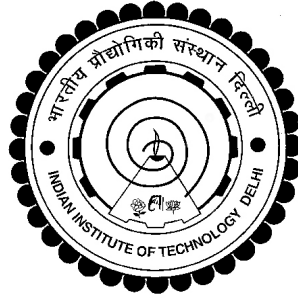


IMAGE ANNOTATION AND TAG PROPAGATION

Anurag Tripathi



DEPARTMENT OF ELECTRICAL ENGINEERING
INDIAN INSTITUTE OF TECHNOLOGY DELHI
MAY 2021

IMAGE ANNOTATION AND TAG PROPAGATION

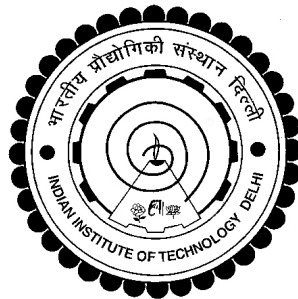
by

Anurag Tripathi

DEPARTMENT OF ELECTRICAL ENGINEERING

Submitted

*in fulfillment of the requirements of the degree of Doctor of Philosophy
to the*



INDIAN INSTITUTE OF TECHNOLOGY DELHI

MAY 2021

Certificate

This is to certify that the thesis entitled “**IMAGE ANNOTATION AND TAG PROP-AGATION**” being submitted by **Mr. ANURAG TRIPATHI** to the Department of Electrical Engineering, Indian Institute of Technology Delhi, for the award of the degree of **Doctor of Philosophy** is the record of the bonafide research work carried out by him under my supervision. In my opinion, the thesis has reached the standards fulfilling the requirements of the regulations relating to the degree.

Prof. Brejesh Lall
Department of Electrical Engineering
Indian Institute of Technology Delhi

Prof. Santanu Chaudhury
Department of Computer Science and Engineering
Indian Institute of Technology Jodhpur

Acknowledgements

I express my sincere thanks to my respected supervisor, Prof. Santanu Chaudhary and Prof. Brejesh Lall for their valuable advice and guidance. I would also like to thank my friends Dr. Siddharth Srivastava, Dr. Abhinav Gupta and Dr. Prerana Mukherjee for their support at all the times. I would like to sincerely thank everyone who guided, helped and supported me in this part of my life's journey.

I would like to mention that this dissertation would not have been possible without support of my parents, my wife, my sister and kids. I would like to express my heartiest gratitude towards my wife, who stood by me patiently through all my efforts and encouraged me to pursue Ph.D.

I dedicate this thesis to my spiritual guru, Sri Dr. Rajendra Das Devacharya, for imparting strength and patience making it possible in balancing through the ups and downs of during my life and Ph.D.

Anurag Tripathi

Abstract

In the era of connected devices and internet large amount of digital data is generated which keeps on growing day by day, with images and videos contributing a majority to it. Such humongous amount of data presents two challenges. First is it's management and the second is extracting meaningful information from this data. However both of these challenges are dependent on each other. For example, if one can interpret the meaning of the data, the organization of data becomes easier and vice-versa. Moreover, text based interaction is more convenient to communicate with the visual data hence its an obvious direction to develop similar technologies for archiving and retrieving visual data. To this end, this thesis first aims towards tagging of images which aids effective organization of data. Further, we utilize these annotations in understanding images. Specifically, we work on prediction of relationship among objects within a scene. Subsequently, based on our contributions in image annotations and understanding of relationship between objects, we propose to generate novel natural scenes in a controlled manner.

Our first work is in annotation of images, where we have two contributions, first approach is probabilistic component analysis (PLCA) based framework wherein we represent the image objects as latent components. Prior work in literature assumed that tags are mutually exclusive, however our analysis shows that the tags are highly correlated. For example correlation between car and road is high as compared to car and chair. The second approach is a joint framework which performs both image classification and annotation. Here, we used the correlation between an image class and their corresponding annotation.

Once we have the annotations, next step is to investigate the semantic information within

objects of an image. We can establish the relationship between the objects by analyzing the interaction between them. This relationship between the objects or labels aids towards deeper understanding of the scene. Therefore, the third contribution of this thesis is to detect relationship using visual information i.e. classes of objects within an image. For example if subject and object in an image are *person* and *bike*, and interaction between them can be a predicate *riding on* from the visual cues. Such visual relationship can be formulated as triplets in the form of $(subject, predicate, object)$. The relationship in a given image can be a verb (e.g. ride), or preposition (e.g. by), spatial phrase (e.g. in the front of), or comparative phrase (e.g. taller than). These relationships can be used for application such as image captioning etc.

