

Connecting Vision and Language



Raymond Ptucha,
Rochester Institute of Technology



~ GCCIS PhD Colloquium ~
~ Rochester Joint Chapter of the IEEE
Computer and Computational
Intelligence Societies ~



Sept 08, 2017
GOL-1610

Modern Technology is Awesome!

- Pack 7 billion transistors on a single chip;
- Routinely fly commercial passenger aircraft at 70% the speed of sound;
- Construct buildings $\frac{1}{2}$ mile tall;
- 5 billion videos are watched on YouTube every day;
- Trade 2 billion shares per day;
- Lab-grown body parts are being used for life-saving transplants.



Burj Khalifa in
Dubai- 2,717 ft tall

Evolution of Change

- 30 years ago computers were for data entry and simple processing.
- 20 years ago cell phones were for voice and very limited texting.
- 10 years ago smart phones were primitive, buggy, and expensive.
- 5 years ago deep learning was slow and ineffective.



Ptucha '17

5



Ph.D. Students Need to Change as Well!

- | | |
|---|-----|
| • Year 1: So hungry for knowledge, wanted to work on everything. | 1.5 |
| • Year 2: Lots of applications oriented research- Wanted to build things to change the world! | 2.5 |
| • Year 3: Going after top-tier venues. Reviewers are brutal! Adding more algorithms and math. | 3 |
| • Year 4: Reviewers don't care about applications, they want math/science contributions. | 4 |
| • Year 5: Build on foundation. Focus on publishing. | 9 |

Ptucha '17

6



Ph.D., Doctor of Philosophy, Golisano College
of Computing and Information Sciences

- Unless you are hungry for knowledge, you should not be getting a Ph.D.
- Listen to your advisor on both research topics and how to write papers.
- Find a topic you have passion for- stay focused!



Your Ph.D. is a 24/7 job...

- If you are not feeling challenged/stressed, you are not working hard enough.
- Don't work in a vacuum, actively seek out advice from your peers and experts in the field.
- Be proud of everything you do.

Ptucha '17

7

A Little Background...



Ptucha '17

8

A Little Background...



Impact of Machine Learning

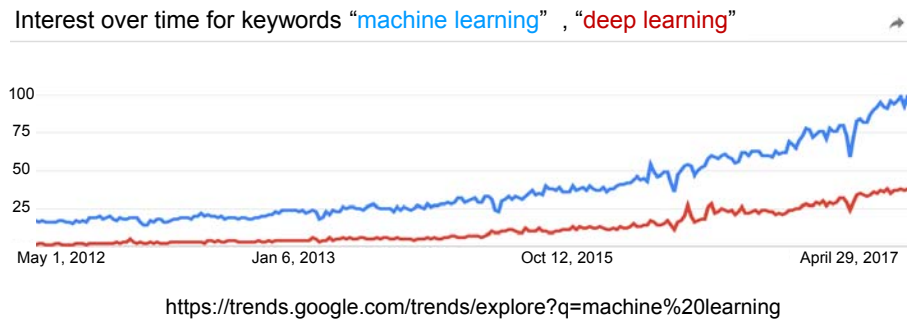
- Machine learning is giving computers the ability to analyze, generalize, think/reason/behave like humans.
- Machine learning is transforming medical research, financial markets, international security, and generally making humans more efficient and improving quality of life.
- Inspired by the mammalian brain, deep learning is machine learning on steroids- bigger, faster, better!



Ptucha '17

10

Machine Learning is a Hot Topic!



- Machine learning, cs229 is most popular course at Stanford.
- Deep learning class, cs231 went from 150 in 2015 to 350 in 2016 to 750 in 2017!

Ptucha '17

11

AI Trends

- Supervised learning has generated billions of marketable products- targeted advertising, click through rates on web pages, and driver assistance are just a few examples.
- Despite all the AI hype today, humans are still much more capable than machines at general tasks.

..however, for targeted tasks, things are different!

- Rule of thumb: Anything that a human can do with $\leq$ one second of thought can be probably now or soon be automated with AI.

Ptucha '17

16



Deep Learning

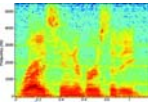
Hottest topic in pattern recognition

Vision



"Calista
Flockhart"

Speech



"Greetings
ladies and
gentlemen"

Network security



Medical diagnosis



Financial markets



Ptucha '17

17

Vision Tasks

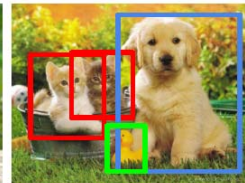
Classification Classification + Localization Object Detection Instance Segmentation



CAT



CAT



CAT, DOG, DUCK



CAT, DOG, DUCK

Ptucha '17

18

Speech Recognition

- Conversion of speech to spectrograms, then from spectrograms to words.

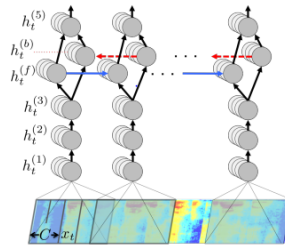


Figure 1: Structure of our RNN model and notation.

Hannun, Awni, et al. "Deep speech: Scaling up end-to-end speech recognition." *arXiv preprint arXiv:1412.5567* (2014).

Ptucha '17

19

OCR/ICR Tasks

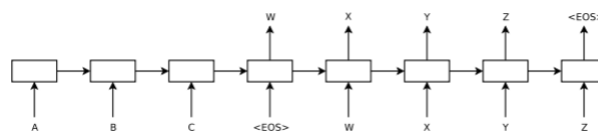


Ptucha '17

20

Language Translation

- In 2014, Sutskever, Vinyals, and Le (from Google) showed that a simple encoder-decoder framework was just as good as sophisticated Statistical Machine Translation (SMT) systems, and almost as good as SMT systems paired with nnets.

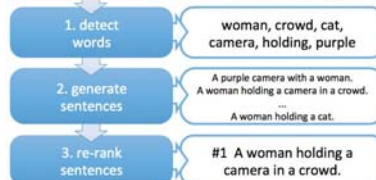
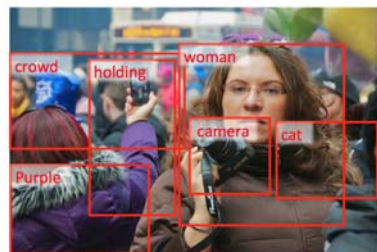


Sutskever et al., NIPS '14

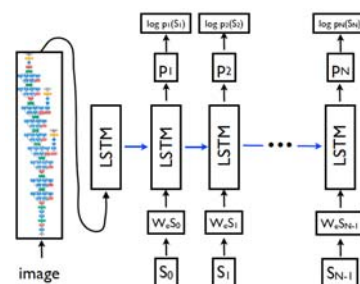
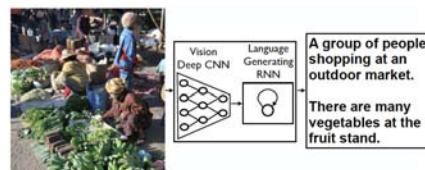
Ptucha '17

21

Image/Video Captioning



[Fang et al. CVPR15]



[Vinyals et al. CVPR15]

Ptucha '17

22

Two Most Important Deep Learning Fields

- Convolutional Neural Networks (CNN)
 - Examine high dimensional input, learn features and classifier simultaneously
- Recurrent Neural Networks (RNN)
 - Learn temporal signals, remember both short and long sequences

Ptucha '17

23

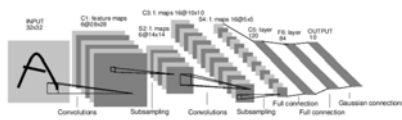
Two Most Important Deep Learning Fields

- Convolutional Neural Networks (CNN)
 - Examine high dimensional input, learn features and classifier simultaneously
- Recurrent Neural Networks (RNN)
 - Learn temporal signals, remember both short and long sequences

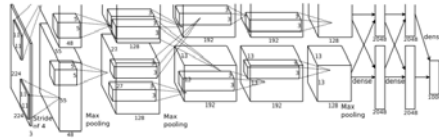
Ptucha '17

24

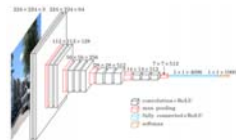
Many Flavors of CNNs...



