **Monotonicity with respect to the Best Moment**, and also that $R@K, \theta$ satisfies the first axiom only. We also empirically examine how AxIoU agrees with $R@K, \theta$, as well as its stability with respect to change in the test data and human-annotated temporal boundaries.

1. Introduction

Video Moment Retrieval (VMR) has been explored to find relevant fragments of videos (*i.e.* video moments) based on a user’s textual query [10, 15]. Most existing VMR systems [10, 22, 39, 40, 43] cast the problem of finding video moments into a ranking problem. For evaluating ranked lists of video moments, $R@K, \theta$ is widely adopted in the literature [10]. $R@K, \theta$ for a query q is defined as 1 if at least one relevant video moment in the top K of the ranked list has an Intersection over Union (IoU) larger than θ with the ground truth for q .

$R@K, \theta$ has two disadvantages as illustrated in Figure 1.

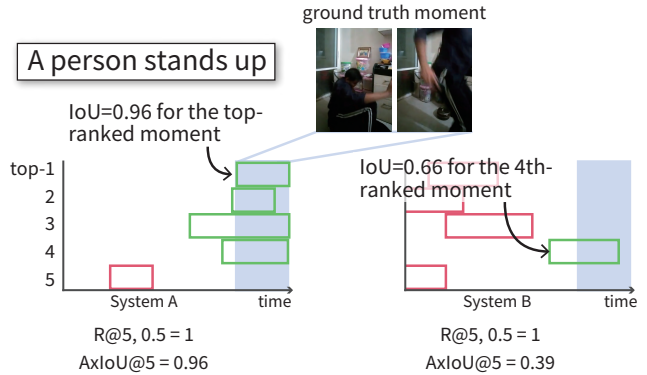


Figure 1. The system on the left shows a moment with a large overlap with the ground truth (blue band) in the top of the ranked list, and the system on the right illustrates a moment with a much smaller overlap with the ground truth at rank 4. According to $R@5, 0.5$, the two systems are equally effective. We propose AxIoU whose measurements reflect the localisation quality (*i.e.* IoU) and the rank of successfully retrieved video moments. The photo in the figure is taken from Charades-STA [10].

First, it is *rank-insensitive*, as the video moments in the top- K ranked list are treated as a set, and their ranks are not considered. Second, it is *localisation-insensitive*, *i.e.*, the exact position (start and end points) of the video moment does not affect the measurement as it binarizes the IoU of each video moment using threshold θ . Thus, $R@K, \theta$ only provides a binary measurement for a ranked list in an all-or-nothing manner, ignoring the ranking and localisation quality of top- K predicted video moments. As we shall demonstrate in this paper, these properties of $R@K, \theta$ are problematic for reliable evaluation. $R@K, \theta$ cannot distinguish ranked lists with different quality due to the binary property, while leading to instability under a small number of evaluation samples and label ambiguity [2, 15, 27, 36]. Moreover, $R@K, \theta$ evaluates rather different aspects of system quality depending on a parameter setting and can conflict with each other; for example, a ranked list whose second moment achieves $\text{IoU} = 0.71$ is measured to be 1.0 in terms of

$R@2, 0.7$ but to be 0.0 in terms of $R@1, 0.7$ when the first moment has $IoU = 0.69$. These undesirable properties of $R@K, \theta$ should be carefully considered for future studies because conclusions drawn from potentially unstable measures may not generalise well. In practice, the instability of $R@K, \theta$ implies that we may underestimate a VMR method by adopting a non-best model based on validation $R@K, \theta$.

In this paper, we propose an alternative measure for evaluating VMR systems, called *Average Max IoU* (AxIoU), which does not suffer from the problems above with $R@K, \theta$. To evaluate evaluation measures, we take an axiomatic approach [3, 9, 34, 35] and introduce two important axioms that an effectiveness measure for VMR must satisfy, namely, *invariance against redundant moments* and *monotonicity with respect to the best moment*. We show that $R@K, \theta$ only satisfies the first axiom. We also empirically investigate the properties of AxIoU in practical terms, namely, agreement with conventional $R@K, \theta$ and the stability to the size of a dataset and label ambiguity.

2. Related Work

Most of the prior studies of VMR adopted the $R@K, \theta$ measure [7, 10, 11, 19, 22, 38, 39, 40, 42, 43]. Gao *et al.* [10] proposed the use of this measure for VMR by referring to the work of Hu *et al.* [16], which is an early study on an object retrieval task with textual queries. The values of K and θ are chosen for each dataset. For example, the combinations of $K = 1, 5, 10$ and $\theta = 0.3, 0.5, 0.7$ are widely adopted in Charades-STA [10], ActivityNet [5, 19], and DiDeMo [15]. In the TACoS dataset [29], relatively relaxed values of θ (*i.e.* $\theta = 0.1, 0.3, 0.5$) are used. For evaluating methods that output only one moment per query [12, 14, 41], $R@K, \theta$ with $K = 1$ is often adopted. Lei *et al.* recently proposed a new retrieval task called video corpus moment retrieval [20], in which a system requires to retrieve relevant moments from multiple videos. Owing to a large number of candidate moments, they utilise large values of K such as $K = 100$. However, the common practice of reporting multiple settings of $R@K, \theta$ is controversial. As we shall demonstrate in this paper, different parameter settings often lead to different system rankings, from which it may be difficult to draw useful conclusions from the evaluation.

Prior studies suggested that the inter-rater agreement of human-annotated temporal boundaries is often not strong [2, 15, 27, 36]. Hendricks *et al.* found that there are multiple video moments, which can be described by a textual query [15]; to alleviate this label ambiguity, they developed a user interface. Sigurdsson *et al.* and Alwassel *et al.* also reported that human-annotated temporal regions do not agree well with each other [2, 36]. Otani *et al.* observed high label ambiguity in Charades-STA and ActivityNet [27]. Nevertheless, the binarization of IoU values

in $R@K, \theta$ introduces potential instability to the change of labels. In particular, a large value of θ requires exactly located temporal regions and thereby being noisy, inheriting label ambiguity.

In the context of object detection, in which evaluation measures often rely on a threshold parameter for spatial IoU, prior studies have discussed the disadvantages of a fixed threshold [13, 26, 28]. On the MSCOCO [21] dataset, an average of measures over IoU threshold values is adopted for evaluating fine-grained localisation quality; the measure is called COCO mean average precision (mAP). Oksuz *et al.* proposed an object detection measure to directly quantify the bounding box tightness by introducing IoU values without thresholding in their measure [26]. Hall *et al.* have explored a way to improve the spatial quality evaluation of detected regions beyond the conventional box-based IoU while reducing the parameters in evaluation measures [13]. In temporal localisation tasks for videos (*e.g.* action detection), Alwassel *et al.* used a COCO mAP-like measure for evaluation [2].

In contrast to rank-insensitive set retrieval measures (*e.g.* precision and recall), ranked retrieval measures have been explored for evaluating the quality of a list of ranked items, such as normalised discounted cumulative gain (nDCG) [17]. Such measures often have weights for the rank positions in a retrieval result; for example, the discount function in nDCG can be regarded as the importance of each position. Based on the interpretation of the position weights from the viewpoint of *user models*, prior studies have developed various evaluation measures [6, 24, 30, 32].

The evaluation of evaluation measures is often challenging as it requires the true evaluation results a priori. One approach to verify the experiments based on an evaluation measure is to collect human manual assessments for *search engine result pages* (SERPs) [33]. For new applications such as VMR, it is often costly to establish a reliable environment to collect the gold data that aligns well with the “true” quality; we may need to study such as human effects on the reliability of the gold data [18]. An axiomatic approach is another direction for the verification of evaluation measures [3, 9, 34, 35]. By formally defining requirements that a measure should satisfy, we can analytically confirm the validity of measures. Such requirements inevitably depend on a number of assumptions. However, this is also true for the assessment-based approach because guidelines for assessors implicitly involve assumptions on users’ behaviours [33].

In this paper, we propose an alternative VMR measure, AxIoU, which is an instantiation of normalised cumulative utility (NCU) [31, 32], which is a wide class of information retrieval measures including AP. Our proposed measure considers the rank positions and IoU values of video moments. To confirm the properties of measures, we take

an axiomatic approach. The derivation of AxIoU is related to COCO mAP, whereas AxIoU analytically reduces the binarization process for IoU values. Through empirical experiments, we confirm the numerical properties of AxIoU while showing the undesirable behaviours of $R@K, \theta$.

3. Preliminaries

3.1. Notations

Our goal is to develop a measure $\mu(q, \sigma)$ that estimates the retrieval effectiveness of a system σ based on a test query q . We also denote by $\mu(\mathcal{Q}, \sigma) = (1/|\mathcal{Q}|) \sum_{q \in \mathcal{Q}} \mu(q, \sigma)$ the mean of the measurements based on a test query set $\mathcal{Q}$. For a query $q \in \mathcal{Q}$, the system σ sorts the set $\mathcal{M}_q$ of candidate moments and creates the ranked list σ_q . We also denote by $\sigma_q(k) \in \mathcal{M}_q$ the moment ranked at position k in σ_q . Let $r_q(m) \in [0, 1]$ be the relevance score of a moment $m \in \mathcal{M}_q$, computed as the temporal IoU (Intersection over Union) between m and the ground truth region for q . Where there is no ambiguity, we will also denote it by $r(m)$.

