

Malicious or Benign? Towards Effective Content Moderation for Children’s Videos

Syed Hammad Ahmed Muhammad Junaid Khan

Dept. of Computer Science Dept. of Computer Science
University of Central Florida University of Central Florida
Orlando, USA Orlando, USA
hammad.ahmed@knights.ucf.edu junaid.k@knights.ucf.edu

H. M. Umer Qaisar

Royal Cyber
Karachi, Pakistan
hafiz.qaisar@royalcyber.com

Gita Sukthankar

Dept. of Computer Science
University of Central Florida
Orlando, USA
gitar@eecs.ucf.edu

Abstract

Online video platforms receive hundreds of hours of uploads every minute, making manual content moderation impossible. Unfortunately, the most vulnerable consumers of malicious video content are children from ages 1-5 whose attention is easily captured by bursts of color and sound. Scammers attempting to monetize their content may craft malicious children’s videos that are superficially similar to educational videos, but include scary and disgusting characters, violent motions, loud music, and disturbing noises. Prominent video hosting platforms like YouTube have taken measures to mitigate malicious content on their platform, but these videos often go undetected by current content moderation tools that are focused on removing pornographic or copyrighted content. This paper introduces our toolkit (**Malicious or Benign**) for promoting research on automated content moderation of children’s videos. We present 1) a customizable annotation tool for videos, 2) a new dataset with difficult to detect test cases of malicious content and 3) a benchmark suite of state-of-the-art video classification models.

Introduction

Online Video Platforms (OVP) facilitate the sharing of videos from a wide variety of genres; YouTube and TikTok are currently the top video hosting platforms with over 2.2 and 1 billion active users, respectively (Doyle 2022). In addition to these dedicated OVPs, social networks like Facebook, Twitter, and Instagram also allow users to post videos. According to Ruby (2022), 720,000 hours of video are uploaded to YouTube daily to be consumed by users from a wide range of demographics including children under 12. Children’s videos can be prime targets for content creators seeking to monetize their work; out of the ten most watched videos on YouTube, five are children’s cartoon videos (Ceci 2022). One of these videos was disliked by many parents who considered it to be “hypnotizing” and “creepy” (Sinelschikova 2020).

There are many YouTube channels, such as Number Blocks, that use cartoon characters to deliver tailored educational content in order to improve the communication and reasoning skills of children ages 1-5. By designing videos that are superficially similar in style, scammers can get their videos recommended based based on similarity of titles,

hashtags, thumbnails, or meta-data. *Malicious* videos intentionally include audiovisual features which captivate the viewer’s attention such as rapid repetitive motion, disgusting and scary characters, disturbing noise, and loud sounds. Kaushal et al. (2016) note that a child who selects a recommended malicious video is more prone to continue watching a series of similar videos based on the video recommender system. This problem is of grave concern since 1) these children are unable to distinguish the appropriateness of the video content by themselves and 2) their parents are unable to constantly supervise their viewing (Ovide 2021).

Overexposure to malicious videos is a significant risk factor for the cognitive development of preschool children. Psychological studies conclude that overexposure to cartoons with fast motions may lead to a deterioration in the performance of commonplace tasks (Lillard and Peterson 2011). Similarly, exposure to loud noises inhibits a child’s brain development especially reading, writing and comprehension skills (Klatte, Bergström, and Lachmann 2013). Violent cartoons can directly affect the behavior of pre-school children, causing aggression and anxiety (Hapkiewicz 1979). An emerging genre of video content actively posted on OVPs include video game walkthroughs and gameplays. Although these are not explicitly malicious, gameplay videos often contain similar fast-motion and violent activities which may jeopardize the cognitive development of preschoolers.

YouTube is trying its best to mitigate the availability of these videos to kids. They introduced the “made for kids” flag which must be set by the video uploader in compliance with the Children’s Online Privacy Protection Act (COPPA). YouTube also checks uploaded video content using AI techniques to moderate the videos. Although many malicious videos are automatically removed by the YouTube automation, there have been instances where the platform overruled the creator’s flag and mistakenly set videos as having been made for kids even when marked as inappropriate by the creator (Amadeo 2022). Only videos marked as “made for kids” can be viewed on the kid-friendly application launched by YouTube called YouTube Kids (YTK); although this is very useful, malicious content remains viewable for long periods of time on YTK. Figure 1 shows an example of a video still available on YTK that includes violent behaviors executed by the Number Blocks characters.