LeNet-5, LeCun 1989



AlexNet, Krizhevsky 2012



VGGNet, Simonyan 2014



GoogLeNet (Inception), Szegedy 2014



ResNet, He 2015



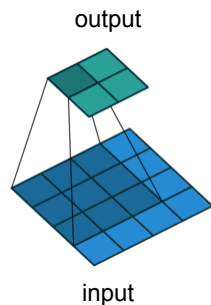
DenseNet, Huang 2017

Ptucha '17

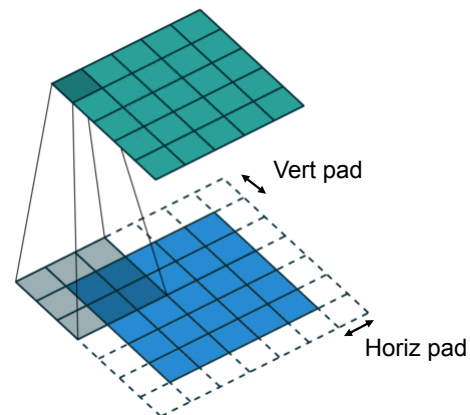
25

Image Convolution

By padding $(\text{filterWidth}-1)/2$, output image size matches input image size



3x3 filter sliding over input image

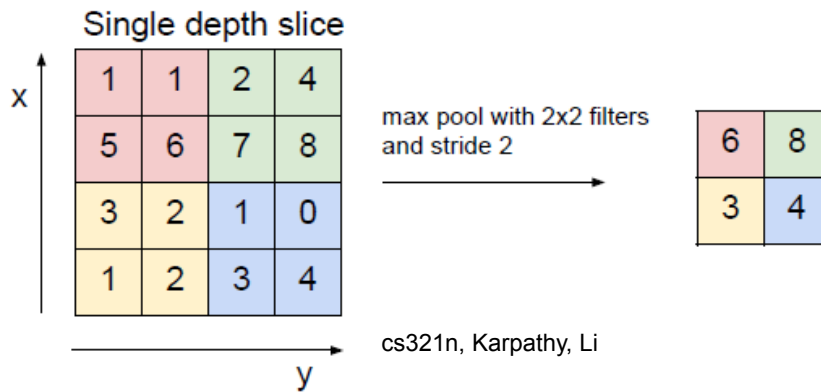


https://github.com/vdumoulin/conv_arithmetic

Ptucha '17

26

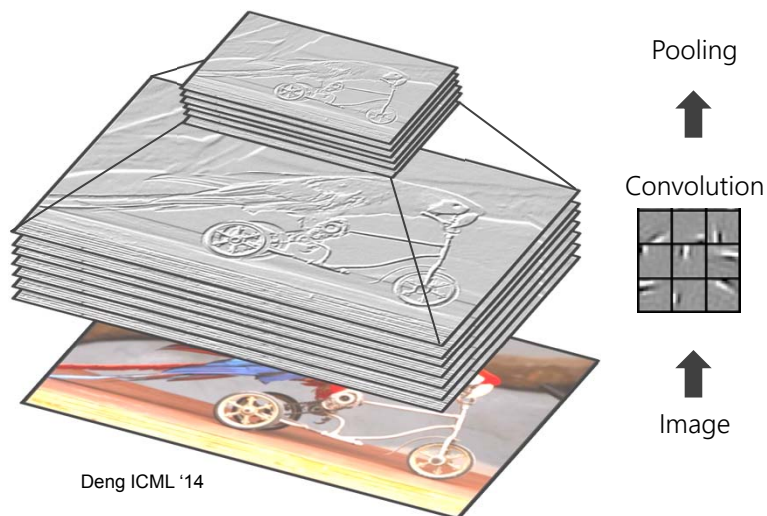
Max Pooling- Reducing the Size of an Image



Ptucha '17

31

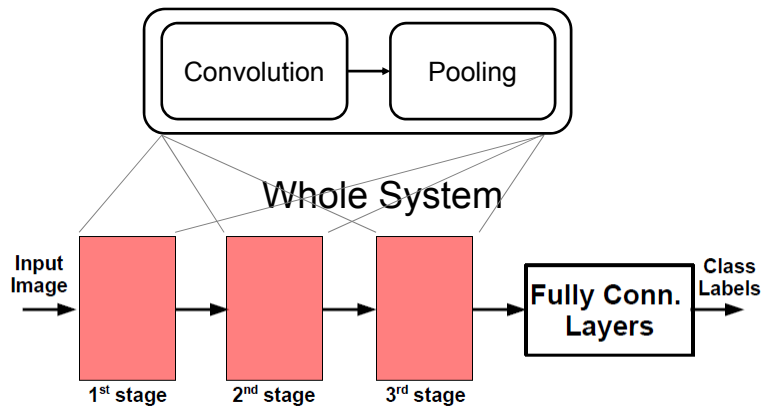
Convolution Neural Network (CNN) Building Block



Ptucha '17

32

Putting it All Together



Ptucha '17

34

Learning Filters

32 Learned filters, each 5×5

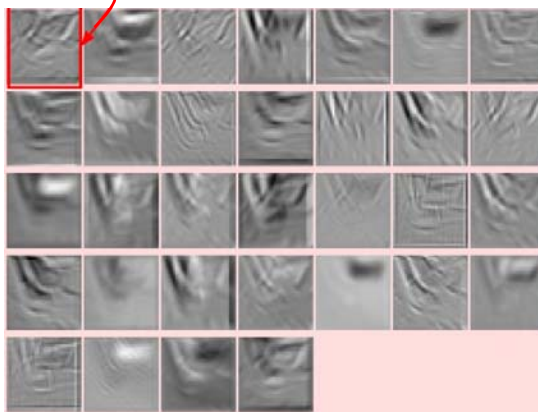


32 Filtered images (activation maps), each is 28×28

Input image
28×28



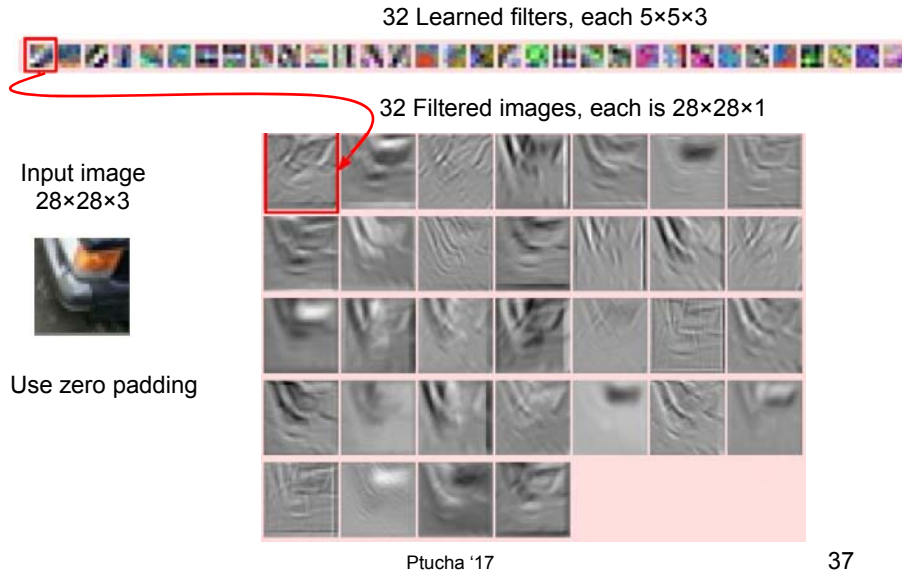
Use zero padding



Ptucha '17

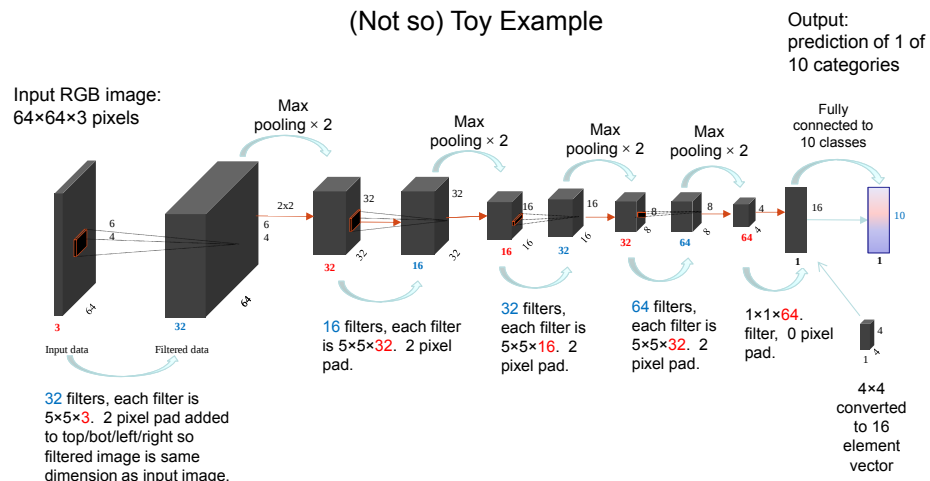
35

Learning Filters



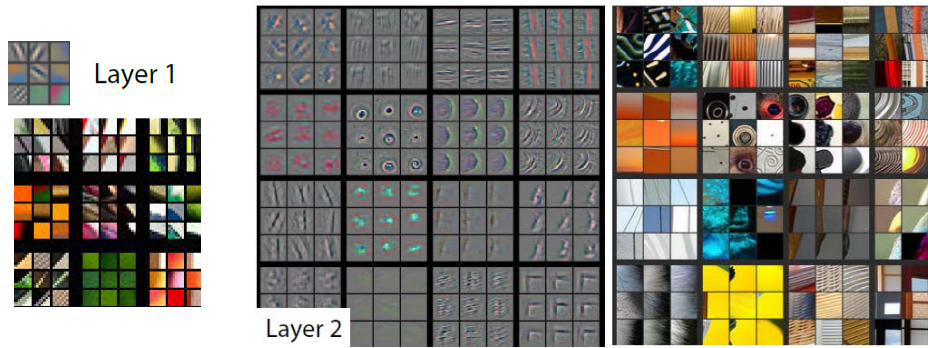
CNN Architecture

(Not so) Toy Example



CNN Visualization

...amazingly, these filters form an abstract hierarchy...