3.2. $R@K, \theta$

First, we formally define the conventional measure, $R@K, \theta$ [10], and clarify what it quantifies as well as its limitations. Here, we denote by $\mathbb{1}: \mathbb{B} \mapsto \{0, 1\}$ the indicator function for the boolean variable X that takes 1 if X is true and 0 if X is false. We express $R@K, \theta$ and Mean $R@K, \theta$ as follows.

$$\begin{aligned} R@K, \theta(q, \sigma) &:= \mathbb{1} \left\{ \sum_{k=1}^K \mathbb{1} \{r(\sigma_q(k)) > \theta\} > 0 \right\} \\ &= \mathbb{1} \left\{ \max_{1 \leq k \leq K} r(\sigma_q(k)) > \theta \right\}. \end{aligned} \quad (1)$$

The value of $R@K, \theta$ depends entirely on whether the most relevant moment in the top- K retrieved results exceeds the θ threshold. It is clear from this that $R@K, \theta$ does not reward redundancy: the retrieved moments other than the most relevant one in the SERP do not count, even if they also exceed θ . We shall refer to such relevant moments as *redundant* moments.

The above property of $R@K, \theta$ is a desirable feature, since real VMR system users probably do not care about redundant moments in their SERPs. However, it is clear from Eq. 1 that $R@K, \theta$ has two potential shortcomings. First, $R@K, \theta$ is unchanged by the rank positions of the relevant moments: it is a set retrieval measure rather than a ranked retrieval measure. For $K > 1$, it cannot distinguish between a system that retrieves a perfectly relevant moment at rank 1, and a system that retrieves the same moment at rank K . Second, it binarizes the IoU of each moment using the θ threshold, and thereby ignores the degree of relevance

of each retrieved moment. Choosing an appropriate value of θ is practically problematic, especially given that K also needs to be chosen at the same time.

4. Proposed Measure

4.1. Average Max IoU Measure

To design a measure for VMR, we adopt the framework of a wide class of retrieval effectiveness measures, *normalised cumulative utility (NCU)* [31, 32]. NCU assumes that there is a population of users who scan a ranked list, starting from the top, and abandons the ranked list on a certain rank position k . Here, NCU for a query q and a system σ can be expressed as follows:

$$NCU(q, \sigma) = \sum_{k=1}^{|\mathcal{M}_q|} P_A(k) U(\sigma_q, k), \quad (2)$$

where $P_A(k)$ is the abandonment probability at rank position k (i.e. the population of users who stop at k), and $U(\sigma_q, k)$ is the utility of the ranked list σ_q at k . As we do not want to reward redundancy in VMR, we follow the approach of $R@K, \theta$ (Eq. 1) to instantiate our utility function:

$$U(\sigma_q, k) = \max_{1 \leq j \leq k} r(\sigma_q(j)). \quad (3)$$

Based on this, we can obtain *normalised cumulative max IoU (NCxIoU)* measure as follows:

$$NCxIoU(q, \sigma) := \sum_{k=1}^{|\mathcal{M}_q|} P_A(k) \max_{1 \leq j \leq k} r(\sigma_q(j)). \quad (4)$$

With VMR, we do not have any prior knowledge on $P_A(k)$. Therefore, given a SERP containing K moments, we assume that the users are uniformly distributed over the K moments: that is, that $1/K$ of the user population abandons the list at rank k ($1 \leq k \leq K$). Note that the Average Precision (AP), an NCU measure widely used in information retrieval evaluation with manual relevance assessments, assumes that the users are uniformly distributed over *all relevant* documents [30]. In the case of VMR, we consider only the top- K items (following $R@K, \theta$), and assume that each retrieved moment is at least somewhat relevant, where the degree of relevance is represented by the IoU of each moment.

Our proposed measure for VMR is also an instantiation of NCU, which we call average max IoU (AxIoU):

$$AxIoU@K(q, \sigma) := \frac{1}{K} \sum_{k=1}^K \max_{1 \leq j \leq k} r(\sigma_q(j)). \quad (5)$$

As the uniform assumption on $P_A(k)$ may not hold when with a large K , we can use a more realistic distribution for $P_A(k)$ such as the expected reciprocal rank (another NCU measure) [6], although we leave this as future work.

4.2. Interpretation of the AxIoU Measure

In this section, we describe the relationship between our proposed measure and $R@K, \theta$. We first consider the marginalisation of $R@K, \theta$ in terms of K and θ . In practical terms, because we do not have any knowledge regarding the distribution of θ for each dataset, each query, or each set of systems to be evaluated, we assume that $\theta \sim \text{Uni}(0, 1)$ and then obtain the following equation:

$$\begin{aligned} & \mathbb{E}_k \mathbb{E}_\theta \left[\frac{1}{|\mathcal{Q}|} \sum_{q \in \mathcal{Q}} R@k, \theta(q, \sigma) \right] \\ &= \mathbb{E}_k \mathbb{E}_\theta \left[\frac{1}{|\mathcal{Q}|} \sum_{q \in \mathcal{Q}} \mathbb{1} \left\{ \max_{1 \leq j \leq k} r(\sigma_q(j)) > \theta \right\} \right] \\ &= \frac{1}{|\mathcal{Q}|} \sum_{q \in \mathcal{Q}} \mathbb{E}_k \mathbb{E}_\theta \left[\mathbb{1} \left\{ \max_{1 \leq j \leq k} r(\sigma_q(j)) > \theta \right\} \right] \\ &= \frac{1}{|\mathcal{Q}|} \sum_{q \in \mathcal{Q}} \mathbb{E}_k \left[\max_{1 \leq j \leq k} r(\sigma_q(j)) \right]. \end{aligned} \quad (6)$$

In Eq (6), by assuming $\theta \sim \text{Uni}(0, 1)$, we can obtain the following:

$$\mathbb{E}_\theta \left[\mathbb{1} \left\{ \max_{1 \leq j \leq k} r(\sigma_q(j)) > \theta \right\} \right] = \max_{1 \leq j \leq k} r(\sigma_q(j)). \quad (7)$$

Because we assume a uniform distribution for k on $1 \leq k \leq K$ and $P_A(k) = 1/K$, we obtain the following:

$$\begin{aligned} (\text{RHS of Eq. (6)}) &= \frac{1}{|\mathcal{Q}|} \sum_{q \in \mathcal{Q}} \frac{1}{K} \sum_{k=1}^K \max_{1 \leq j \leq k} r(\sigma_q(j)) \\ &= \frac{1}{|\mathcal{Q}|} \sum_{q \in \mathcal{Q}} \text{AxIoU}@K(q, \sigma). \end{aligned} \quad (8)$$

That is, the mean $\text{AxIoU}@K$ can be considered as a marginalisation of the mean $R@K, \theta$ without any assumption for θ and with a weak assumption for K . The mean $R@K, \theta$ with a fixed value for each K and θ evaluates certain aspects of systems' behaviour and thus requires to examine multiple settings of the parameters for evaluation. We argue that AxIoU is a reasonable approach to avoiding the dependence on the θ threshold while considering the rank position of the best moment in a top- K ranked list.

5. Requirements for Effectiveness Measures

To evaluate the evaluation measures, we take an axiomatic approach. We first set the following requirements for the design of our VMR measure based on the properties of $R@K, \theta$ in Section 3.2: (1) It should ignore redundant moments in a ranked list, (2) it should consider the IoU value between a ranked moment and the ground truth

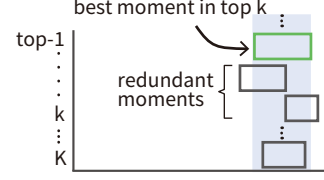


Figure 2. **INV-k** requires that a measure should be invariant to redundant moments which have smaller IoU and lower rank position than the best moment in the top- k ($1 \leq k \leq K$) ranked list.

moment, and (3) it should consider the rank position of relevant moments for evaluating top- K retrieval effectiveness of a system. We show that our Mean AxIoU satisfies all requirements while Mean $R@K, \theta$ satisfies only Requirement (1). To investigate VMR measures based on these requirements, we define two axioms for an effectiveness measure for VMR.

Invariance against Redundant Moments A measure should be unchanged to redundant moments in a ranked list. We define this requirement as the following axiom.

Axiom 1 (Invariance against Top- k Non-Best Moment (**INV-k**)). Suppose that two systems σ and σ' such that σ' differs from σ only for the k -th moment in the ranked lists for q . The measurement of σ' must not change from that of σ (i.e. $\mu(\mathcal{Q}, \sigma) = \mu(\mathcal{Q}, \sigma')$) when the k -th moment in σ' has a better IoU value than the k -th moment in σ but is not the most relevant within the top k of σ' .

Figure 2 depicts the concept of **INV-k**. $R@K, \theta$ satisfies this requirement because it utilises only the moment with the maximum IoU value in a ranked list (see Eq. (1)). AxIoU can also handle the redundant moments by inheriting the property of $R@K, \theta$. On the other hand, $\text{AP}@K, \theta$, which is a ranked retrieval measure widely adopted in computer vision [8], does not satisfy **INV-k**. Similarly, while an information retrieval measure for graded relevance such as DCG [17] would be a straightforward choice for evaluating ranked lists while avoiding the binarization by θ , it does not satisfy **INV-k** either. The formal definition of the axiom and proofs are given in our supplementary material.

Monotonicity with respect to the Best Moment The VMR measure score should monotonically increase with the maximum IoU value in a ranked list. More specifically, we require that at any rank k , the measurement based on the top- k moments of the SERP should monotonically increase with the maximum IoU observed within the top k . This requirement can be defined through the following axiom.

Axiom 2 (Strict Monotonicity for Top- k Best Moment (**MON-k**)). Suppose that two systems σ and σ' such that

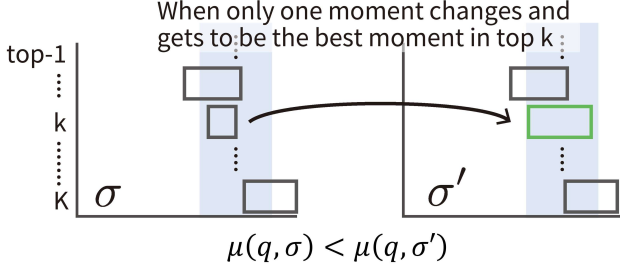


Figure 3. **MON-k** requires that a measure should be sensitive to the IoU value of the best moment in a top- k ranked list.