A variety of approaches have been used to tackle this

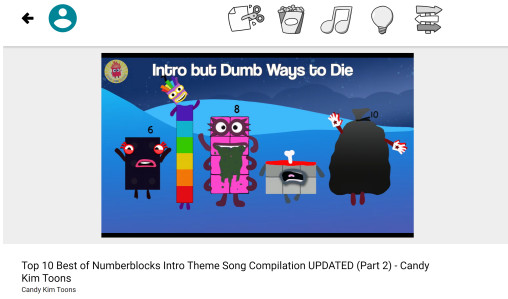


Figure 1: An example of a malicious video still viewable on YouTube Kids that is superficially similar to Number Blocks.

problem including leveraging user comments (Alshamrani et al. 2021), metadata (Gkolemi et al. 2022; Papadamou et al. 2021) and video features (Yousaf and Nawaz 2022; Singh et al. 2019; Ishikawa, Bollis, and Avila 2019); others have employed multimodal techniques rather than relying on a single data source (Chuttur and Nazurally 2022; Papadamou et al. 2021; Singh et al. 2019; Kaushal et al. 2016). We focus on analyzing and annotating subtle audio-visual features present in cartoon videos which are not conducive for the overall mental and physical development of toddlers and pre-school kids. This paper makes the following contributions:

1. Identifying a set of malicious features that are harmful to preschool children but remain unaddressed.
2. Publishing a labeled dataset (**Malicious or Benign**) of cartoon video clips. For each video clip, the presence or absence of each feature was marked by the annotators. We will make our MOB dataset publicly available to stimulate further research in this area.
3. Providing a user-friendly web-based video annotation tool which can easily be customized and used for video classification tasks with any number of ground truth classes.
4. A thorough analysis of baseline results for a few SOTA video classification benchmarks.

Background

Age Groups of Children

We assume that children have an age range from newborn to 17 years, after which one is considered as an adult. Children’s age groups should be clearly defined when considering appropriate or inappropriate content. In terms of the content available on OVPs, a video considered appropriate for a specific age group may not be suitable for another, and vice versa. Age groups are defined as follows: infants (newborn-1 year old), toddlers (1-3 years), pre-schoolers (4-5 years), middle childhood (6-11 years), and teenagers (12-17 years old) (CDC 2019). The age groups most vulnerable to malicious video content are toddlers and preschoolers, whom we focus on in this paper.

Malicious or Benign Classification

Malicious This paper assumes that video content which may not be suitable for viewing by toddlers and pre-schoolers includes a set of clearly defined, trivial and intuitive features as well as some complex and subtle audio and video features. The former includes forms of obscenity and violence on a higher level e.g. nudity, gore, etc. while the latter includes elements such as fast repetitive motion, loud music, disgusting and scary characters, smashing people or things, forms of aggression, loud music, screaming or shouting, gunshots and explosions. Furthermore, most anime are made for kids 10 and above, but a few anime series which are widely viewed by kids may contain “fantasy violence”. Similarly, certain video game gameplays are not appropriate as they may contain physical violence, strong language, drug and substance use, etc. Videos which depict famous cartoon characters dressed inappropriately and performing strange and obscene activities, as were propagated during Elsgate, are also considered unsuitable for toddlers and pre-schoolers, and possibly for older kids as well.

Benign Educational videos and videos of nursery rhymes are usually considered to be appropriate for the toddlers and pre-schoolers; in fact some experts recommend letting kids watch them for a limited number of hours (CDC 2019). Benign videos are characterized by a slower tempo, softer music or sound effects, moderate-paced motions, and soft-toned conversations. Most importantly, benign video content should not contain any indicators of malicious content as discussed earlier in this section. The absence of all malicious features in a video qualifies it to be marked as benign by our annotators.

Privacy and Rating Standards

Children’s Online Privacy Protection Act (COPPA) US Federal Trade Commission’s (FTC) COPPA clearly defines strict privacy and content type guidelines for online content publishers specifically for 13-year-old kids or younger. In 2019, COPPA added more clauses specifically for YouTube to ensure kids watch safe content.