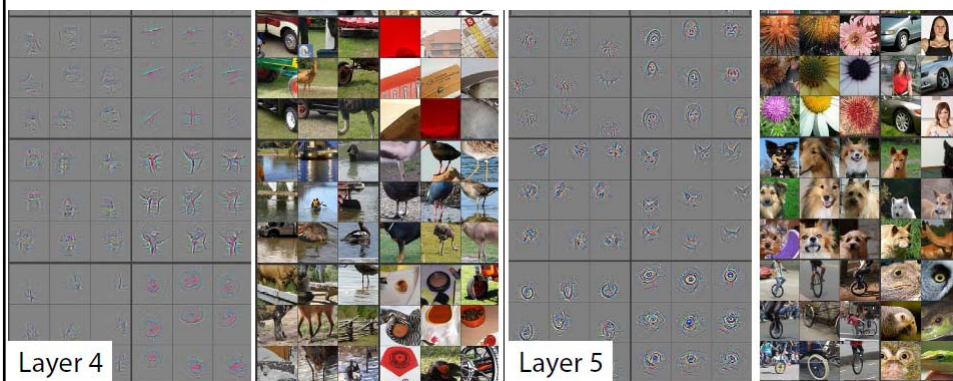


Zeiler, Fergus, 2014

Ptucha '17

40

CNN Visualization



Zeiler, Fergus, 2014

Ptucha '17

41

CNN Results

- Handwritten characters
 - MNIST: 0.17% error, Ciresan et al '11
 - Arabic & Chinese: Ciresan '12
- CIFAR-10 (60K images of 10 classes)
 - 9.3% error, Wan et al. '13
- Traffic Sign Recognition
 - 0.56% error vs 1.16% for humans, Ciresan '11



CIFAR, Krizhevsky '09



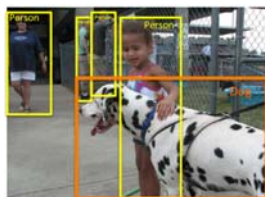
Ciresan '11

Ptucha '17

43

ImageNet

- Amazon Turk did bulk of labeling
- 14M labeled images
- 20K classes



Russakovsky et al., 2015

IMAGENET Large Scale Visual Recognition Challenge (ILSVRC)

- 1.2M images, 1000 categories
- Image classification, object localization, video detection

Ptucha '17

44

A large, dense grid of 100 small images showing various tools, equipment, and materials used in construction and manufacturing, arranged in a 10x10 pattern. The images include: power tools like drills, saws, and grinders; hand tools like hammers, wrenches, and sockets; construction materials like rebar, pipes, and concrete; and various pieces of machinery and equipment. The images are arranged in a grid, with some images showing tools in use or in storage, and others showing materials or components. The overall theme is construction and manufacturing.

45

Top-5 error on ImageNet



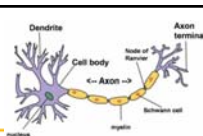
16

Two Most Important Deep Learning Fields

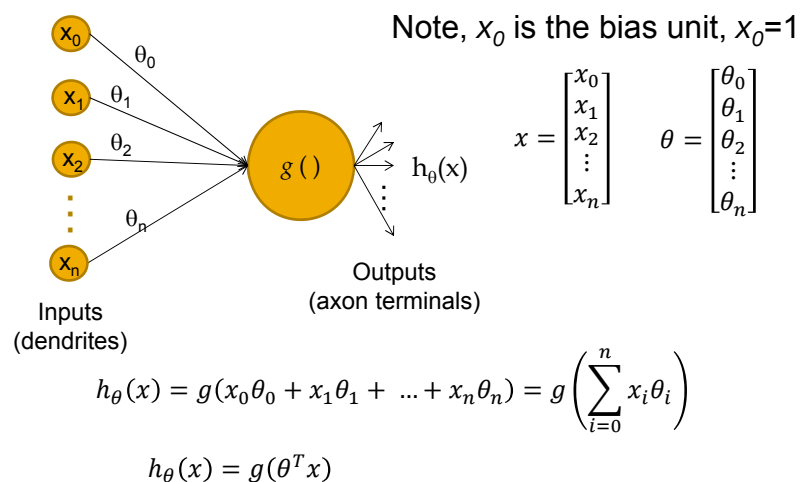
- Convolutional Neural Networks (CNN)
 - Examine high dimensional input, learn features and classifier simultaneously
- Recurrent Neural Networks (RNN)
 - Learn temporal signals, remember both short and long sequences

Ptucha '17

47



Artificial Neuron

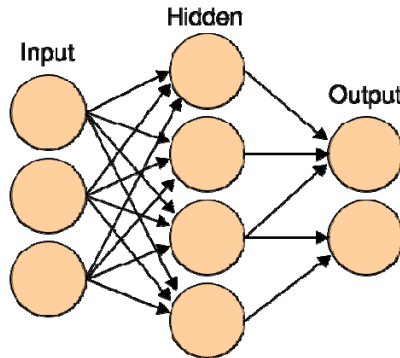


Ptucha '17

48

Artificial Neural Networks

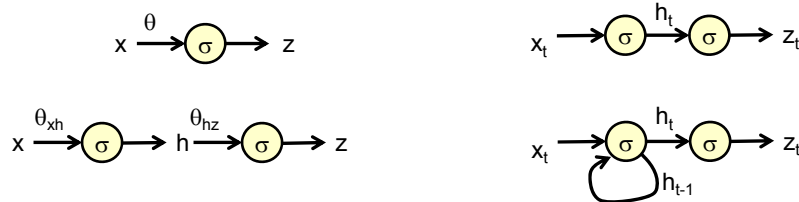
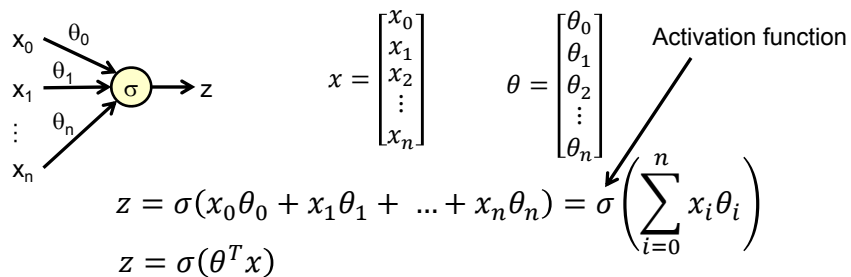
- Artificial Neural Network (ANN) – A network of interconnected nodes that “mimic” the properties of a biological network of neurons.



Ptucha '17

49

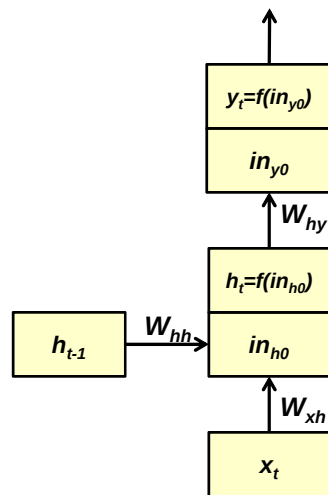
Adding Recurrence



Ptucha '17

50

Recurrent Networks



$$in_{h0} = (W_{xh}x_t + W_{hh}h_{t-1})$$

$$h_t = f(in_{h0})$$

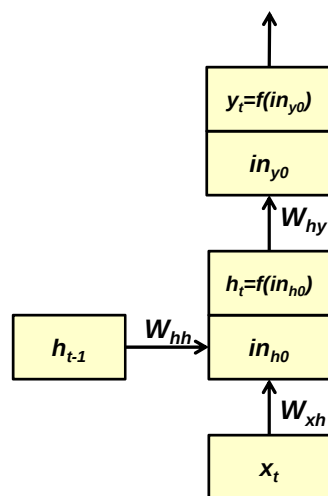
Where:

- f is some activation function
- x_t and h_t are current input and current output values
- W_{xh} is the weight matrix for input, hidden and output stages respectively
- in_{h0} is the input to activation function in hidden and output layers

Ptucha '17

51

Recurrent Networks



$$in_{h0} = (W_{xh}x_t + W_{hh}h_{t-1})$$

$$h_t = f(in_{h0})$$

$$in_{y0} = W_{hy}h_t$$

$$y_t = f(in_{y0})$$

Where:

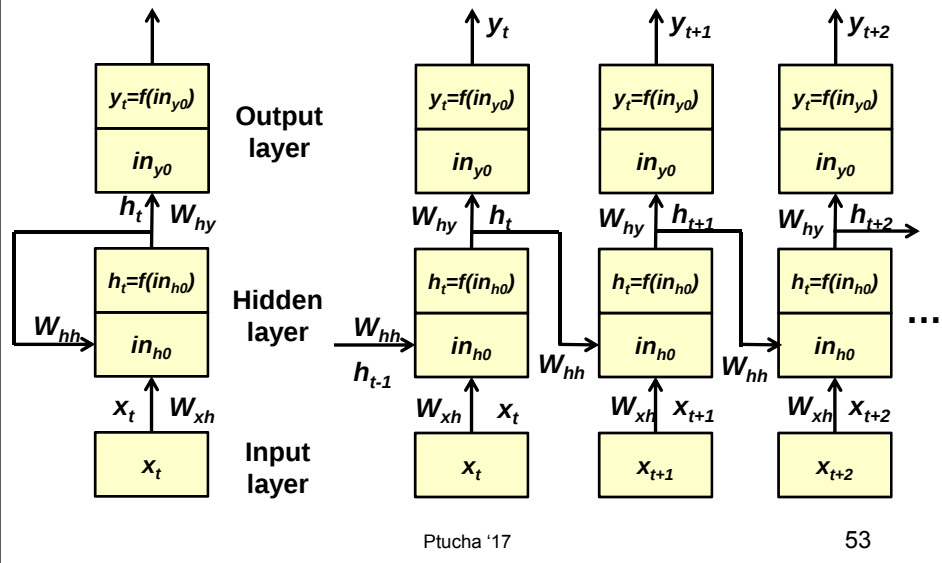
- f is some activation function
- x_t, h_t, h_{t-1} and y_t are current input, hidden, previous hidden and current output values
- W_{xh}, W_{hh} and W_{hy} are the weight matrices for input, hidden and output stages respectively
- in_{h0} and in_{y0} are the inputs to activation function in hidden and output layers

Ptucha '17

52

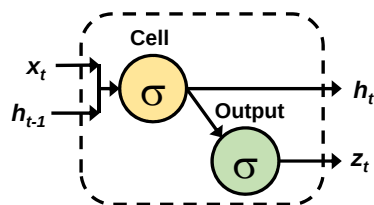
Recurrent Networks

Both figures represent the same architecture



Recurrent Networks

Recurrent Neural Network "neuron"



$P(\text{next event} \mid \text{previous events})$

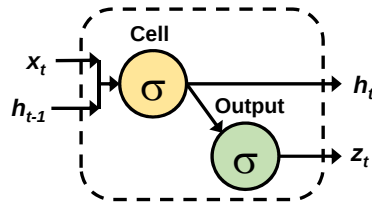
- Unfortunately, these vanilla RNNs don't always work.
- Can't store info over long periods of time.
- Suffer from vanishing and/or exploding gradients.

Ptucha '17

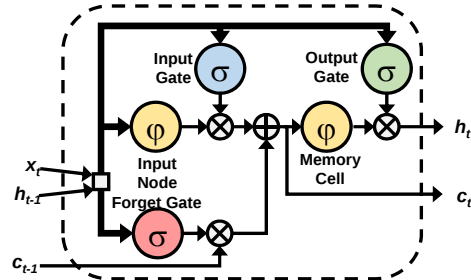
55

Recurrent Networks

Recurrent Neural Network "neuron"



Long Short Term Memory "neuron"



Donahue et al., 2015

- LSTM's allow read/write/reset functions to neurons.
- Remember past to predict the future- (over long time periods).
- Can have many hidden neurons per layer and many layers.

Ptucha '17

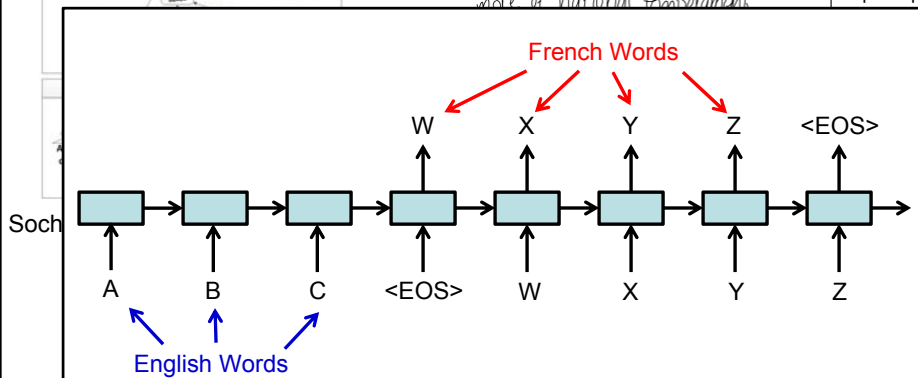
56

Recurrent Applications



more of national temperament
more of national temperament
more of national temperament

Top is input,



Sutskever et al., 2014

Ptucha '17

57

Learning Shakespeare

- LSTMs can learn structure and style in the data
- Karpathy downloaded all the works of Shakespeare and concatenated them into a single (4.4MB) file.
- Train a 3-layer LSTM with 512 hidden nodes on each layer.
- After we train the network for a few hours Karpathy obtained samples such as:

Ptucha '17

58

PANDARUS:

Alas, I think he shall be come approached and the day
When little strain would be attain'd into being never fed,
And who is but a chain and subjects of his death,
I should not sleep.

Second Senator:

They are away this miseries, produced upon my soul,
Breaking and strongly should be buried, when I perish
The earth and thoughts of many states.

DUKE VINCENTIO:

Well, your wit is in the care of side and that.

Second Lord:

They would be ruled after this chamber, and
my fair nues begun out of the fact, to be conveyed,
Whose noble souls I'll have the heart of the wars.

Clown:

Come, sir, I will make did behold your worship.

<http://karpathy.github.io/2015/05/21/rnn-effectiveness/>

Ptucha '17

59

Two Most Important Deep Learning Fields

- Convolutional Neural Networks (CNN)
 - Examine high dimensional input, learn features and classifier simultaneously
- Recurrent Neural Networks (RNN)
 - Learn temporal signals, remember both short and long sequences

Ptucha '17

62

Two Most Important Deep Learning Fields

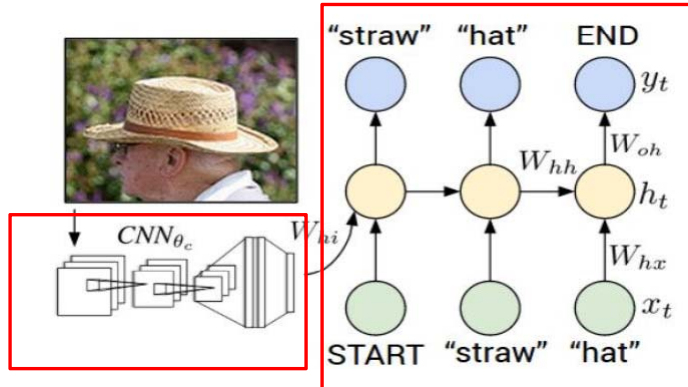
- Convolutional Neural Networks (CNN)
 - Examine high dimensional input, learn features and classifier simultaneously
- Recurrent Neural Networks (RNN)
 - Learn temporal signals, remember both short and long sequences
- Combining Visual CNNs with Language RNNs for Image Captioning

Ptucha '17

63

Image Captioning

RNN takes in a latent representation of an image, and generates a sequence.