σ' differs from σ only for the k -th moment in the ranked lists for q . The measurement of σ' strictly increases from that of σ (i.e. $\mu(Q, \sigma) < \mu(Q, \sigma')$) when the k -th moment in σ' has a better IoU value than the k -th moment in σ and is the most relevant within the top k of σ' .

Figure 3 depicts the concept of **MON-k**. $R@K, \theta$ with a fixed parameter setting for K and θ does not satisfy this requirement. $\mu(Q, \sigma) < \mu(Q, \sigma')$ is not guaranteed since $R@K, \theta$ binarizes the relevance using θ . By contrast, ranked retrieval measures for graded relevance, such as $DCG@K$ and our $AxIoU@K$, satisfies this property because these consider the ranked position and IoU value of each moment in a ranked list. The formal definition of the axiom and proofs are provided in the supplementary material.

6. Experiments

While we analytically showed the properties of $AxIoU@K$ in terms of the axioms, we also examine the measures empirically in this section. We first investigate the agreement between the evaluation results based on the measures to confirm the compatibility of $AxIoU@K$ with $R@K, \theta$. To examine the effect of θ , we also discuss the stability of the measures with respect to change in the test data. Moreover, we demonstrate the advantages of $AxIoU@K$ as the criterion for model selection.

6.1. Experimental Setup

Datasets Following the experimental settings of Otani *et al.* [27], we utilise two popular datasets for our experiments, Charades-STA [10] and ActivityNet [5, 19]. Each dataset contains a set of manually annotated temporal regions for query-video pairs that indicate the relevant moment in a video as ground truth. Charades-STA is built upon Charades [37] and contains 9,848 videos, each of which is associated with multiple natural language sentences. The number of test queries is 3,720. ActivityNet contains 19,209 YouTube videos. Each video is associated with the captions and their temporal locations. The number of the test queries is 17,031.

Retrieval Systems for Evaluation In our experiment, we utilise multiple VMR systems to evaluate the measures; for example, we create two rankings of the systems based on two measures, and then compute the similarity of the rankings (i.e. Kendall's τ - b [1]) as agreement between the two measures. To examine each measure in a realistic setting, we employ real VMR systems trained on each dataset. Throughout this paper, we used three conventional methods, **Action-Aware Blind** (Blind) [27], **SCDM** [40] and **2DTAN** [44]. In addition, we include the variants of 2DTAN, i.e., (1) **2DTAN nonms**, a variant without Non-maximum suppression (NMS) [25], (2) **2DTAN rand**, a variant with randomisation of video frames proposed by Otani *et al.* [27], and (3) **2DTAN rand+nonms**, a variant without NMS and with randomisation.

Figure 5 compares the effectiveness of the above six systems according to different measures on Charades-STA (left) and ActivityNet (right). In each graph, the systems have been sorted by Mean $R@5, 0.5$. For Charades-STA, the $R@10, 0.3$ score of Blind (green solid line), which is a video-agnostic baseline, is almost one; as the Charades-STA dataset is a relatively easy dataset, $R@10, 0.3$ is a too relaxed measure even for Blind. This result suggests that a inappropriate choice of K and θ leads to uninformative evaluation results.

6.2. Agreement between Measures

Figure 4 shows the agreement between each pair of measures among $R@K, \theta$ ($K = 1, 5, 10, \theta = 0.3, 0.5, 0.7$) and the $AxIoU@K$ ($K = 1, 5, 10$) on Charades-STA and ActivityNet datasets, respectively. To assess the agreement between two measures, we first rank the six systems using each measure. We then compute Kendall's τ - b [1], which considers the ties in a ranking. Hereafter, we shall refer to τ - b simply as τ . A high τ value means that the rankings according to the two measures are similar [4].

In the Charades-STA dataset, $AxIoU@10$ agrees well with all instances of $R@K, \theta$ ($0.36 \leq \tau \leq 0.87$). The values of $AxIoU@K$ with different values of K agree reasonably well with one another ($0.36 \leq \tau \leq 0.73$). By contrast, different instances of $R@K, \theta$ can conflict with one another; $R@5, 0.7$ agrees well with $R@1, \theta$ ($\tau = 0.64$) whereas $R@10, \theta$ does not agree with $R@1, \theta$ instances ($-0.21 \leq \tau \leq 0.21$). Probably, the main reason for this result is that $R@K, \theta$ is rank-insensitive. On the other hand, $AxIoU@K$, which satisfies **MON-k**, aligns well with itself for different values of K . Although $R@5, 0.5$ and $R@5, 0.7$, which are popular instances of $R@K, \theta$, agree relatively well with the other $R@K, \theta$ instances ($0.21 \leq \tau \leq 0.73$ for $R@5, 0.5$ and $0.2 \leq \tau \leq 0.64$ for $R@5, 0.7$), the agreement between the two measures is $\tau = 0.60$ despite the small difference in the setting of θ ; remarkably, $AxIoU@10$ agrees with $R@5, 0.5$ and $R@5, 0.7$ with

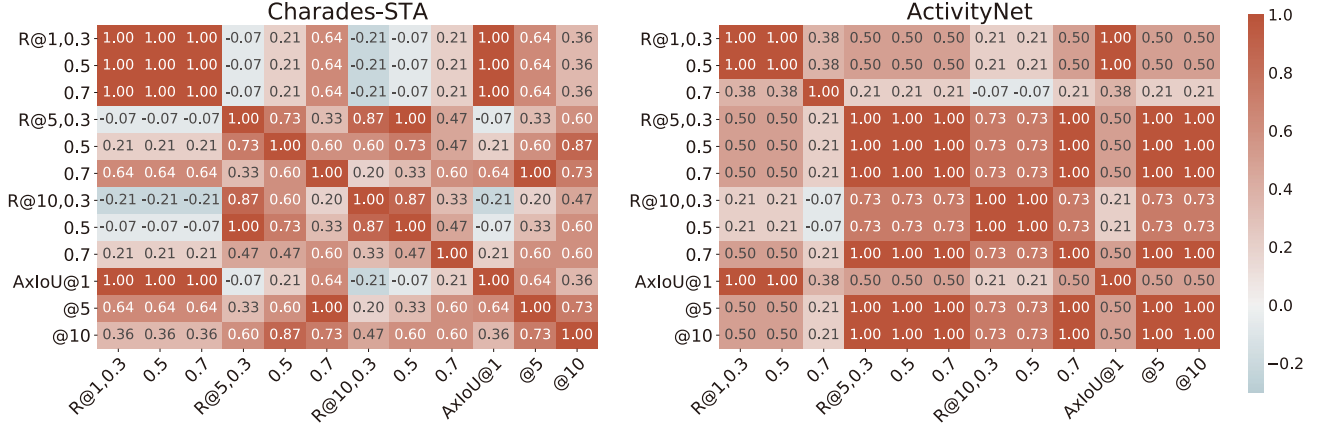


Figure 4. Agreement between two measures on Charades-STA (left) and ActivityNet (right).

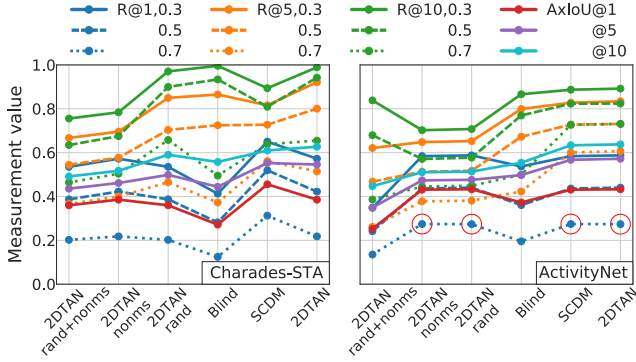


Figure 5. Effectiveness of each system on Charades-STA and ActivityNet datasets according to $R@K, \theta$ and $AxIoU@K$.

$\tau = 0.87$ and $\tau = 0.73$, respectively.

Because the ActivityNet dataset has a much larger number of test queries than that of the Charades-STA dataset, most of the measures agree well with each other. Nevertheless, $R@1, 0.7$, which is a widely adopted instance of $R@K, \theta$, does not agree with the other instances ($-0.07 \leq \tau \leq 0.38$). It is worth mentioning that $R@1, 0.7$ is highly demanding (*i.e.* requiring systems to return a highly relevant moment at rank 1). Thus, $R@1, 0.7$ lacks the sensitivity to distinguish between the systems. As shown in Fig. 5 (red circles in the right side), the scores by $R@1, 0.7$ are low for all six systems, and the scores for four out of six systems are all 0.274 (1,019/3,720). $AxIoU@10$ achieves strong agreement ($\tau \geq 0.5$) with all instances of $R@K, \theta$ except $R@1, 0.7$.

6.3. Stability of Measures against the Choice of Evaluation Data

In this section, we investigate the stability of the measures, *i.e.*, the consistency of evaluation results based on a measure on different test datasets [4]. The stability of an

effective measure against different test datasets is one of the essential properties: If a measure is unstable, the conclusion drawn for a certain test dataset may not generalise well. We evaluate each measure based on Kendall's τ -b between the system rankings based on two different subsets of a dataset as the *self-agreement* of the measure. To examine the stability with respect to the choice and size of test data we investigate the self-agreement on conjoint query set pairs with different sizes.

Figure 6 visualises the effect of reducing the size of the query subsets on self-agreement. We experimented with 5,000 trials for each query subset size. The horizontal axes represent the size of each query subset. The top graphs show the means of the self-agreement τ 's; the bottom graphs show the variances.