Now as an inverse problem our fourth contribution is a framework to generate images containing scenes corresponding to specified objects aligned with the textual annotations/classes. The proposed technique generates scene which has a high correlation with the provided set of input objects while also maintaining the natural placement of objects within the scene. This is in contrast with earlier works where the objective is to generate a natural scene from a noise vector or conditioning the network over a variable. However, such methods have limitations in their ability to control the objects within the generated images. On the contrary, we show that by training a Generative Adversarial Network with raw image pixels as input, we can generate scenes which constitute the objects as well as generate the surrounding environment suitable for the combination of the input objects.

We have evaluated our proposed framework with the relevant dataset and metrics used in similar settings. On image annotation we have achieved state-of-the art results as compared to other method. On image generation task we have achieved state-of-the art results using the inception score as metrics.

सार

कनेक्टेड डिवाइस और इंटरनेट के युग में बड़ी मात्रा में डिजिटल डेटा उत्पन्न होता है जो दिन-प्रतिदिन बढ़ता रहता है, जिसमें छवियों और वीडियो का बड़ा योगदान होता है।

इतनी भारी मात्रा में डेटा दो चुनौतियां पेश करता है। पहला है इसका प्रबंधन और दूसरा है इस डेटा से सार्थक जानकारी निकालना। हालाँकि ये दोनों चुनौतियाँ एक दूसरे पर निर्भर हैं। उदाहरण के लिए, यदि कोई डेटा के अर्थ की व्याख्या कर सकता है, तो डेटा का संगठन आसान हो जाता है और इसके विपरीत। इसके अलावा, दृश्य डेटा के साथ संवाद करने के लिए पाठ आधारित बातचीत अधिक सुविधाजनक है, इसलिए दृश्य डेटा को संग्रहीत करने और पुनर्प्राप्त करने के लिए समान तकनीकों को विकसित करने की यह एक स्पष्ट दिशा है। इसके लिए, इस थीसिस का उद्देश्य सबसे पहले छवियों को टैग करना है जो डेटा के प्रभावी संगठन में सहायता करता है। इसके अलावा, हम छवियों को समझने में इन एनोटेशन का उपयोग करते हैं। विशेष रूप से, हम एक दृश्य के भीतर वस्तुओं के बीच संबंध की भविष्यवाणी पर काम करते हैं।

इसके बाद, छवि एनोटेशन और वस्तुओं के बीच संबंधों की समझ में हमारे योगदान के आधार पर, हम उपन्यास प्राकृतिक दृश्यों को नियंत्रित तरीके से उत्पन्न करने का प्रस्ताव करते हैं।

हमारा पहला काम छवियों के एनोटेशन में है, जहां हमारे दो योगदान हैं, पहला दृष्टिकोण संभाव्य घटक विश्लेषण (पीएलसीए) आधारित ढांचा है जिसमें हम छवि वस्तुओं को गुप्त घटकों के रूप में प्रस्तुत करते हैं। साहित्य में पहले के काम ने माना कि टैग परस्पर अनन्य हैं, हालांकि हमारे विश्लेषण से पता चलता है कि टैग अत्यधिक सहसंबद्ध हैं। उदाहरण के लिए कार और सड़क के बीच सहसंबंध कार और कुर्सी की तुलना में अधिक है।

दूसरा दृष्टिकोण एक संयुक्त ढांचा है जो छवि वर्गीकरण और एनोटेशन दोनों करता है।

यहां, हमने एक छवि वर्ग और उनके संबंधित एनोटेशन के बीच सहसंबंध का उपयोग किया।

एक बार हमारे पास एनोटेशन हो जाने के बाद, अगला कदम एक छवि की वस्तुओं के भीतर अर्थ संबंधी जानकारी की जांच करना है। हम वस्तुओं के बीच की बातचीत का विश्लेषण करके उनके बीच संबंध स्थापित कर सकते हैं। वस्तुओं या लेबल के बीच यह संबंध दृश्य की गहरी समझ की दिशा में सहायता करता है। इसलिए, इस थीसिस का तीसरा योगदान दृश्य जानकारी यानी छवि के भीतर वस्तुओं के वर्गों का उपयोग करके संबंधों का पता लगाना है। उदाहरण के लिए यदि किसी छवि में विषय और वस्तु {व्यक्ति} और {बाइक} हैं, और उनके बीच की बातचीत दृश्य संकेतों से एक विधेय {सवारी} हो सकती है। इस तरह के दृश्य संबंध को {विषय, विधेय, वस्तु} के रूप में त्रिगुणों के रूप में तैयार किया जा सकता है।