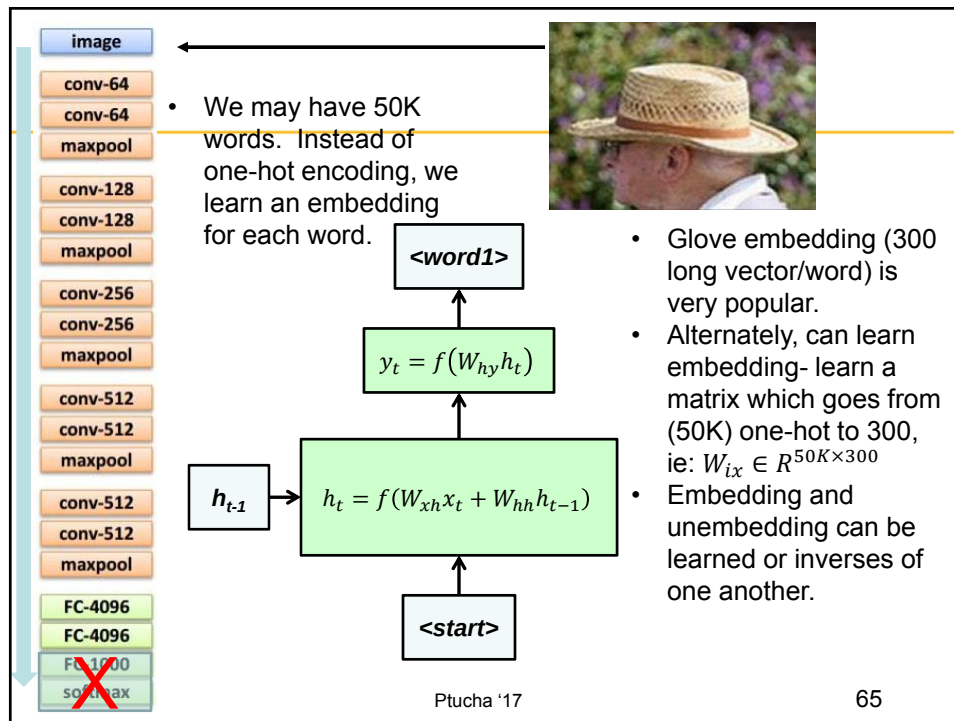


CNN helps represent an image as a numeric value. (image2vec)

Karpathy & Li, CVPR'15

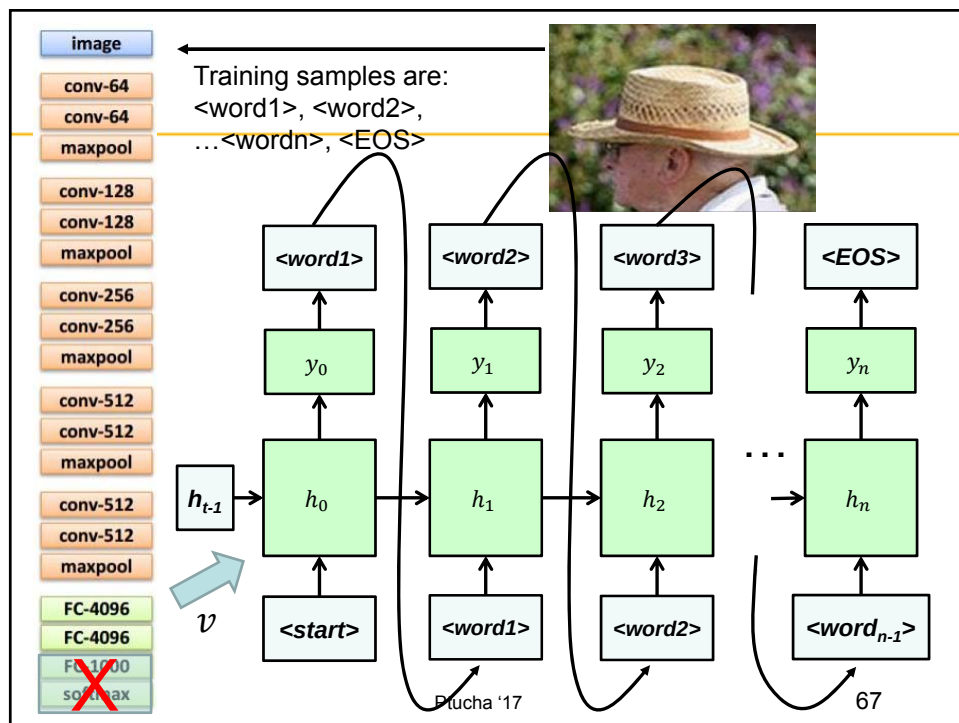
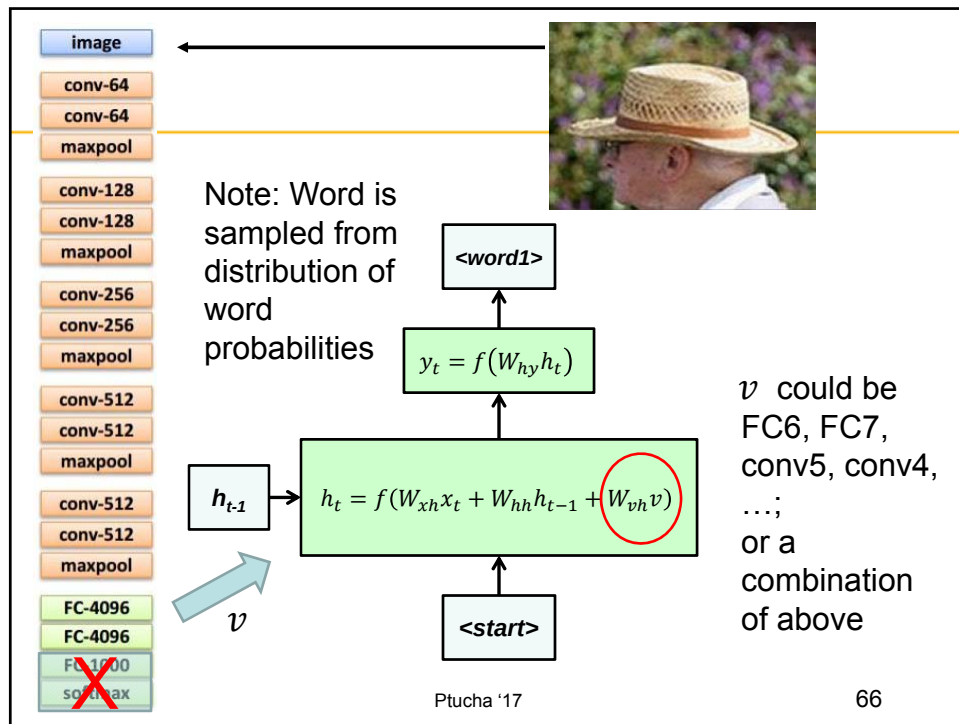
Ptucha '17