For Charades-STA in the first to third columns, it can be observed that, for each K , the $R@K, \theta$ instances with $\theta = 0.3, 0.5$ substantially underperform the other in terms of mean and variance of τ ; on the other hand, the $R@K, 0.7$ (red dotted line) instances are consistently stable. The $AxIoU@K$ instances outperform most of the $R@K, \theta$ instances with the same value of K whereas it performs relatively poorly with small query sets for $K = 5$. The most robust batch of measures for this dataset are $R@1, 0.7$, $R@5, 0.7$, $R@10, 0.7$, $AxIoU@1$ and $AxIoU@10$. Also for ActivityNet in the forth to fifth columns, the $AxIoU@K$ instances outperform most of the $R@K, \theta$; the $R@K, 0.7$ instances also perform well.

6.4. Stability against Label Ambiguity

In this section, we evaluate the measures in terms of the stability to label ambiguity (*i.e.* disagreement between human annotations). We generate a testing sample based on a simple noise model by following steps; (1) we consider each annotation in an original testing dataset as a low-noise sample and denote one by $(s^*, e^*) \in \mathbb{R}^2$ where s^* and e^* are the start and end points of a temporal boundary; (2) we draw

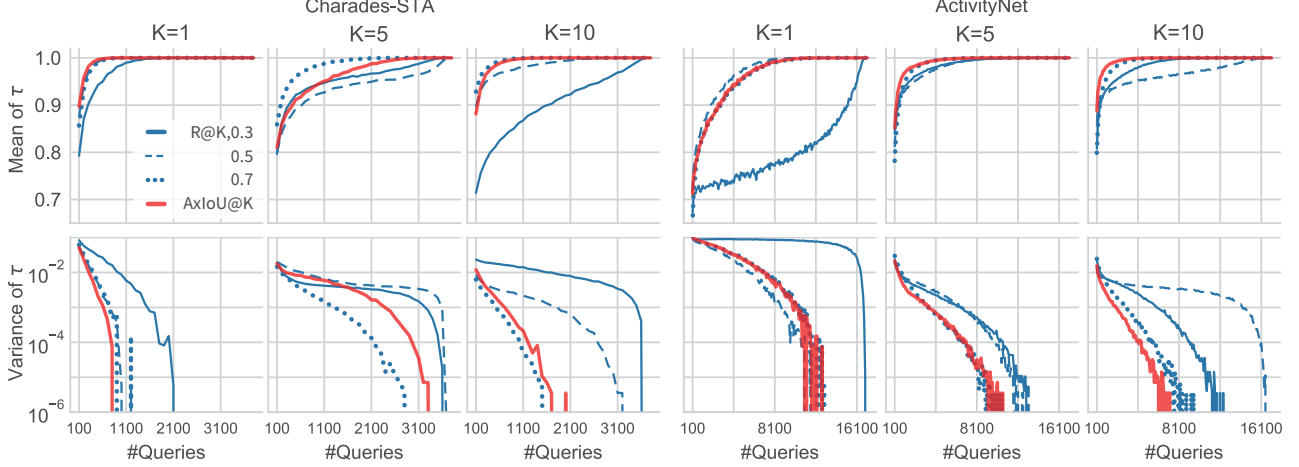


Figure 6. Effect of reducing the size of query subsets on means and variances of self-agreement for Charades-STA and ActivityNet.

a start point s by a normal distribution with mean s^* and variance β^2 ; (3) we then draw a length l by an exponential distribution with mean $e^* - s^*$; and (4) we obtain the drawn sample $(s, s + l)$ as a noisy one. For each testing sample in Charades-STA and ActivityNet, we independently draw five samples from the noise model and then create a final testing annotation by taking medians of s and $s + l$. Here, it should be noted that the variance parameter β^2 can be considered as the quality (i.e. noise level) of five raters who annotate temporal boundaries to one sample. We generate datasets with different noise levels by varying β^2 in $\{1, 2, 3, 4\}$. The IoU between the mean of the median IoU values between the original and a drawn annotation for each noise level $\{1, 2, 3, 4\}$ is respectively 0.906, 0.870, 0.835 and 0.802 for Charades-STA, and 0.846, 0.778, 0.712 and 0.650 for ActivityNet; note that, the noise levels are in a realistic range as the previously reported IoU agreement between human annotations is around 0.725 in Charades-STA [36] and 0.641 in ActivityNet [2]. We generate independent 100 testing datasets for each dataset and each noise level. To evaluate the effect of label ambiguity for each measure, we compute the root mean squared error (RMSE) between the measurements based on the original dataset and 100 noisy datasets for each of six systems used in the above experiments.

Figure 7 shows the effect of the label noise on the evaluation based on the measures. The x- and y-axes indicate the noise level and Mean RMSE for each measure. In all datasets and all K , the AxIoU instances show lower errors than the $R@K, 0.7$ instances but higher errors than $R@K, 0.3$ instances in a wide range of noise levels. In particular, the $R@K, 0.7$ instances shows severely high errors. This is because $R@K, \theta$ with large IoU threshold requires exactly localised moments thereby drastically changing evaluation results even with small perturbation in a ground truth. Therefore, the use of a large θ assumes the

low-noise condition of human annotations, which is difficult to ensure [2, 15, 27]. On the other hand, the instances of $R@K, \theta$ with $\theta = 0.3, 0.5$ show comparable or lower errors than the AxIoU instances because these $R@K, \theta$ instances ignore localisation quality.

6.5. Summary: Agreement and Stability

We demonstrated the undesirable properties of $R@K, \theta$ from various aspects; (1) the $R@K, 0.3$ and $R@1, \theta$ instances (i.e. non-demanding measures) often disagree with other $R@K, \theta$ instances (Section 6.2); (2) the $R@K, 0.3$ and $R@K, 0.5$ instances are unstable to the change of the size of a dataset (Section 6.3); and (3) the $R@K, 0.7$ instances are unstable to label ambiguity and potentially noisy (Section 6.4). By contrast, our AxIoU measure reconciles the agreement and stability while reducing the hyper-parameter θ , which is difficult to tune. Moreover, it should be noted that the cut-off parameter K for AxIoU@ K is easier to handle than that of $R@K, \theta$ as it considers the ranking quality of top- K ranked lists; the agreement between the AxIoU instances (Section 6.2) is also an evidence for this.

6.6. Model Selection

As discussed in Section 6.3, the stability with respect to the choice of test queries is vital for avoiding inconsistent evaluation on different dataset splits. This is true also in the process of model selection; when we select the best model based on a validation split and evaluate it on a test split, the measurement on the validation split should be consistent with that on the test split.

This section investigates the effectiveness of AxIoU as the criterion for model selection. To this end, we first created 640 variants of the 2DTAN system (See Section 6.1) by varying its hyper-parameters such as the learning rate and the threshold of NMS. Then, based on each instance of

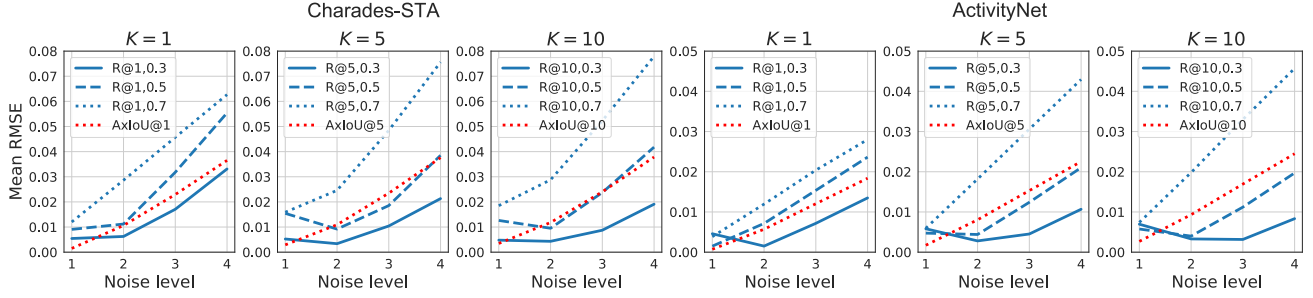


Figure 7. Effect of the ambiguity of annotated temporal regions.

$R@K, \theta$ as well as AxIoU (12 measures in total), we select the best model using the validation split. Finally, we evaluate the above 12 models using $R@K, \theta$ on the test split. As the $R@10, 0.3$ scores saturate easily (see also Figure 5), we omitted it from the test measures for visibility of figures; thus, we utilise 8 test measures in total. For each of the 8 test measures, we compute the Z-scores of the 12 models so that the average of the 12 scores equals zero.

Figures 8 (a)–(d) show the results for Charades-STA. The x-axis in each figure shows the test measure; each of the 12 lines represents a validation measure; the y-axis shows the Z-scores of “all” the 12 models for each test measure. Since each line represents a single model selected by a particular validation measure, if the line is straight and horizontal, it would imply that the validation measure is useful for effective model selection. $R@10, 0.3$ (blue line in (c)) and $R@5, 0.5$ (orange line in (b)) perform poorly as validation measures: when the models selected according to these measures are evaluated with $R@1, 0.7$ on the test data, these systems are actually the worst among the 12 systems by far. Similarly, $R@10, 0.7$ (green in (c)), $R@1, 0.3$ (blue in (a)), and AxIoU@1 (blue in (d)) perform relatively poorly: for example, when the model selected according to AxIoU@1 is evaluated with $R@10, 0.5$ on the test data, this system is one of the worst performers among the twelve. On the other hand, it can be observed that AxIoU@5, AxIoU@10 and some other $R@K, \theta$ instances such as $R@10, 0.5$ (orange in (c)) perform well: that is, the models selected based on these measures generally perform well regardless of what the test measure is. Only AxIoU@10 (green in (d)) could select a system that is above average in terms of all the test measures.

The above result is consistent to the insights obtained in Sections 6.2–6.4. However, the disagreement and instability of each $R@K, \theta$ instance is severe for model selection because we must use a single evaluation measure to determine the best model on a validation split. As we cannot know the best setting of K and θ in the validation phase, AxIoU, which is an expectation of $R@K, \theta$ (Section 4.2), is a reasonable measure for model selection.