Video Content Rating Standards The Motion Picture Association (MPA) regulates movie content ratings. It provides a comprehensive rating system where the “G” (general audience) rating category is considered suitable for kids. Similarly, for video games, the ESRB (Entertainment Software Rating Board) provides well-defined rating categories and content features. According to ESRB, the appropriate video game content for kids younger than 5 years is the “E” (everyone) category. Pan-European Game Information (PEGI) is another such regulatory body devoted to video game ratings. Since video game gameplays and walkthroughs are a common genre on OVPs, we have included their rating guidelines to help us identify malicious videos. The YouTube Kids platform includes video game gameplay videos for games which are rated as not appropriate for kids according to these rating standards. Gameplay videos for Super Mario 3D All-Stars, rated PEGI-7 by PEGI for including frightening or violent content, are still viewable on YouTube Kids platform for kids of all age groups.

Related Work

Many solutions have been proposed to the problem of detecting video content that is inappropriate for children. A general dichotomy between these techniques is whether they rely on language classification models and leverage video meta-data, user comments, and sub-titles as inputs (Gkolemi et al. 2022; Papadamou et al. 2021; Alshamrani et al. 2021), use the video stream alone (Yousaf and Nawaz 2022; Singh et al. 2019; Ishikawa, Bollis, and Avila 2019), or use a multimodal approach that combines both vision and language cues (Chuttur and Nazurally 2022; Le, Tandon, and Oinar 2022; Tahir et al. 2019; Kaushal et al. 2016).

Only a small set of prior work has focused on cartoon videos (Chuttur and Nazurally 2022; Kaushal et al. 2016; Ishikawa, Bollis, and Avila 2019), whereas most other authors have studied videos with human characters, movie clips, news, and sports. Elsgate (Placido 2017), the misuse of popular cartoon characters to create clickbait content, triggered a surge of interest in content moderation for cartoons. Ishikawa, Bollis, and Avila (2019) proposed a fusion model where they extract features from frames and the related motion vectors. They focus specifically on Elsgate videos and have published a related dataset.

Recurrent neural networks are a popular approach for identifying inappropriate videos, and we utilize an LSTM as a benchmark. Yousaf and Nawaz (2022) classify cartoon videos as safe or unsafe using an LSTM based model with around 96% accuracy. KidsGUARD (Singh et al. 2019) used an LSTM autoencoder to classify four classes of video: violent, sexual, both, and safe.

There is no single labeling strategy that is uniformly agreed upon for classifying inappropriate content. Papadamou et al. (2021) proposed a meta-data based solution for identifying inappropriate videos. They presented results for both binary and multi-class scenarios in which the latter consists of four classes: suitable, disturbing, restricted, and irrelevant. In their work, gaming videos were marked as irrelevant whereas it is an important genre of kids videos which may contain inappropriate content. Our dataset also includes these gameplay videos.

Gkolemi et al. (2022) focus on a slightly different problem: instead of identifying malicious videos they focus on detecting channels which post malicious videos. An inherent assumption in their work is that if even one disturbing video is detected, all videos posted by that channel are flagged as disturbing, making it vulnerable to false positives.

Le, Tandon, and Oinar (2022) introduced the Samba framework which is a multimodal solution that uses subtitles as input in addition to video metadata. Since neither music nor sound effects are depicted in the subtitles, a video which has malicious audio content will still be classified as safe by Samba. We manually checked a few videos from the Samba dataset and found that most of the inappropriate content was not cartoons and included news, Hollywood movies, and TV shows.

Vijaykrishnan et al. (2018) performed a comprehensive survey on the impact of cartoons on kids; they found that children are more enticed by cartoon characters vs. non-cartoon characters. Since less attention has been devoted to

cartoons, we focus on cartoon videos only and have collected a dataset of malicious cartoon videos that are difficult to detect. Chuttur and Nazurally (2022) proposed a multimodal solution using text (user comments and captions), along with image data. As images do not capture the motion related features it is likely that malicious spatio-temporal patterns will be ignored by their method.

Liu et al. (2022) try to address the problem of finding inappropriate advertisements shown on appropriate videos. They only present a statistical analysis of the data gathered and did not implement a machine learning classifier. The analysis performed by Kaushal et al. (2016) reveals that malicious content creators have strong connections among each other and they promote each other's contents by commenting, liking, and sharing playlists. Tahir et al. (2019) propose a solution which distinguishes between original and fake cartoon videos. They also further identify the malicious fake cartoons as being violent or explicit.