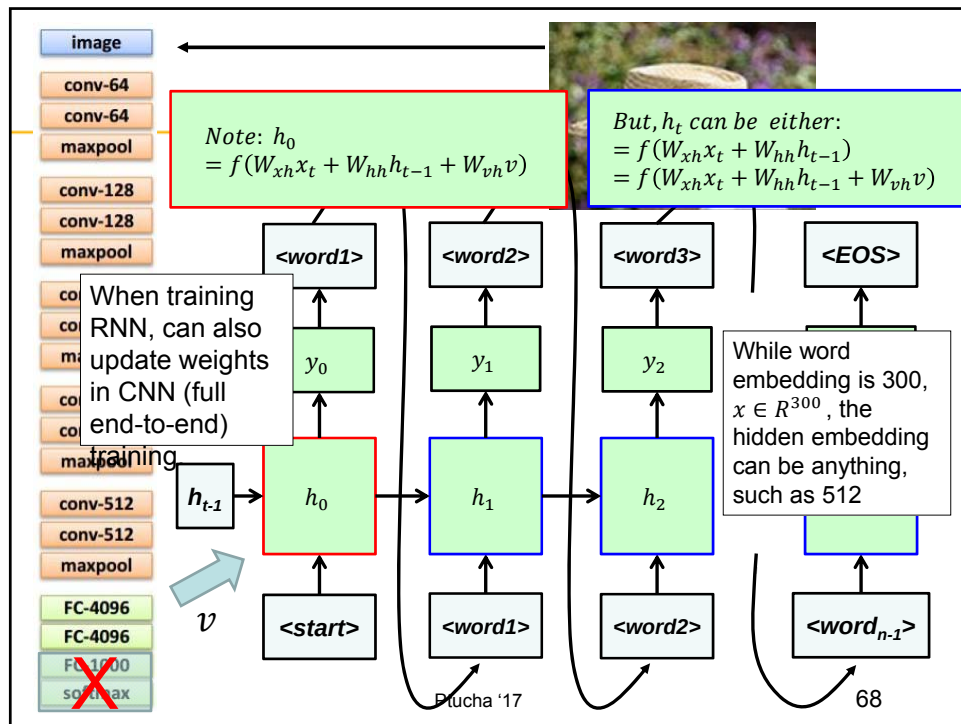
64



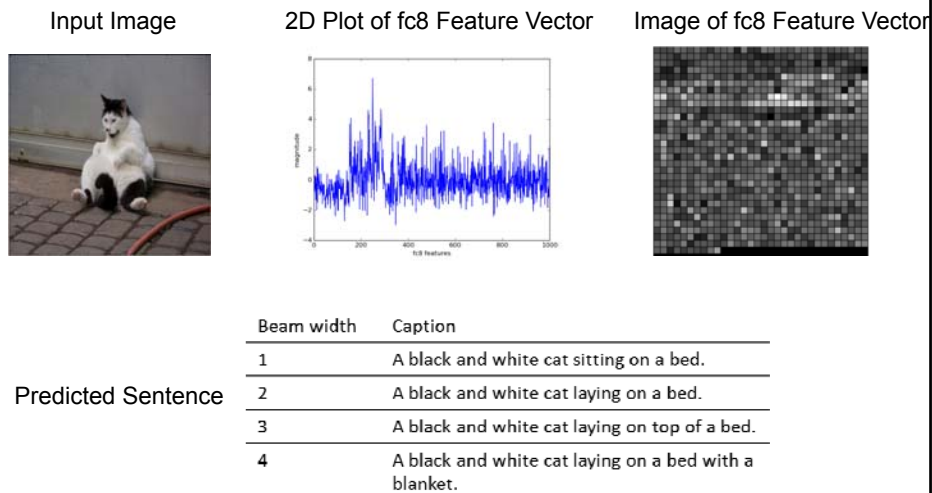
Ptucha '17

65





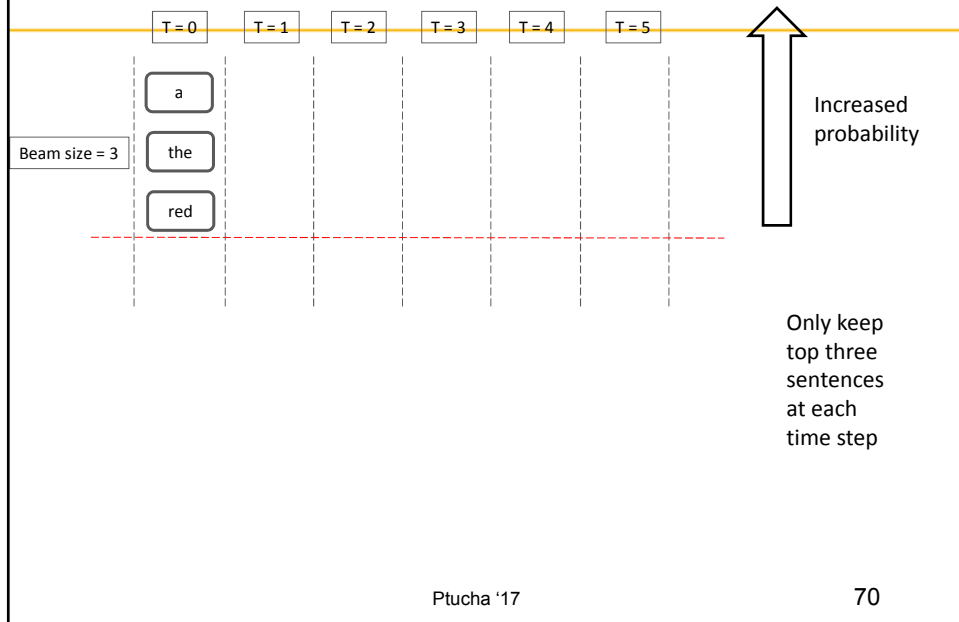
Caption Prediction Example



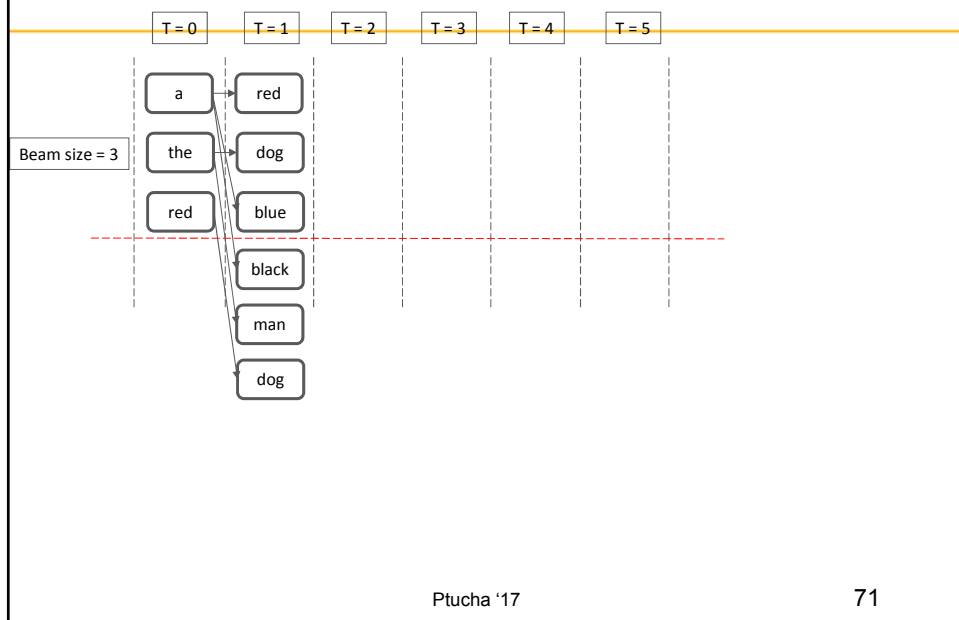
Ptucha '17

69

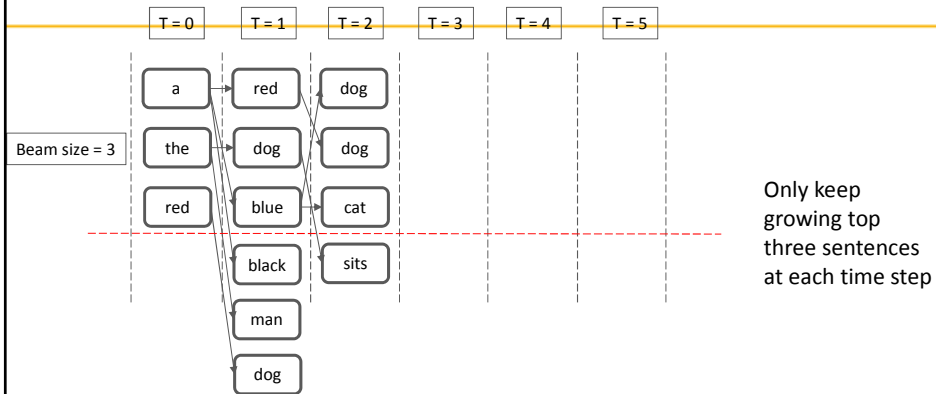
Beam Size Example



Beam Size Example



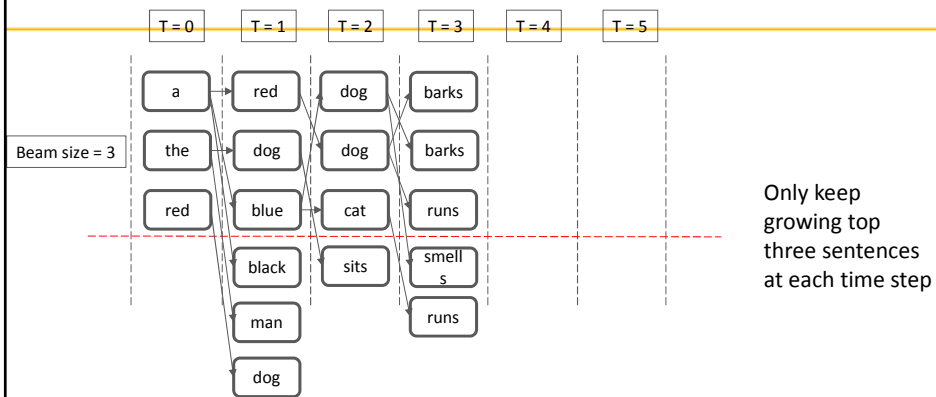
Beam Size Example



Ptucha '17

72

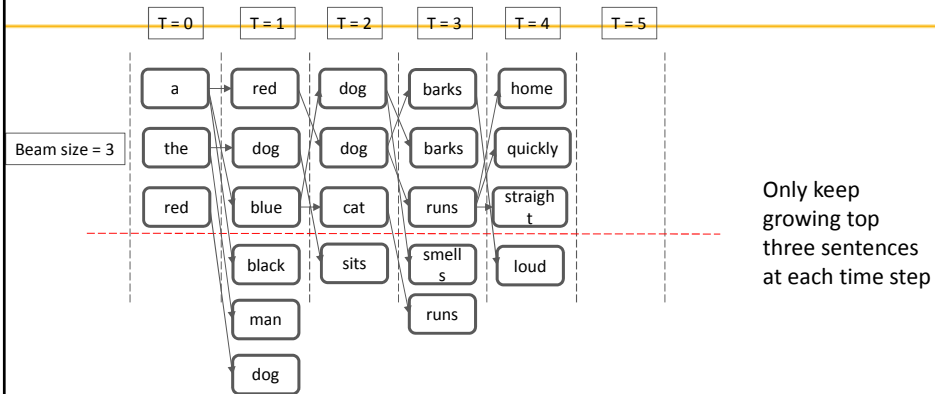
Beam Size Example



Ptucha '17

73

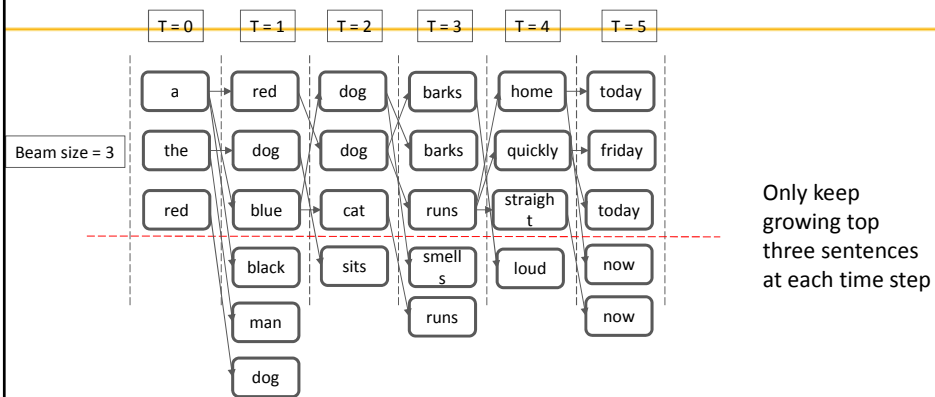
Beam Size Example



Ptucha '17

74

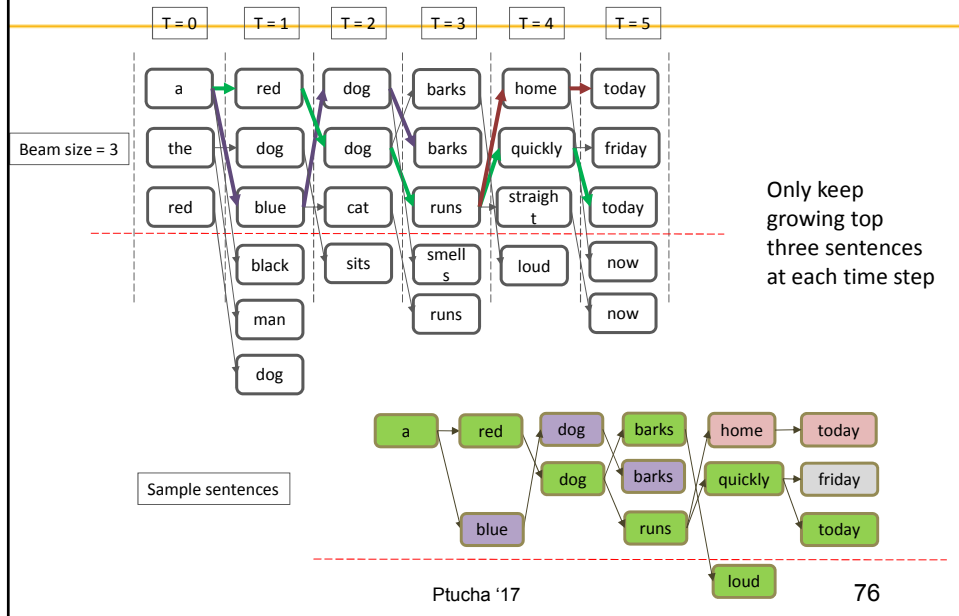
Beam Size Example



Ptucha '17

75

Beam Size Example

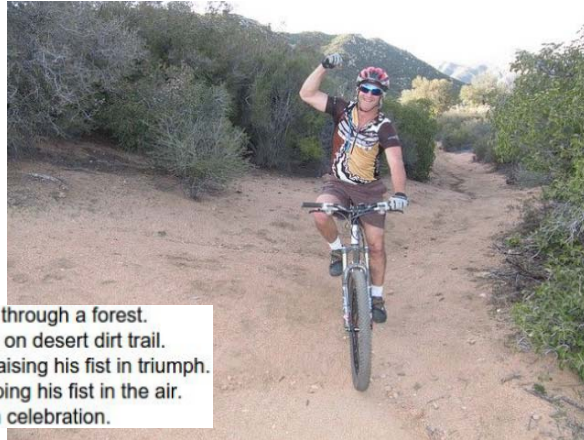


Data for Captioning

- Flickr8K
 - 8,000 images, from Flickr website, each with five captions
 - <http://nlp.cs.illinois.edu/HockenmaierGroup/8k-pictures.html>
- Flickr30K
 - 31,783 images, from Flickr website, each with five captions
 - <http://shannon.cs.illinois.edu/DenotationGraph/>
- MSCOCO
 - 80,000 training images, each with five captions
 - <http://mscoco.org/>

Captioning Datasets

Amazon
mechanical
turkers do all
labeling
<https://www.mturk.com/mturk/welcome>



a man riding a bike on a dirt path through a forest.
bicyclist raises his fist as he rides on desert dirt trail.
this dirt bike rider is smiling and raising his fist in triumph.
a man riding a bicycle while pumping his fist in the air.
a mountain biker pumps his fist in celebration.

Ptucha '17

78

MSCOCO Dataset



- A pile of wooden boxes filled with fruits and vegetables.
- An assortment of fruit in buckets for sale in a shop.
- An outdoor fruit stand with various types of fruits for sale.