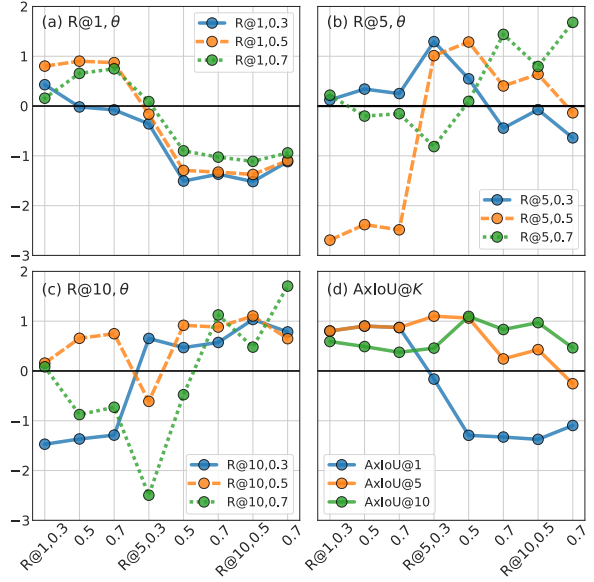


Figure 8. The effect of the validation measure for model selection on effectiveness on the test split.

7. Conclusion

In this paper, we proposed an evaluation measure, AxIoU, for video moment retrieval. AxIoU can offer consistent evaluation compared to $R@K, \theta$ without the threshold parameter θ in $R@K, \theta$, which is the main cause of the insensitivity of $R@K, \theta$. We analytically examined the properties of AxIoU through an axiomatic approach and empirically showed that AxIoU@10 can provide stable evaluation while maintaining the similarity to $R@K, \theta$ instances. We also demonstrated that AxIoU@10 is a reliable measure for model selection, even if the final test measures are $R@K, \theta$ instances. As future work, we will explore a more sophisticated distribution for abandonment position k , $P_A(k)$ [6].

8. Acknowledgement

This work was partly supported by JST CREST Grant No. JPMJCR20D3, FOREST Grant No. JPMJFR216O, and Academy of Finland project number 324346.

A. Proofs

A.1. Definition of Axioms

We formally define the definition of each axiom and each measure in this supplementary material.

Axiom 1 (Invariance against Top- k Non-Best Moment (**INV-k**)). *For any query $q \in \mathcal{Q}$ and any rank position k ($k > 1$), and for all systems σ and σ' such that σ' differs from σ only for the k -th moment in the ranked lists for q , $\mu(\mathcal{Q}, \sigma) = \mu(\mathcal{Q}, \sigma')$ holds when the k -th moments satisfy the following conditions.*

Condition A.1 (Inequality of relevance scores). The relevance scores of the k -th moment in σ_q and σ'_q satisfy $r(\sigma_q(k)) < r(\sigma'_q(k))$.

Condition A.2 (Non-maximum relevance score of the top- k moment). The k -th moment returned by system σ is less relevant than that returned by system σ' . That is, $r(\sigma'_q(k)) \leq \max_{1 \leq j < k} r(\sigma'_q(j))$.

Axiom 2 (Strict Monotonicity for Top- k Best Moment (**MON-k**)). *For any query $q \in \mathcal{Q}$ and any rank position k , and for all systems σ and σ' such that σ' differs from σ only for the k -th moment in the ranked lists for q , $\mu(\mathcal{Q}, \sigma) < \mu(\mathcal{Q}, \sigma')$ holds whenever the k -th moment satisfies Condition A.1 and the following condition.*

Condition A.3 (Maximum relevance score of the top- k moment). The k -th moment returned by σ' is the most relevant within the top k . That is, $r(\sigma'_q(k)) > \max_{1 \leq k' < k} r(\sigma'_q(k'))$ if $k > 1$.

Note that, Condition A.3 is necessary to avoid the contradiction between **INV-k** and **MON-k**.

A.2. Properties of $R@K, \theta$

Mean $R@K, \theta$, is defined as the ratio of queries for which a system successfully retrieves at least one relevant moment with a sufficient IoU with respect to threshold θ [10].

$$\begin{aligned} \text{Mean } R@K, \theta(\mathcal{Q}, \sigma) \\ = \frac{1}{|\mathcal{Q}|} \sum_{q \in \mathcal{Q}} \mathbb{1} \left\{ \sum_{k=1}^K \mathbb{1} \{r(\sigma_q(k)) > \theta\} > 0 \right\}. \end{aligned} \quad (9)$$

Property 1. *Mean $R@K, \theta$ does not satisfy **MON-k** (Axiom 2).*

Proof. For two systems σ and σ' such that σ' differs from σ only for k -th moment in the ranked list for q , the difference

of the measurements can be expressed as follows:

$$\begin{aligned} \text{Mean } R@K, \theta(\mathcal{Q}, \sigma) - \text{Mean } R@K, \theta(\mathcal{Q}, \sigma') \\ = \frac{1}{|\mathcal{Q}|} \mathbb{1} \{ \mathbb{1} \{r(\sigma_q(k)) > \theta\} + C > 0 \} \\ - \frac{1}{|\mathcal{Q}|} \mathbb{1} \{ \mathbb{1} \{r(\sigma'_q(k)) > \theta\} + C > 0 \}, \end{aligned} \quad (10)$$

where $C = \sum_{1 \leq j \leq K \wedge j \neq k} \mathbb{1} \{r(\sigma_q(j)) > \theta\}$. Here, when $r(\sigma'_q(k)) \leq \theta$ holds, it also holds that $r(\sigma_q(k)) \leq \theta$ by utilising Condition A.1. Then, the k -th moments do not contribute to the measurements, $\mathbb{1} \{r(\sigma_q(k)) > \theta\} = \mathbb{1} \{r(\sigma'_q(k)) > \theta\} = 0$. Here, because $\theta \geq r(\sigma'_q(k)) \geq \max_{1 \leq j < k} r(\sigma'_q(j))$ holds by Condition A.3, there is no moment that has a sufficient relevance score in the ranked lists σ_q and σ'_q and $C = 0$ holds. Therefore, combining these and Eq. (10), when $r(\sigma'_q(k)) \leq \theta$, we obtain $\text{Mean } R@K, \theta(\mathcal{Q}, \sigma) = \text{Mean } R@K, \theta(\mathcal{Q}, \sigma')$, which proves our proposition. $\square$

This problem results from the thresholding of temporal IoUs in the measure. This leads to the information loss of the retrieval effectiveness by binarizing the relevance score of moments and thus to the insensitivity of the measure. Property 1 suggests that the measure may ignore the improvement of systems when utilising a large value of θ .

Remarkably, $R@K, \theta$ obviously does not satisfy **MON-k** even with assuming $r(\sigma_q(k)) > \theta$ in the case of $K > 1$; when $r(\sigma_q(j)) > \theta$ holds for any rank position j ($1 \leq j \leq K \wedge j \neq k$), $C \geq 1$ in Eq. (10) holds, and thus $\text{Mean } R@K, \theta(\mathcal{Q}, \sigma) - \text{Mean } R@K, \theta(\mathcal{Q}, \sigma') = (1/|\mathcal{Q}|)(1 - 1) = 0$ holds. Therefore, setting a small value of θ , it also leads to information loss.

Property 2. *Mean $R@K, \theta$ satisfies **INV-k** (Axiom 1).*

Proof. Because we may assume that Condition A.2 holds, when $r(\sigma'_q(k)) > \theta$, there is at least one moment in a position j that satisfies $r(\sigma'_q(j)) \geq r(\sigma'_q(k)) > \theta$, and thus, $C \geq 1$ holds in Eq. (10). When $r(\sigma'_q(k)) \leq \theta$, the k -th moment does not contribute to the measurement, and $\mathbb{1} \{r(\sigma'_q(k)) > \theta\} = 0$ holds. Therefore, by utilising Condition A.1, $r(\sigma_q(k)) \leq r(\sigma'_q(k)) \leq \theta$, we have,

$$\begin{aligned} \text{Mean } R@K, \theta(\mathcal{Q}, \sigma) - \text{Mean } R@K, \theta(\mathcal{Q}, \sigma') \\ = \frac{1}{|\mathcal{Q}|} \mathbb{1} \{C > 0\} - \frac{1}{|\mathcal{Q}|} \mathbb{1} \{C > 0\} = 0, \end{aligned} \quad (11)$$

regardless of $r(\sigma'_q(k)) > \theta$ or $r(\sigma'_q(k)) \leq \theta$. Thus, we obtain $\text{Mean } R@K, \theta(\mathcal{Q}, \sigma) = \text{Mean } R@K, \theta(\mathcal{Q}, \sigma')$, which proves our proposition. $\square$

This result suggests that the thresholding and indicator function in $R@K, \theta$ play a vital role in ensuring invariance against the redundant moments in the lower rank positions.

Although these mechanisms are indispensable as the invariance is required under the problem settings of VMR, they are the main causes of information loss (See Property 1).

A.3. Properties of AP Measures

Using the average precision (AP) measure is one approach to consider the rank of relevant moments [23]. AP and Mean AP (a.k.a. mAP) can be expressed as follows:

$$\text{AP@}K, \theta(q, \sigma) := \frac{1}{K} \sum_{k=1}^K \frac{1}{k} \sum_{j=1}^k \mathbb{1}\{r(\sigma_q(j)) > \theta\}. \quad (12)$$

$$\begin{aligned} \text{Mean AP@}K, \theta(\mathcal{Q}, \sigma) \\ := \frac{1}{|\mathcal{Q}|} \sum_{q \in \mathcal{Q}} \frac{1}{K} \sum_{k=1}^K \frac{1}{k} \sum_{j=1}^k \mathbb{1}\{r(\sigma_q(j)) > \theta\}. \end{aligned} \quad (13)$$

As the AP measure is for binary relevance grades, it also requires a thresholding process for IoU values.