To summarize, this paper tries to address the following general limitations in previous work: 1) no psychological studies were used in prior work to support why the content was unsuitable for kids of ages 1-5; 2) annotators were directly asked about their opinion whether the videos are appropriate or inappropriate for young children which makes the dataset more subjective; 3) spatio-temporal audiovisual features that are not conducive to children's mental and physical health are omitted from the malicious examples in other datasets. We hope that our annotation tool, dataset, and benchmarks will be a helpful resource for authors seeking to benchmark their video content moderation approaches.

Methodology

The MOB Dataset

Our dataset¹ focuses on cartoon videos which contain activities or events not suitable for viewing by children of ages 1-5 (toddlers and pre-schoolers); examples include rapid repetitive motions, scary or disgusting characters, violent actions (e.g. hitting, smashing, biting, kicking), obscene appearances, loud music or noise, screaming, shouting, explosion and gunshot sounds, and use of any offensive language. Our dataset contains a more diverse set of cartoon videos and video game gameplays than has been previously publicly released.

Relevant Datasets There are only a few existing datasets suitable for studying this research problem. From among these, limited datasets have been published publicly. Papadamou et al. (2021) and Singh et al. (2019) provide authorized access to their datasets upon email request. To the best of our knowledge, only Ishikawa, Bollis, and Avila (2019) focus specifically on cartoon videos. However, they only include Elsgate related videos and have not extended their analysis to other types of malicious cartoon content which we have tried to highlight. Liu et al. (2022) address the problem of identifying inappropriate advertisements that are watched by children when shown during kid-appropriate

¹<https://github.com/syedhammadahmed/mob>

videos; their dataset has a list of inappropriate advertisements instead of videos. We have used their list of appropriate channels to generate seed videos for our dataset. While Gkolemi et al. (2022) and Le, Tandon, and Oinar (2022) share their dataset of malicious and benign videos, they include non-cartoon YouTube videos. For instance news, sports, Hollywood movie clips, etc. are included as part of the dataset as being malicious and not suitable for kids. Furthermore, a number of the videos in these datasets have been removed by YouTube and hence an up-to-date dataset is required. Most importantly, none of the datasets have annotated the existence of absence of specific malicious audiovisual features which have been identified as harmful by psychologists.

Spatio-temporal Audiovisual Features The spatio-temporal, audiovisual features that we have emphasized in our dataset have been left unconsidered whereas their impact on toddlers’ and preschoolers’ physical and mental development is significant. These features entail both audio and video contents and can be enumerated as follows:

Video	Audio
fast repetitive motions	loud music/noise
scary/disgusting appearance	screaming or shouting
hurting/destruction/killing activity	explosion or gunshot sounds
obscene/indecent activity	use of offensive language

Table 1: Malicious features backed by psychology studies

The Collection Process Our first step was to create a list of channels for both of our classes: 1) benign and 2) malicious. For benign, we used the channels list from Liu et al. (2022) which comprised of 50 channels. These 50 channels included around 47K videos which were extracted using the YouTube Data API. But as our scope of the problem was cartoon videos only, we performed a random non-exhaustive inspection of around 5 videos from each channel and manually sifted out 16 non-cartoon channels. From the remaining 34 benign channels around 11 videos were selected from each, resulting in 363 videos. For each video we made 10-second clips and randomly selected 5 clips. A few of the videos had less than 5 clips as they were shorter than a minute. We ended up with a total of 594 benign clips. Similarly, for the malicious class, we started with 32 channels which we gathered from a walkthrough of related videos starting with search results using keywords like “angry number blocks”, “angry mario”, “scary peppa”, etc. Although there were around 14K videos extracted from this seed pool of channels, we selected the top 680 videos according to the view count and ended up with 1281 10-second clips of malicious videos. The statistics of seed videos before annotation are summarized in Table 2. Table 3 includes the dataset details after the annotation was performed on the seed videos. We excluded anime clips for consideration as well as videos that are neither cartoon videos nor videogame gameplay.

Class	Channels	Videos	Clips
Benign	50	363	594
Malicious	32	680	1281
Total	82	1043	1875

Table 2: Statistics of MOB dataset seed videos.

Class	Clips
Benign	1046
Malicious	519
Excluded	310
Total	1875

Table 3: Statistics of annotated MOB dataset.