किसी दी गई छवि में संबंध एक क्रिया (जैसे सवारी), या पूर्वसर्ग (जैसे द्वारा), स्थानिक वाक्यांश (जैसे सामने में), या तुलनात्मक वाक्यांश (जैसे इससे लंबा) हो सकता है। इन रिश्तों का उपयोग इमेज कैप्शनिंग आदि जैसे एप्लिकेशन के लिए किया जा सकता है।

अब एक उलटी समस्या के रूप में हमारा चौथा योगदान पाठ्य एनोटेशन/वर्गों के साथ संरेखित निर्दिष्ट वस्तुओं के अनुरूप दृश्यों वाले चित्र उत्पन्न करने के लिए एक ढांचा है। प्रस्तावित तकनीक दृश्य उत्पन्न करती है जिसमें दृश्य के भीतर वस्तुओं के प्राकृतिक स्थान को बनाए रखते हुए इनपुट वस्तुओं के प्रदान किए गए सेट के साथ उच्च सहसंबंध होता है। यह पहले के कार्यों के विपरीत है जहां उद्देश्य एक शोर वेक्टर या एक चर पर नेटवर्क कंडीशनिंग से प्राकृतिक दृश्य उत्पन्न करना है। हालांकि, ऐसी विधियों में उत्पन्न छवियों के भीतर वस्तुओं को नियंत्रित करने की उनकी क्षमता में सीमाएं हैं। इसके विपरीत, हम दिखाते हैं कि इनपुट के रूप में कच्ची छवि पिक्सेल के साथ एक जनरेटिव एडवरसैरियल नेटवर्क को प्रशिक्षित करके, हम ऐसे दृश्य उत्पन्न कर सकते हैं जो वस्तुओं का निर्माण करते हैं और साथ ही इनपुट ऑब्जेक्ट्स के संयोजन के लिए उपयुक्त परिवेश उत्पन्न करते हैं।

हमने समान सेटिंग्स में उपयोग किए गए प्रासंगिक डेटासेट और मीट्रिक के साथ हमारे प्रस्तावित ढांचे का मूल्यांकन किया है। छवि एनोटेशन पर हमने अन्य पद्धति की तुलना में अत्याधुनिक परिणाम प्राप्त किए हैं।

इमेज जनरेशन टास्क पर हमने इंसेप्शन स्कोर को मेट्रिक्स के रूप में इस्तेमाल करते हुए अत्याधुनिक परिणाम हासिल किए हैं।

Contents

Certificate	i
Acknowledgements	ii
Abstract	iii
List of Figures	xi
List of Tables	xiii
1 Introduction	1
1.1 Problems of Interest	2
1.2 Basic Terminology	3
1.2.1 Image Annotation	4
1.2.2 Scene Graph	4
1.2.3 Visual Relations	6
1.2.4 Generative Adversarial Network (GAN)	6
1.3 Contributions	9
1.4 Thesis Organization	10
2 Literature Survey	13
2.1 Image Annotation	14

2.2	Visual Relationship Detection	17
2.3	Image Generation	20
2.4	Motivation	22
3	EM based Image Annotation	25
3.1	Introduction	25
3.2	Tag Prediction using Latent Components	27
3.2.1	Preliminaries	28
3.2.2	Extracting Visual Words using PLCA	28
3.2.3	Visual and Textual Co-occurrence	30
3.2.4	Transmedia Association Model	31
3.2.5	Tag Prediction	33
3.3	Experimental Results	34
3.4	Conclusion	36
4	Image annotation and classification using Iterative learning	37
4.1	Introduction	37
4.2	Methodology	39
4.2.1	Loss Function	41
4.2.2	Tag Prediction	43
4.2.3	Classification	43
4.2.4	Training and Testing	44
4.3	Experiments	45
4.3.1	Experimental Setup	45
4.3.2	Comparison with different models	46
4.3.3	Results	47
4.4	Conclusion	53

5	Scene graph based visual relationship prediction	54
5.1	Introduction	54
5.2	Methodology	57
5.2.1	Preliminaries	57
5.2.2	Object Detection	58
5.2.3	Spatially aware word embeddings	58
5.2.4	Network Architecture	59
5.2.5	Testing	62
5.3	Experiments	63
5.3.1	Dataset	63
5.3.2	Network Settings	63
5.3.3	Metrics	63
5.3.4	Comparison with state-of-the-art	64
5.3.5	Ablation Studies	65
5.4	Conclusion	67
6	Generating images using Image Annotations	68
6.1	Introduction	68
6.2	Preliminaries of GAN	70
6.3	Methodology	71
6.3.1	Problem Formulation	72
6.3.2	Guided Compositional GAN	74
6.4	Experiments	76
6.4.1	Experimental Setup	76
6.4.2	Dataset	77
6.4.3	Parameteric Settings	77
6.4.4	Metric	78
6.4.5	Quantitative Results	78