Property 3. Mean AP@K, θ does not satisfy **INV-k** (Axiom 1).

Proof. For two systems σ and σ' such that σ' differs from σ only for k' -th moment in the ranked list for q , the difference of the measurements can be expressed as follows:

$$\begin{aligned} \text{Mean AP@}K, \theta(\mathcal{Q}, \sigma') - \text{Mean AP@}K, \theta(\mathcal{Q}, \sigma) \\ = \frac{1}{|\mathcal{Q}|K} \sum_{k=1}^K \frac{1}{k} \sum_{j=1}^k (\mathbb{1}\{r(\sigma'_q(j)) > \theta\} - \mathbb{1}\{r(\sigma_q(j)) > \theta\}) \\ = \frac{1}{|\mathcal{Q}|K} \sum_{k=1}^K \frac{1}{k} (\mathbb{1}\{r(\sigma'_q(k')) > \theta\} - \mathbb{1}\{r(\sigma_q(k')) > \theta\}). \end{aligned} \quad (14)$$

To derive the second equality, we assume that the top- $(k' - 1)$ ranked lists of σ_q and σ'_q are identical, and the partial ranked lists from the $(k' + 1)$ -th position are also identical. When $r(\sigma'_q(k')) > \theta \geq r(\sigma_q(k'))$ holds, we have the following: $\mathbb{1}\{r(\sigma'_q(k')) > \theta\} - \mathbb{1}\{r(\sigma_q(k')) > \theta\} = 1 - 0 = 1$. Therefore, we can obtain the following:

$$\begin{aligned} \text{Mean AP@}K, \theta(\mathcal{Q}, \sigma') - \text{Mean AP@}K, \theta(\mathcal{Q}, \sigma) \\ = \frac{1}{|\mathcal{Q}|K} \sum_{k=1}^K \frac{1}{k} > 0 \\ \iff \text{Mean AP@}K, \theta(\mathcal{Q}, \sigma') > \text{Mean AP@}K, \theta(\mathcal{Q}, \sigma). \end{aligned} \quad \square$$

AP cannot handle the redundant moments in a ranked list because each top- K ranked relevant moment contributes to the measurement as an equally relevant one; in other words, AP is concerned with the number of the relevant moments in a ranked list. It suggests that a system without NMS can unfairly take an advantage in the evaluation based on AP.

Property 4. Mean AP@K, θ does not satisfy **MON-k** (Axiom 1).

Proof. In Eq. (14), when $\theta \geq r(\sigma'_q(k')) > r(\sigma_q(k'))$ and Condition A.3 hold, $\mathbb{1}\{r(\sigma'_q(k')) > \theta\} - \mathbb{1}\{r(\sigma_q(k')) > \theta\} = 0 - 0 = 0$. Therefore, we obtain the following:

$$\begin{aligned} \text{Mean AP@}K, \theta(\mathcal{Q}, \sigma') - \text{Mean AP@}K, \theta(\mathcal{Q}, \sigma) \\ = 0 \\ \iff \text{Mean AP@}K, \theta(\mathcal{Q}, \sigma') = \text{Mean AP@}K, \theta(\mathcal{Q}, \sigma). \end{aligned} \quad \square$$

Although AP is rank-sensitive, it has the threshold θ as in R@K, θ and can ignore the improvement of IoU values of relevant moments.

A.4. Properties of DCG-type Measures

The naïve approach to remove thresholding parameter θ while considering the rank positions of relevant moments is to utilise the measures for multiple relevance grades, such as normalised discounted cumulative gain (nDCG) [17] because an IoU value can be considered as a continuous relevance score. A DCG-type measure can be expressed as follows:

$$\text{DCG@}K(q, \sigma) := \sum_{k=1}^K \frac{g(r(\sigma(q)_i))}{d(k)}. \quad (15)$$

$$\text{Mean DCG@}K(\mathcal{Q}, \sigma) := \frac{1}{|\mathcal{Q}|} \sum_{q \in \mathcal{Q}} \sum_{k=1}^K \frac{g(r(\sigma(q)_i))}{d(k)}, \quad (16)$$

where $g(r) \geq 0$ and $d(k) > 0$ denotes the gain and discounting functions that are strictly monotonically increasing with respect to the relevance score r and rank position k , respectively. However, any DCG-type measure obviously does not satisfy **INV-k** (Axiom 1).

Property 5. Mean DCG@K does not satisfy **INV-k** (Axiom 1).

Proof. For any query $q \in \mathcal{Q}$ and any rank position k , and for all systems σ and σ' such that σ' differs from σ for only the k -th moment in the ranked lists for q , we can obtain the following equation.

$$\begin{aligned} \text{Mean DCG@}K(\mathcal{Q}, \sigma') - \text{Mean DCG@}K(\mathcal{Q}, \sigma) \\ = \frac{1}{|\mathcal{Q}|} \sum_{q \in \mathcal{Q}} \sum_{j=1}^K \frac{g(r(\sigma'_q(j)))}{d(j)} - \frac{1}{|\mathcal{Q}|} \sum_{q \in \mathcal{Q}} \sum_{j=1}^K \frac{g(r(\sigma_q(j)))}{d(j)} \\ = \frac{1}{|\mathcal{Q}|} \frac{g(r(\sigma'_q(k))) - g(r(\sigma_q(k)))}{d(k)}. \end{aligned} \quad (17)$$

Based on Condition A.1, $r(\sigma_q(k)) < r(\sigma'_q(k))$ and the strict monotonicity of the gain function $g(\cdot)$, we have the following: $\text{Mean DCG@}K(\sigma') - \text{Mean DCG@}K(\sigma) > 0$, which completes the proof. $\square$

Because $\text{DCG}@K$ is defined as the sum of element-wise discounted gain values of moments in a ranked list, it cannot handle the redundant moments in a ranked list appropriately. Moreover, it is difficult to utilise DCG-type measures with normalisation (nDCG) for VMR as the definition of the ideal list is not trivial owing to **INV-k**; thus, $\text{DCG}@K$ is under-normalised and can take a large value for a single query.

Property 6. *Mean DCG@K satisfies **MON-k** (Axiom 1).*

Proof. In Eq. (17), by the strict monotonicity of the gain function $g(\cdot)$ and non-negativity of the discount function $d(\cdot)$, $g(r(\sigma'_q(k))) - g(r(\sigma_q(k))) > 0$, and thus, $\text{Mean DCG}@K(\sigma') - \text{Mean DCG}@K(\sigma) > 0$ always holds, which completes the proof. $\square$

$\text{DCG}@K$ is thresholding-free and rank-sensitive, and thereby satisfies **MON-k**; however, due to these properties, it does not satisfy **INV-k**.

By considering the properties of $\text{R}@K$, θ and $\text{DCG}@K$, it is challenging for conventional information retrieval measures to satisfy **INV-k** and **MON-k** simultaneously.

A.5. Properties of AxIoU Measure

In this section, we demonstrate that AxIoU is thresholding-free and rank-sensitive while satisfying both **INV-k** and **MON-k**. AxIoU can be expressed as follows:

$$\text{AxIoU}@K(q, \sigma) := \frac{1}{K} \sum_{k=1}^K \max_{1 \leq j \leq k} r(\sigma_q(j)), \quad (18)$$

$$\text{Mean AxIoU}@K(\mathcal{Q}, \sigma) := \frac{1}{|\mathcal{Q}|} \sum_{q \in \mathcal{Q}} \frac{1}{K} \sum_{k=1}^K \max_{1 \leq j \leq k} r(\sigma_q(j)). \quad (19)$$

Property 7. *Mean AxIoU@K satisfies **INV-k** (Axiom 1).*

Proof. For two systems σ and σ' such that σ' differs from σ only for the k' -th moment in the ranked list for q ,

$$\begin{aligned} & \text{Mean AxIoU}@K(\mathcal{Q}, \sigma') - \text{Mean AxIoU}@K(\mathcal{Q}, \sigma) \\ &= \frac{1}{|\mathcal{Q}|K} \sum_{k=k'}^K \left(\max_{1 \leq j \leq k} r(\sigma'_q(j)) - \max_{1 \leq j \leq k} r(\sigma_q(j)) \right) \end{aligned} \quad (20)$$

By utilising $r(\sigma'_q(k')) \leq \max_{1 \leq j < k'} r(j)$ and $r(\sigma_q(k')) < r(\sigma'_q(k'))$, it holds that $r(\sigma_q(k')) \leq \max_{1 \leq j < k'} r(\sigma_q(j))$ because the top- $(k' - 1)$ lists of σ and σ' are identical. Therefore, $\max_{1 \leq j \leq k} r(\sigma'_q(j)) = \max_{1 \leq j \leq k'} r(\sigma_q(j))$ holds. In addition, because the partial ranked lists of σ and σ' from the $(k' + 1)$ -th position are identical, we have

$\max_{1 \leq j \leq k} r(\sigma'_q(j)) = \max_{1 \leq j \leq k} r(\sigma_q(j))$ for any position $k (1 \leq k \leq K)$. Therefore, we have the following:

$$\begin{aligned} & \sum_{k=k'}^K \left(\max_{1 \leq j \leq k} r(\sigma'_q(j)) - \max_{1 \leq j \leq k} r(\sigma_q(j)) \right) = 0 \\ & \iff \text{Mean AxIoU}@K(\mathcal{Q}, \sigma') = \text{Mean AxIoU}@K(\mathcal{Q}, \sigma). \end{aligned}$$

$\square$

AxIoU determines the contribution of the relevant moments in a ranked list by comparing the relevance of these moments. Therefore, it can handle the redundant moments without any binarisation and thresholding processes.