Data Preparation After collecting the list of video-IDs of YouTube videos, we downloaded the videos using the pytube API at a frame rate of 25. Then for each video we made clips of duration 10 seconds using the moviepy API which is a wrapper for the ffmpeg framework. The nomenclature for a video clip # x for a video-id y was kept as y_x. Similarly, we selected 5 random clips from the set of clips generated from a video and for each of the clips extracted frames. As the frame rate was 25, each of the 10-second clip resulted in 250 frames. Each frame was saved in a directory bearing the name of the clip, where the frame names were simply the frame numbers ranging from 0-249 but padded to string of length 4 with *0000.jpg* being the first frame’s name whereas *0249.jpg* being the last. These frames were provided as input to the different benchmarks.

Ethics Another principle of law is “fair use” which in special cases allows others to use existing artifacts that have ownership protection, without the explicit consent of the owner (YouTube Help b). Within the US, research is one of the tasks taken into consideration as fair use (YouTube Help a). We did not gather any user’s web usage activity, personal information or any kind of demographics for YouTube, YouTube Kids, or for the annotation tool. For the annotation tool web application, we generated anonymous login credentials for each user and hence the annotator cannot be identified.

Annotation Process

Our MOB dataset was annotated using a web-application, MAWA (MOB Annotation Web Application), which was custom-built for annotating malicious videos. There is a user authentication feature to keep track of progress of each user’s annotation. However, to ensure privacy of users, we created generic credentials e.g. user1, user2, and so on. Each user was randomly assigned a pool of 125 unannotated clips from the database. The user interface is very user-friendly and easy to use with negligible training time. Users can perform annotations over multiple sessions at their convenience. The user interface is inspired by the interface of Kay et al. (2017) which was used to annotate human actions.

Annotation Methodology We enumerated the list of all possible malicious video and audio features which are unsuitable for kids of ages 1-5 to watch, as shown in Table 1. For every video clip, we formulate questions based on the features and ask the annotator whether they noticed the presence of that feature. From these set of questions we further made groups as shown in Figure 2 following the effective strategy of multi-class annotation as proposed by Deng et al. (2014). Using this strategy we have split our question set into 3 hierarchy groups:

- Group 1: to find the category of the video clip (cartoon, videogame gameplay, anime, or non-cartoon)
- Group 2: only asked if the video clip is a cartoon
- Group 3: only asked if there is audio in the clip

Questions in Group 1 belong to top hierarchy level and will be asked for each video clip. However, if there is no cartoon character in the video clip, Group 2 will be skipped as it is not categorized as a cartoon video. Similarly, if there is no audio in the video clip being annotated none of the questions in Group 3 will be asked from the annotator. Each video is viewed by one annotator; as each question depicts one feature this makes the questions simple and easy to comprehend, hence not requiring the annotator to make any inferences. Therefore, the possibility of the annotation responses being highly subjective is negligible. However, we did perform a random validation cycle after the annotation phase in which random videos were selected from the set of annotated video clips and one user performed annotation on those again. The validation accuracy was more than 90%.

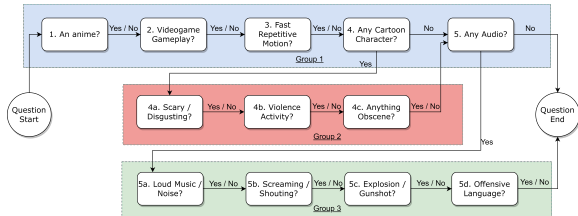


Figure 2: Annotation questions grouped together for fast annotation

Annotation Tool The MOB Annotator Web Application (MAWA) was specially developed considering the intricate requirements of the dataset annotation process. We used the following development tools and technologies: HTML, CSS, Javascript Bootstrap library, FASTAPI (Python) for Web services, and MySQL database. Each user is authenticated before the annotation tasks. For each feature’s annotation, the response is saved to the database server and hence allows user to restart if the user wishes to complete the task in multiple sessions. The source code of the current version is available by request.

MAWA comes with a highly intuitive graphical interface as shown in Figure 3. There are flexible controls available for the annotator to play, pause, rewind, view in full screen, and adjust volume for the video clips. Below the video panel the annotator can view the count of total videos annotated as

well as the total videos assigned, as well as the questions and buttons to give response. On the right side, is the instruction panel where the annotator can read the comprehensive self-explanatory instructions to become acquainted with the environment and understand what action each button performs.