6.4.6 Qualitative Results	79
6.5 Conclusions	80
7 Conclusions and Future Work	83
Bibliography	86
List of Publications	97
Biodata of Author	98

List of Figures

1.1	Illustration of Automatic Image Annotation.	2
1.2	Overview of the sub problems related to the image annotation that are addressed	4
1.3	Automatic Image Annotation	5
1.4	Scene Graph of an Image	5
1.5	Visual Relations	6
1.6	Generative Adversarial Network	7
1.7	Network Architecture of Deep Convolutional Generative Adversarial Network	8
2.1	Left side figure represents holistic view and right side figure represents compositional view	18
3.1	Block diagram of proposed transmedia annotation model.	29
3.2	Computing image to label correspondence using textual co-occurrence.	31
3.3	Illustrating tag prediction using proposed transmedia annotation model.	33
4.1	Illustration of relationship among classes, images and annotations	38
4.2	Workflow of iterative classification and annotation. Upper part of figure (Training) shows the pipeline for training where the arrows depict the iterative optimization. Bottom part figure (Testing) shows the testing pipeline which uses the parameters learned during training.	40

4.3	Confusion matrix (top) and variation in classification accuracy w.r.t number of iterations (bottom) on LabelMe dataset.	48
4.4	Annotation Error (top) and Classification Error (bottom) for the LabelMe at 40 th epoch.	49
4.5	Confusion matrix (top) and variation in classification accuracy w.r.t number of iterations (bottom) on UIUC dataset.	50
4.6	Annotation Error (top) and Classification Error (bottom) for the UIUC at 40 th epoch.	51
5.1	Process of Visual Relationship Detection	55
5.2	The Proposed Framework for Scene Graph based Visual Relation Detection	57
5.3	Spatially aware scene graph embeddings as t-SNE plot (top). The bottom image shows a zoomed in subsection of the t-SNE plot of the spatially aware word embedding	62
5.4	Relationship distribution in dataset	64
6.1	Representative diagram showing the workflow of the proposed technique. The input is an object which is passed to a neural network generating a scene with the context generated for the objects and the relationship among them.	70
6.2	Architecture of Guided Compositional Generative Adversarial Network (gcGAN)	72
6.3	Object to Image Generation: First row shows the input images to gcGAN (DCGAN) while the second row shows the corresponding generated images.	76

6.4	Object to Image Generation: First row shows the input images to gcGAN (MSG-GAN) while the second row shows the corresponding generated images. To test the applicability to data collected from other sources, the input objects are obtained from randomly collected images from internet.	77
6.5	Evolution of the generated images at different epochs for gcGAN (DC-GAN).	81
6.6	Object detection with YOLO for images generated with gcGAN (DC-GAN). The YOLO detections and classes are shown in rectangular boxes.	81
7.1	Query: Children playing with chess under tree	85
7.2	Query: Children playing with chess	85
7.3	Query: Children playing under tree	86
7.4	Style transfer from <i>horse</i> $\rightarrow$ <i>zebra</i> by [1], Style gets transfer over the person	86

List of Tables

2.1	Summarization of state-of-the-art methods for visual relationship detection on Visual Genome Dataset	20
3.1	Performance of various tag prediction approaches	35
4.1	Performance comparison of tag prediction LabelMe	52
4.2	Performance of various classification approaches on LabelMe & UIUC (%)	52
4.3	Performance comparison of tag prediction on UIUC	52
5.1	Comparison with state-of-the-art for predicate and relationship detection on Visual Genome dataset. Ours (S <i>or</i> O) and Ours (S <i>and</i> O represent <i>or</i> , <i>and</i> contextual alignment of subjects and objects.	66
5.2	Results on VG with Scene Graph Embedding generated with subject and object individually aligned on various pooling methods and combination of features for predicate detection. Learning Rate = 10^{-5} , DropOut = 0.5. P is Predicate, SG is scene graph, O is Object and S is Subject features. N is set to 50.	66
5.3	Analysis of variation in hyper parameters by individually aligning the subject and object. DO represents dropout, LR represents learning rate, ZS_R represents recall in zero-shot setting.	67

5.4	Analysis of varying number of nearest neighbors (NN) at test time for predicate detection. The dropout rate is set to 0.6 and learning rate to 10^{-6} with average pooling.	67
6.1	Performance evaluation on MSCOCO. Higher is better.	79
6.2	Inception Score on MSCOCO for validation set with gcGAN(DCGAN)	79