Property 8. *Mean AxIoU@K satisfies **MON-k** (Axiom 2).*

Proof. For two systems σ and σ' such that σ' differs from σ only for k' -th moment in the ranked list for q , the evaluation measures can be expressed as follows:

$$\begin{aligned} & \text{Mean AxIoU}@K(\mathcal{Q}, \sigma') - \text{Mean AxIoU}@K(\mathcal{Q}, \sigma) \\ &= \frac{1}{|\mathcal{Q}|K} \sum_{k=k'}^K \left(\max_{1 \leq j \leq k} r(\sigma'_q(j)) - \max_{1 \leq j \leq k} r(\sigma_q(j)) \right) \\ &= \frac{1}{|\mathcal{Q}|K} \left(r(\sigma'_q(k')) - \max_{1 \leq j \leq k'} r(\sigma_q(j)) \right. \\ & \quad \left. + \sum_{k=k'+1}^K \left(\max_{1 \leq j \leq k} r(\sigma'_q(j)) - \max_{1 \leq j \leq k} r(\sigma_q(j)) \right) \right) \end{aligned} \quad (21)$$

In the second equality, we utilised $r(\sigma'_q(k')) = \max_{1 \leq j \leq k'} r(\sigma'_q(j))$ to derive the first term in the right hand side. By utilising $r(\sigma'_q(k')) > \max_{1 \leq j < k'} r(\sigma'_q(j))$ and $r(\sigma'_q(k')) > r(\sigma_q(k'))$, $r(\sigma'_q(k')) - \max_{1 \leq j \leq k'} r(\sigma_q(j)) > 0$ holds in the right hand side of the second equality. For the second term, because we may assume that the partial ranked lists of σ and σ' from the $(k' + 1)$ -th position are identical, $\max_{1 \leq j \leq k} r(\sigma'_q(j)) \geq \max_{1 \leq j \leq k} r(\sigma_q(j))$ holds for any position $k (k' < k \leq K)$. Thus, the following inequality holds.

$$\begin{aligned} & r(\sigma'_q(k')) - \max_{1 \leq j \leq k'} r(\sigma_q(j)) > 0 \\ & \wedge \sum_{k=k'+1}^K \left(\max_{1 \leq j \leq k} r(\sigma'_q(j)) - \max_{1 \leq j \leq k} r(\sigma_q(j)) \right) \geq 0 \\ & \iff \text{Mean AxIoU}@K(\mathcal{Q}, \sigma') > \text{Mean AxIoU}@K(\mathcal{Q}, \sigma), \end{aligned}$$

which completes the proof. $\square$

As a summary, AxIoU reflects the rank positions of the relevant moments in a ranked list as in AP and considers IoU values as in DCG, while it can handle the redundant moments as in $\text{R}@K$, θ .

B. Analysis of Number of Tied Results

To demonstrate the behaviours of $R@K, \theta$ and $AxIoU@K$, we investigate the number of queries for which the 6 systems have the exactly same score for each VMR measure; we define the ratio of such queries in all test queries as *all-tied query ratio* of a measure. Figure 9 shows the all-tied query ratio of each measure on Charades-STA and ActivityNet, respectively. From Figure 9, the $R@K, \theta$ instances with relaxed or demanding settings, such as $R@10, 0.3$, $R@1, 0.7$, show higher all-tied query ratios than the other instances. It indicates that these measures cannot distil any information from the evaluation results based on a large number of queries. $R@5, 0.7$ performs well in both Charades-STA and ActivityNet. On the other hand, the $AxIoU@K$ instances show substantially lower all-tied query ratios for $K = 1, 5, 10$. It is remarkable that, with a larger K , $AxIoU@K$ performs well whereas $R@K, \theta$ with $\theta = 0.3, 0.5$ becomes worse. Probably, it is because $AxIoU@K$ can leverage the information of the lower positions in ranked lists owing to its rank-sensitivity, whereas $R@K, \theta$, which is a set retrieval measure, becomes insensitive when with a large K and requires a large θ to detect the difference of systems. This suggests that the setting of θ is rather difficult when K is large such as in the TVR dataset [20]; an extremely large θ may be required although it can make difficult queries uninformative.

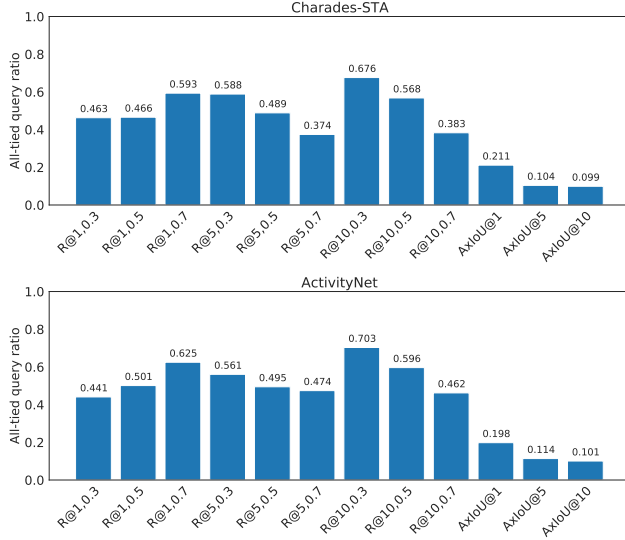


Figure 9. All-tied query ratio of each measure.

C. On the Stability to Label Ambiguity

To show the stability to label ambiguity of the measures, we showed the behaviour of the measures through numerical experiments (Section 6.4). In this section, we discuss the effect of the IoU thresholding on the estimation stability.

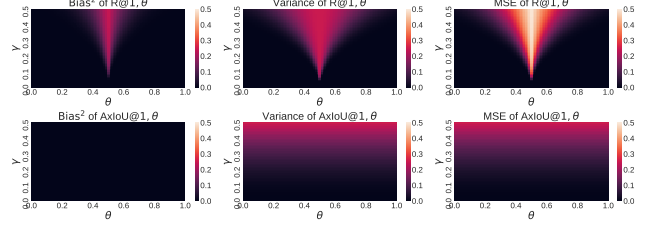


Figure 10. Effect of θ and γ on estimation errors.

We here show the case of $K = 1$ for a simple example. Let r be the IoU value of the top-1 moment for the true unobservable ground truth and $\hat{r}$ be that for the noisy ground truth. Under a Gaussian noise model $\hat{r} = r + \epsilon$, where $\epsilon \sim N(0, \gamma^2)$, noisy IoU $\hat{r}$ also obeys a normal distribution $\hat{r} \sim N(r, \gamma^2)$. The expected difference between the true and observed $AxIoU@1$ (*i.e.* bias) is obtained as $\mathbb{E}_{\hat{r}}[r - \hat{r}] = r - \mathbb{E}_{\hat{r}}[\hat{r}]$. Hence, $AxIoU@1$ is unbiased (*i.e.* $\mathbb{E}_{\hat{r}}[r - \hat{r}] = 0$) because $\hat{r}$ is an unbiased estimator of r (*i.e.* $\mathbb{E}_{\hat{r}}[\hat{r}] = r$). The variance of $AxIoU@1$ is exactly that of $\hat{r}$ (*i.e.* $\mathbb{V}[\hat{r}] = \gamma^2$).

On the other hand, the bias of $R@1, \theta$ can be obtained as

$$\begin{aligned} \mathbb{E}_{\hat{r}}[\mathbb{1}\{\hat{r} \geq \theta\} - \mathbb{1}\{r \geq \theta\}] &= \mathbb{E}_{\hat{r}}[\mathbb{1}\{\hat{r} \geq \theta\}] - \mathbb{1}\{r \geq \theta\} \\ &= P(\hat{r} \geq \theta) - \mathbb{1}\{r \geq \theta\}. \end{aligned}$$

If $\theta \leq r$ holds, because the true $R@1, \theta$ is one, the bias is then $P(\hat{r} \geq \theta) - 1 = -P(\hat{r} < \theta)$. If $\theta > r$ holds, the bias is $P(\hat{r} \geq \theta)$. Therefore, $R@1, \theta$ is statistically biased; that is, it has the error even in the expectation. Because $\mathbb{1}\{\hat{r} \geq \theta\}$ obeys the Bernoulli distribution $\text{Bern}(P(\hat{r} \geq \theta))$, The variance of $R@1, \theta$ is $P(\hat{r} \geq \theta)P(\hat{r} < \theta)$, which depends on θ and γ . Figure 10 shows the theoretical (squared) bias, variance and mean squared error (MSE) of $AxIoU@1$ and $R@1, \theta$ for different θ and γ under $r = 0.5$. We can observe that both $AxIoU@1$ and $R@1, \theta$ have large estimation errors (*i.e.* MSE) when noise level γ is large. In addition to this, $R@1, \theta$ suffers from a severe error even with small γ , particularly when θ is close to $r = 0.5$. This is an undesirable property because we often need to discriminate competitive VMR methods and thus to use θ around the boundary, which leads to estimation errors under label noise.