- A display of crates of fruit on a city street.
- There are many crates with fruit and vegetables.



- A cat stands on a counter while a dog stands on the floor.
- A cat on the kitchen counter is looking down at a dog.
- A cat is looking at a dog rummage in the garbage.
- A cat on the counter and a dog on the ground in the kitchen.
- A cat stalking a dog on the kitchen floor.

Ptucha '17

79

			
"man in black shirt is playing guitar."	"construction worker in orange safety vest is working on road."	"two young girls are playing with lego toy."	"boy is doing backflip on wakeboard."
			
"a young boy is holding a baseball bat."	"a cat is sitting on a couch with a remote control."	"a woman holding a teddy bear in front of a mirror."	"a horse is standing in the middle of a road."
Ptucha '17	Karpathy'15	80	

Video Data for Captioning

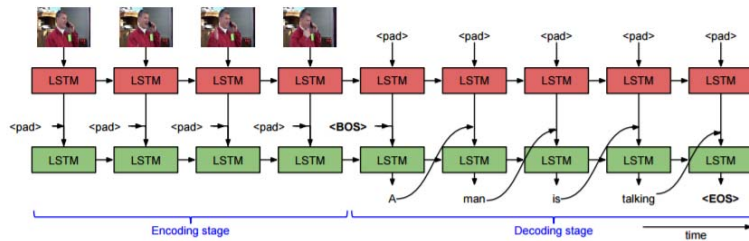
MSVD: Microsoft Video Description Dataset
 MSR-VTT: Microsoft Research -Video to Text
 M-VAD: Movie description dataset M-VAD

	MSVD	MSR-VTT	M-VAD
#sentences	80,827	200,000	54,997
#sent. per video	~42	20	~1-2
vocab. size	9,729	24,282	16,307
avg. length	10.2s	14.8s	5.8s
#train video	1,200	6,513	36,921
#val. video	100	497	4,651
#test video	670	2,990	4,951

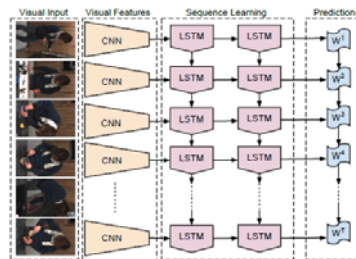
Ptucha '17

81

Image and Language Applications



Venugopalan, 2015

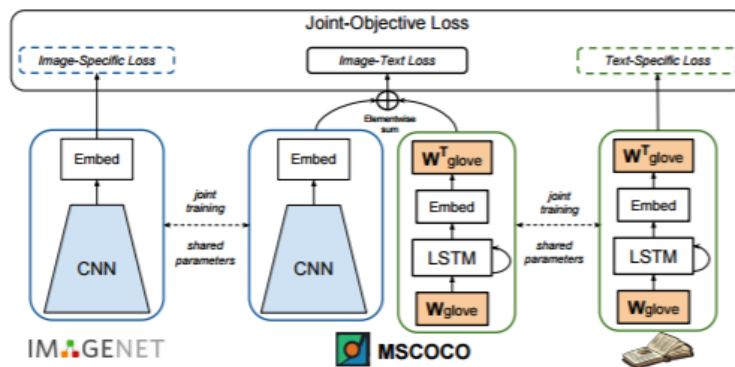


Donahue, et al., 2015

Ptucha '17

82

The Novel Object Captioner (NOC)



Venugopalan 2017

- Visual network (left), language network (right), and caption network (center) are all trained simultaneously with different objectives, but with shared parameters.