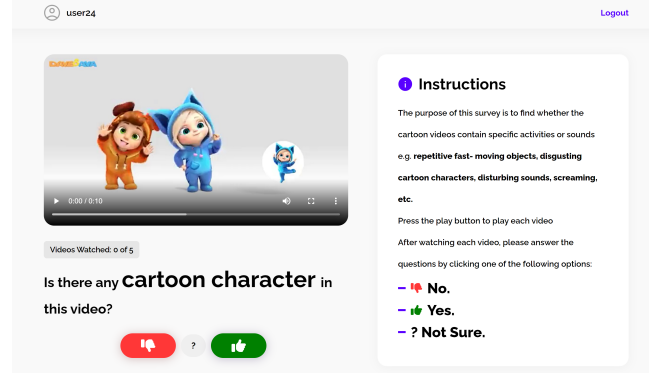


Figure 3: Layout of MAWA (MOB Annotation Web Application)

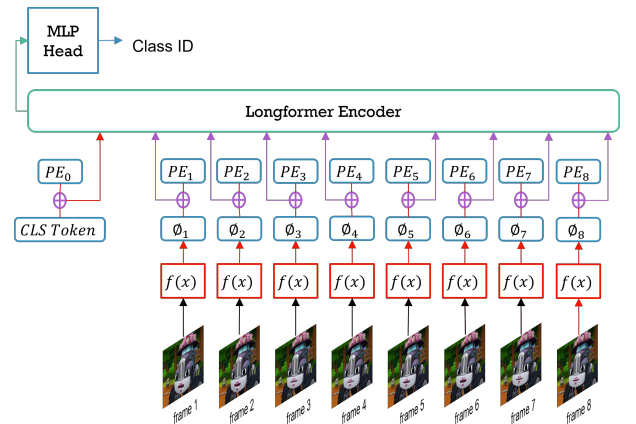


Figure 4: Architecture of Video Transformer Network (VTN) as applied to the frames of our MOB dataset.

Benchmark Evaluation

To stimulate research on video content moderation, we have also included results from a competitive set of video analysis benchmarks. This section describes the deep learning models we selected for our evaluation. We employ three different approaches: 1) first, for the baseline experiment, we consider a 3D convolutional neural network (CNN) namely I3D (Carreira and Zisserman 2017); 2) then we have a recurrent architecture i.e., Convolutional Long Short-term Memory (ConvLSTM) (Donahue et al. 2015; Ng et al. 2015); 3) and lastly we adopt a more recent transformer-based approach, the Video Transformer Network (VTN) (Neimark et al. 2021). These approaches are briefly discussed in the following subsections.

Model	Training				Testing			
	Accuracy	Precision	Recall	F1	Accuracy	Precision	Recall	F1
VTN	97.93	98.18	97.91	97.86	77.85	82.27	87.17	82.69
I3D	93.33	95.64	96.79	96.03	72.11	80.42	82.03	80.92
ConvLSTM	72.39	73.26	93.45	80.78	69.71	78.87	85.95	78.07

Table 4: Training and evaluation scores for our benchmarks. The transformer-based VTN achieves the best performance.

VTN

VTN is a deep learning model that utilizes the Longformer (Beltagy, Peters, and Cohan 2020) model which is based on the original Transformer model introduced by Vaswani et al. (2017). It extracts features from a 2D CNN and combines them with positional encoding to learn temporal relationships for video specific tasks such as classification. It is robust to long video sequence, since it is based on Longformer which essentially removes the limitation of the original Transformer and replaces the self-attention mechanism with a more scalable windowed attention along with task specific global attention.

ConvLSTM

ConvLSTMs (Donahue et al. 2015; Ng et al. 2015) combines convolution layers with LSTM layers where convolution layers extract features while LSTM layers process that information over multiple time intervals. They have been used in literature for various tasks ranging from video classifications, time series analysis, anomaly detection, and natural language processing tasks. As we are working with video classification task, they were a natural choice for a benchmark.

I3D

I3D (Carreira and Zisserman 2017) has been used in various approaches as both the main network as well as a backbone network for extracting features. In I3D, the pretrained 2D kernels of a deep network are expanded into 3D kernel which helps with learning spatio-temporal features.