References

- [1] Alan Agresti. *Analysis of ordinal categorical data*, volume 656. John Wiley & Sons, 2010. 5
- [2] Humam Alwassel, Fabian Caba Heilbron, Victor Escorcia, and Bernard Ghanem. Diagnosing error in temporal action detectors. In *Eur. Conf. Comput. Vis.*, pages 256–272, 2018. 1, 2, 7
- [3] Enrique Amigó, Damiano Spina, and Jorge Carrillo-de Albornoz. An axiomatic analysis of diversity evaluation metrics: Introducing the rank-biased utility metric. In *Int. ACM SIGIR Conf. on Research and Devet. in Inform. Retrieval*, pages 625–634, 2018. 2
- [4] Chris Buckley and Ellen M Voorhees. Retrieval evaluation with incomplete information. In *Int. ACM SIGIR Conf. on Research and Devet. in Inform. Retrieval*, pages 25–32, 2004. 5, 6
- [5] Fabian Caba Heilbron, Victor Escorcia, Bernard Ghanem, and Juan Carlos Niebles. Activitynet: A large-scale video benchmark for human activity understanding. In *IEEE Conf. Comput. Vis. Pattern Recog.*, pages 961–970, 2015. 2, 5
- [6] Olivier Chapelle, Donald Metzler, Ya Zhang, and Pierre Grinspan. Expected reciprocal rank for graded relevance. In *ACM Conf. Inform. and Knowledge Management*, pages 621–630, 2009. 2, 3, 8
- [7] Victor Escorcia, Mattia Soldan, Josef Sivic, Bernard Ghanem, and Bryan Russell. Temporal localization of moments in video collections with natural language. *arXiv preprint arXiv:1907.12763*, 2019. 2
- [8] Mark Everingham, SM Ali Eslami, Luc Van Gool, Christopher KI Williams, John Winn, and Andrew Zisserman. The pascal visual object classes challenge: A retrospective. *Int. J. Comput. Vis.*, 111(1):98–136, 2015. 4
- [9] Hui Fang, Tao Tao, and Chengxiang Zhai. Diagnostic evaluation of information retrieval models. *ACM Trans. Inform. Syst.*, 29(2):1–42, 2011. 2
- [10] Jiyang Gao, Chen Sun, Zhenheng Yang, and Ram Nevatia. Tall: Temporal activity localization via language query. In *IEEE Conf. Comput. Vis. Pattern Recog.*, pages 5267–5275, 2017. 1, 2, 3, 5, 9
- [11] Junyu Gao and Changsheng Xu. Fast video moment retrieval. In *Int. Conf. Comput. Vis.*, pages 1523–1532, 2021. 2
- [12] Soham Ghosh, Anuva Agarwal, Zarana Parekh, and Alexander Hauptmann. ExCL: Extractive clip localization using natural language descriptions. In *Conf. the North American Ch. the Assoc. Comput. Linguistics: Human Language Tech.*, pages 1984–1990, Minneapolis, Minnesota, June 2019. Association for Computational Linguistics. 2
- [13] David Hall, Feras Dayoub, John Skinner, Haoyang Zhang, Dimity Miller, Peter Corke, Gustavo Carneiro, Anelia Angelova, and Niko Sünderhauf. Probabilistic object detection: Definition and evaluation. In *Winter Conf. Apl. Comput. Vis.*, 2020. 2
- [14] Dongliang He, Xiang Zhao, Jizhou Huang, Fu Li, Xiao Liu, and Shilei Wen. Read, watch, and move: Reinforcement learning for temporally grounding natural language descriptions in videos. In *AAAI*, volume 33, pages 8393–8400, 2019. 2
- [15] Lisa Anne Hendricks, Oliver Wang, Eli Shechtman, Josef Sivic, Trevor Darrell, and Bryan Russell. Localizing moments in video with natural language. In *Int. Conf. Comput. Vis.*, 2017. 1, 2, 7
- [16] Ronghang Hu, Huazhe Xu, Marcus Rohrbach, Jiashi Feng, Kate Saenko, and Trevor Darrell. Natural language object retrieval. In *IEEE Conf. Comput. Vis. Pattern Recog.*, pages 4555–4564, 2016. 2
- [17] Kalervo Järvelin and Jaana Kekäläinen. Cumulated gain-based evaluation of ir techniques. *ACM Trans. Inform. Syst.*, (4), 2002. 2, 4, 10
- [18] Gabriella Kazai, Jaap Kamps, and Natasa Milic-Frayling. An analysis of human factors and label accuracy in crowdsourcing relevance judgments. *Inform. retrieval*, 16(2):138–178, 2013. 2
- [19] Ranjay Krishna, Kenji Hata, Frederic Ren, Li Fei-Fei, and Juan Carlos Niebles. Dense-captioning events in videos. In *Int. Conf. Comput. Vis.*, 2017. 2, 5
- [20] Jie Lei, Licheng Yu, Tamara L Berg, and Mohit Bansal. TVR: A large-scale dataset for video-subtitle moment retrieval. In *Eur. Conf. Comput. Vis.*, 2020. 2, 12
- [21] Tsung-Yi Lin, Michael Maire, Serge Belongie, James Hays, Pietro Perona, Deva Ramanan, Piotr Dollár, and C Lawrence Zitnick. Microsoft coco: Common objects in context. In *Eur. Conf. Comput. Vis.*, pages 740–755. Springer, 2014. 2
- [22] Meng Liu, Xiang Wang, Liqiang Nie, Xiangnan He, Baoquan Chen, and Tat-Seng Chua. Attentive moment retrieval in videos. In *Int. ACM SIGIR Conf. on Research and Devet. in Inform. Retrieval*, pages 15–24, 2018. 1, 2
- [23] Christopher D Manning, Hinrich Schütze, and Prabhakar Raghavan. *Introduction to information retrieval*. Cambridge university press, 2008. 10
- [24] Alistair Moffat and Justin Zobel. Rank-biased precision for measurement of retrieval effectiveness. *ACM Trans. Inform. Syst.*, 27(1):1–27, 2008. 2
- [25] Alexander Neubeck and Luc Van Gool. Efficient non-maximum suppression. In *Int. Conf. Pattern Recog.*, volume 3, pages 850–855. IEEE, 2006. 5
- [26] Kemal Oksuz, Baris Cam, Emre Akbas, and Sinan Kalkan. Localization recall precision (lrp): A new performance metric for object detection. In *Eur. Conf. Comput. Vis.*, 2018. 2
- [27] Mayu Otani, Yuta Nakahima, Rahtu Esa, and Heikkilä Janne. Uncovering hidden challenges in query-based video moment retrieval. In *Brit. Mach. Vis. Conf.*, 2020. 1, 2, 5, 7
- [28] Rafael Padilla, Wesley L. Passos, Thadeu L. B. Dias, Sergio L. Netto, and Eduardo A. B. da Silva. A comparative analysis of object detection metrics with a companion open-source toolkit. *Electronics*, 10(3), 2021. 2
- [29] Michaela Regneri, Marcus Rohrbach, Dominikus Wetzel, Stefan Thater, Bernt Schiele, and Manfred Pinkal. Grounding action descriptions in videos. *Trans. the Assoc. for Comput. Linguistics*, 1:25–36, 2013. 2
- [30] Stephen Robertson. A new interpretation of average precision. In *Int. ACM SIGIR Conf. on Research and Devet. in Inform. Retrieval*, pages 689–690, 2008. 2, 3
- [31] Tetsuya Sakai. Metrics, statistics, tests. In *PROMISE winter school*, pages 116–163. Springer, 2013. 2, 3
- [32] Tetsuya Sakai and Stephen Robertson. Modelling a user population for designing information retrieval metrics. In *NTCIR*

- Workshop, 2008. 2, 3
- [33] Mark Sanderson, Monica Lestari Paramita, Paul Clough, and Evangelos Kanoulas. Do user preferences and evaluation measures line up? In *Int. ACM SIGIR Conf. on Research and Devet. in Inform. Retrieval*, pages 555–562, 2010. 2
 - [34] Fabrizio Sebastiani. An axiomatically derived measure for the evaluation of classification algorithms. In *Int. Conf. The Theory of Inform. Retrieval*, pages 11–20, 2015. 2
 - [35] Fabrizio Sebastiani. Evaluation measures for quantification: An axiomatic approach. *Inform. Retrieval J.*, 23(3):255–288, 2020. 2
 - [36] Gunnar A Sigurdsson, Olga Russakovsky, and Abhinav Gupta. What actions are needed for understanding human actions in videos? In *Int. Conf. Comput. Vis.*, pages 2137–2146, 2017. 1, 2, 7
 - [37] Gunnar A Sigurdsson, Gül Varol, Xiaolong Wang, Ali Farhadi, Ivan Laptev, and Abhinav Gupta. Hollywood in homes: Crowdsourcing data collection for activity understanding. In *Eur. Conf. Comput. Vis.*, pages 510–526. Springer, 2016. 5
 - [38] Hao Wang, Zheng-Jun Zha, Liang Li, Dong Liu, and Jiebo Luo. Structured multi-level interaction network for video moment localization via language query. In *IEEE Conf. Comput. Vis. Pattern Recog.*, pages 7026–7035, 2021. 2
 - [39] Aming Wu and Yahong Han. Multi-modal circulant fusion for video-to-language and backward. In *Int. Joint Conf. on Artificial Intelligence*, page 1029–1035, 2018. 1, 2
 - [40] Yitian Yuan, Lin Ma, Jingwen Wang, Wei Liu, and Wenwu Zhu. Semantic conditioned dynamic modulation for temporal sentence grounding in videos. In *Adv. Neural Inform. Process. Syst.*, pages 536–546, 2019. 1, 2, 5
 - [41] Yitian Yuan, Tao Mei, and Wenwu Zhu. To find where you talk: Temporal sentence localization in video with attention based location regression. In *AAAI*, volume 33, pages 9159–9166, 2019. 2
 - [42] Runhao Zeng, Haoming Xu, Wenbing Huang, Peihao Chen, Minghui Tan, and Chuang Gan. Dense regression network for video grounding. In *IEEE Conf. Comput. Vis. Pattern Recog.*, pages 10287–10296, 2020. 2
 - [43] Da Zhang, Xiyang Dai, Xin Wang, Yuan-Fang Wang, and Larry S Davis. Man: Moment alignment network for natural language moment retrieval via iterative graph adjustment. In *IEEE Conf. Comput. Vis. Pattern Recog.*, pages 1247–1257, 2019. 1, 2
 - [44] Songyang Zhang, Houwen Peng, Jianlong Fu, and Jiebo Luo. Learning 2d temporal adjacent networks for moment localization with natural language. In *AAAI*, 2020. 5