Training and Evaluation

We trained each network for 100 epochs with a batch size of 8 for VTN and ConvLSTM while a batch size of 16 was used for I3D. Both VTN and I3D were trained with the SGD optimizer while ConvLSTM was trained with Adam optimizer with a learning rate of $1e^{-3}$, which was reduced every 10^{th} epoch, and standard cross entropy loss function. In addition, each model is then evaluated at every 5^{th} epoch. All the models were trained on NVidia RTX 3090 GPU.

Each clip has a length of 10 seconds where each frame has a spatial dimension of 224×224 and eight frames were used for the temporal dimension. In addition, we also used standard data augmentation techniques such as random cropping, color jitters, horizontal flipping, and normalizing.

Results and Discussion

Our empirical analysis shows that the VTN outperforms both the ConvLSTM as well as the I3D models, with an F1 score of 82.69 on the test set. These results are summarized in Table 4. It is unsurprising that the transformer-based approach outperforms the others as they can often learn better spatio-temporal features, can process longer time horizons, and are scalable too, while both ConvLSTM and I3D lack the capability of self-attention and are not as scalable as transformers.

VTN Variations We evaluated the performance of different variations of the VTN, including modifying the hidden dimension size, the number of transformer encoder layers, and other hyper-parameters.

Note that we have used competitive video classification models as our benchmarks. None of them are specialized for our task or our dataset. We hope that they will provide a fair but challenging benchmark for authors interested in doing research on content moderation for children’s cartoons.

Conclusion and Future Work

This paper addresses the problem of effective content moderation for children’s cartoon videos. Consumption of malicious content by children between ages 1-5 adversely affects their behavior and cognitive development. We enumerated a set of fine-grained malicious audiovisual features, and also presented the associated medical studies and findings supporting their exclusion. This paper proposed a comprehensive solution including a dataset, a customizable annotation tool for videos, and a benchmark suite of state-of-the-art video classification models. The dataset and annotation tool will be made available for others to stimulate research in this area.

In future work, we will focus on analyzing the disturbing audio elements present in our dataset as well leveraging meta-data, subtitles, and user comments within a multimodal classification framework. We also plan to explore few-shot learning and transfer learning techniques as mechanism to overcome data limitations.

References

Alshamrani, S.; Abusnaina, A.; Abuhamad, M.; Nyang, D.; and Mohaisen, D. 2021. Hate, obscenity, and insults: Measuring the exposure of children to inappropriate comments in YouTube. In *Companion Proceedings of the Web Conference 2021*, 508–515. Ljubljana Slovenia: ACM.

- Amadeo, R. 2022. YouTube flags horror video as “for kids,” won’t let creator change rating [updated].
- Beltagy, I.; Peters, M. E.; and Cohan, A. 2020. Longformer: The long-document transformer. *arXiv:2004.05150*.
- Carreira, J., and Zisserman, A. 2017. Quo vadis, action recognition? a new model and the Kinetics dataset. In *2017 IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, 4724–4733.
- CDC. 2019. Positive parenting tips.
- Ceci, L. 2022. YouTube: most viewed video.
- Chuttur, M. Y., and Nazurully, A. 2022. A multi-modal approach to detect inappropriate cartoon video contents using deep learning networks. *Multimedia Tools and Applications* 81(12):16881–16900.
- Commission, F. T. 2020. Complying with COPPA: Frequently asked questions.
- Deng, J.; Russakovsky, O.; Krause, J.; Bernstein, M. S.; Berg, A.; and Fei-Fei, L. 2014. Scalable multi-label annotation. In *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems*, 3099–3102. Toronto Ontario Canada: ACM.
- Donahue, J.; Hendricks, L. A.; Guadarrama, S.; Rohrbach, M.; Venugopalan, S.; Darrell, T.; and Saenko, K. 2015. Long-term recurrent convolutional networks for visual recognition and description. In *2015 IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, 2625–2634.
- Doyle, B. 2022. TikTok statistics - everything you need to know.
- Gkolemi, M.; Papadopoulos, P.; Markatos, E.; and Kourtellis, N. 2022. YouTubers Not MadeForKids: Detecting channels sharing inappropriate videos targeting children. In *14th ACM Web Science Conference 2022, WebSci ’22*, 370–381. New York, NY, USA: Association for Computing Machinery.
- Hapkiewicz, W. G. 1979. Children’s reactions to cartoon violence. *Journal of Clinical Child Psychology* 8(1):30–34.
- Ishikawa, A.; Bollis, E.; and Avila, S. 2019. Combating the Elsagate phenomenon: Deep learning architectures for disturbing cartoons. (arXiv:1904.08910). arXiv:1904.08910 [cs].
- Kaushal, R.; Saha, S.; Bajaj, P.; and Kumaraguru, P. 2016. KidsTube: Detection, characterization and analysis of child unsafe content and promoters on youtube. (arXiv:1608.05966). arXiv:1608.05966 [cs].
- Kay, W.; Carreira, J.; Simonyan, K.; and Zhang, B. 2017. The Kinetics human action video dataset. 22.
- Klatte, M.; Bergström, K.; and Lachmann, T. 2013. Does noise affect learning? a short review on noise effects on cognitive performance in children. *Frontiers in Psychology* 4.
- Le, B. M.; Tandon, R.; and Oinar, C. 2022. Samba: Identifying inappropriate videos for young children on youtube. 10.
- Lillard, A. S., and Peterson, J. 2011. The immediate impact of different types of television on young children’s executive function. *Pediatrics* 128(4):644–649.
- Liu, J.; Tandon, R.; Durairaj, U.; Guo, J.; Zahabizadeh, S.; Ilango, S.; Tang, J.; Gupta, N.; Zhou, Z.; and Mirkovic, J. 2022. Did your child get disturbed by an inappropriate advertisement on YouTube? 7.
- Neimark, D.; Bar, O.; Zohar, M.; and Asselmann, D. 2021. Video transformer network. In *2021 IEEE/CVF International Conference on Computer Vision Workshops (ICCVW)*, 3156–3165.
- Ng, J. Y.-H.; Hausknecht, M.; Vijayanarasimhan, S.; Vinyals, O.; Monga, R.; and Toderici, G. 2015. Beyond short snippets: Deep networks for video classification. In *2015 IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, 4694–4702.
- Ovide, S. 2021. The messy truth about kids’ screen time. *The New York Times*.
- Papadamou, K.; Papasavva, A.; Zannettou, S.; Blackburn, J.; Kourtellis, N.; Leontiadis, I.; Stringhini, G.; and Sirivianos, M. 2021. Disturbed YouTube for Kids: Characterizing and detecting inappropriate videos targeting young children. (arXiv:1901.07046). arXiv:1901.07046 [cs].
- Placido, D. D. 2017. YouTube’s “Elsagate” illuminates the unintended horrors of the digital age.
- Ruby, D. 2022. YouTube Statistics (2022) — updated data, facts and figures shared!
- Sinelschikova, Y. 2020. This creepy children’s cartoon has 3.5 billion views (video).
- Singh, S.; Kaushal, R.; Buduru, A. B.; and Kumaraguru, P. 2019. KidsGUARD: fine grained approach for child unsafe video representation and detection. In *Proceedings of the 34th ACM/SIGAPP Symposium on Applied Computing*, 2104–2111. Limassol Cyprus: ACM.
- Tahir, R.; Ahmed, F.; Saeed, H.; Ali, S.; Zaffar, F.; and Wilson, C. 2019. Bringing the kid back into YouTube kids: detecting inappropriate content on video streaming platforms. In *Proceedings of the 2019 IEEE/ACM International Conference on Advances in Social Networks Analysis and Mining*, 464–469. Vancouver British Columbia Canada: ACM.
- Vaswani, A.; Shazeer, N.; Parmar, N.; Uszkoreit, J.; Jones, L.; Gomez, A. N.; Kaiser, L.; and Polosukhin, I. 2017. Attention is all you need.
- Vijayakrishnan, V.; Hari Krishnan, D.; Babu, D.; and Kochi, A. V. V. 2018. Marketing strategy in advertisements using animated characters. *International Journal of Pure and Applied Mathematics* 119(12):2841–2852.
- Yousaf, K., and Nawaz, T. 2022. A deep learning-based approach for inappropriate content detection and classification of YouTube videos. *IEEE Access* 10:16283–16298.
- YouTube Help. Fair use on YouTube.
- YouTube Help. FAQ on fair use.