Ptucha '17

83

Latest Vision-Language Research from the Machine Intelligence Laboratory



Ptucha '17

84

The Machine Intelligence Lab

- Advanced deep learning research-developing new algorithms and deep architectures.
- Research on tomorrow's smart devices: Machine learning, robotics, HCI, computer vision, NLP, ...
- Teaching a robot how to recognize objects, interact with people, or navigate complex environments.
- All students experts in deep learning.



Ptucha '17

85



Machine Intelligence Lab

- Mostly masters students in computer engineering:
 - Includes 2+2+1/2 PhD students
 - Includes students from computer science, electrical engineering, imaging science, and computer engineering
- 34 publications since 2015
- 2 patent applications
- Deep learning tutorials, consultation, online classes, reading groups.

Ptucha '17

86

Summarization of Long Videos: Tradition computer vision is used to detect interesting regions of video, an encoder-decoder uses CNN representation of frames and kicks out textual descriptions. Traditional text summarizes combine all descriptions into paragraphs.



WACV 2017

Semantic Text Summarization of Long Videos

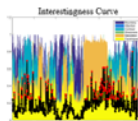
Shagan Sah¹, Sourabh Kulhare¹, Allison Gray¹, Subhashini Venugopalan², Emily Prud'hommeaux¹, Raymond Ptucha¹
¹Rochester Institute of Technology, ²UT Austin
Contact: sss4337@rit.edu



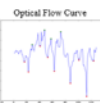
Introduction

- Our work proposes methods to generate both visual and textual summaries of long videos.
- Videos of several hours long are passed into a superframe segmentation algorithm.
- Each superframe segmentation is evaluated for interestingness (boundary motion, superframe motion, contrast, saturation, sharpness, and facial content).
- The annotation module receives temporal segments, and generates captions.
- Captions are then passed into the summarization tool, which outputs a single summary paragraph per video.

Super-Segment Selection



Interestingness is determined by a non-linear combination of scores measuring boundary, attention, contrast, sharpness, saturation, and facial impact.

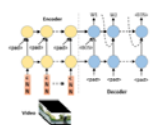


Dataset

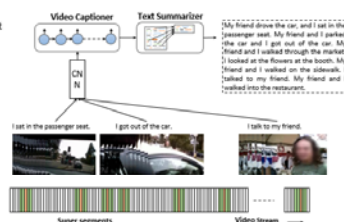
VideoSet dataset with 11 daily life egocentric, Disneyworld egocentric and TV episode videos with duration between 5.5 hrs and 45 mins and captions for every 5-10 seconds. We fine-tune the captioner on 8 videos and test on remaining 3 videos.

Video Captioner

- Encode frames, one frame at a time to the first layer of a two layer LSTM
- Decode this into a natural language sentence one word at a time, feeding the output of one time step in the subsequent time step.



Proposed Method



Results

ROUGE2 scores for machine generated vs ground truth (LSA/LexRank/Sumbank)

Test Video	All Clips	Key Clips
DY01	0.25 / 0.17 / 0.28	0.31 / 0.30 / 0.31
OR03	0.15 / 0.07 / 0.14	0.15 / 0.11 / 0.15
TV04	0.02 / 0.02 / 0.02	0.01 / 0.01 / 0.01

Human evaluations of test summaries (1 (very poor) to 5 (very good))

Video	DY01	OR03	TV04
Video Length	5.5 hours	4 hours	45 min.
Summ. Semantics	3.65	2.35	1.40
Summ. Syntax	3.55	2.40	1.65
Summ. Semantics	3.80	2.35	1.45

Sample Machine Generated Summary

I used my phone while waiting for the train to depart. I looked through the attendant and I rode the train. my friends and I waited for the train to depart. my friends and I stood around the tour guide. my friends and I posed for a group picture. my friends and I talked about our day while walking around the park. my friends and I walked to the <en_sak><en_sak> talking to the theater. my friends and I listened to the tour guide. I talked on my phone while walking around the park. my friends and I talked while moving along the line. my friends and I talked about our foot while walking around the park. my friends and I talked about the camera while walking around the park. my friends and I walked with our group leader through the park while talking. I stood in a dark place and talked to my friends. I watched a mascot entertain I walking. I grabbed some food while moving along the line. my friends and I sat at the table and had dinner. I watched a mascot entertain another group. my friends and I sat at the table and talked...

Ptucha '17

87

Advanced Video Understanding: This paper combines several concepts: Gaussian attention such that input to attention can be varying length; attention steering features for richer steering; global video features; and hierarchical LSTM.



DeepVision
Deep Learning in Computer Vision, RIT

Temporally Steered Gaussian Attention for Video Understanding

Shagan Sah, Thang Nguyen, Miguel Dominguez, Felipe Petroski Such, Raymond Ptucha
Rochester Institute of Technology
Contact: rvp@cs.rit.edu



Attention Steering

- Visual features are computed across the video using ImageNet trained on 4K, classes (extracting 2K pool) layer from ResNet 152.
- Global word embeddings are mean pooled at top EdgeBox locations.
- At each LSTM time step, the model computes an updated frame relevance vector.

Captioning Framework

- We adopt two-layer LSTM hierarchical framework, replacing soft attention with Gaussian attention, and call it HGA.
- Inputs to the first hierarchical layer are clips of 12 frames long.
- Learn parameters, ϕ over each caption with r words, $w_1, \dots, w_r$ ($r=0, r=1$ are special begin and end tokens).

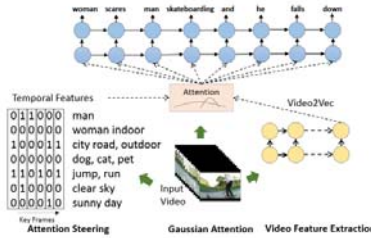
$$\theta^* = \arg \max_{\theta} \sum_{t=1}^T \log(w_t | V_t, S_t, V_{t-1}, w_{t-1}; \theta)$$

Visual word embeddings
Global VideoVec
Prior word

Dataset statistics.

	MSVD	MSR-VTT	MSVD
#frames	10,000	100,000	10,000
frames per video	~42	20	~1.2
length ratio	9.73%	24.26%	16.30%
avg. length	10.2s	14.6s	5.6s
#frames video	1,200	6,313	30,921
Pool. video	100	497	4,651
#total video	670	2,980	4,551

PTUCHA



Gaussian Attention

- To enable the attention mechanism to be independent of video duration, we present a Gaussian attention model which learns a continuous function.

Video Feature Extraction

- A video vector representation is fed to the model to represent global action and context across the entire length of the video.
- Use ActivityNet for human activity understanding to describe over 200 activity classes.
- *RGB and Optical Flow (OF) models used.

Results

MSVD results on the held out test set.

Model	MEV	B-1	B-2	B-3
MF [23]	29.1	-	-	33.2
S2VT [24]	29.8	-	-	-
SA [25]	29.6	-	-	41.9
P-RNN [26]	32.8	81.5	78.4	86.4
HRE-At [27]	33.4	79.2	66.3	35.1
Baseline-GA-5	31.5	80.9	66.9	32.9
HGA	32.8	79.1	65.8	54.8
HGA+vg	30.8	81.8	68.3	35.7
HGA+OF	33.1	77.5	64.1	53.4

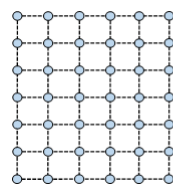
MSR-VTT results on the held out test set.

Model	MEV	REU-1	REU-2	REU-3	REU-4	REU-5	REU-6
Dong et al. [28]	36.9	-	-	-	-	30.3	41.9
McDonald et al. [29]	27.0	-	-	-	-	38.3	41.9
Shen et al. [30]	27.7	-	-	-	-	40.8	40.8
Mean pool	31.4	78.3	66.4	46.4	34.3	39.8	57.7
HGA	31.4	78.7	66.9	47.1	48.9	42.9	58.1
HGA+vg	31.6	78.9	67.2	48.1	47.8	43.8	58.2
HGA+OF	31.6	78.7	66.9	48.0	47.8	43.8	58.2
HGA+OF+vg	31.7	78.9	68.0	48.0	47.8	43.8	58.0
HGA+OF+vg+OF	31.7	78.9	68.0	48.0	47.8	43.8	58.0
HGA+OF+vg+OF+vg	31.7	78.9	68.0	48.0	47.8	43.8	58.0

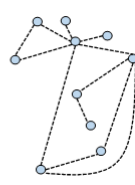
Ptucha '17

88

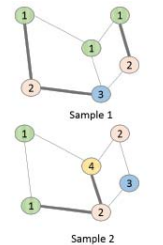
Most of the World's Problems Not Homogeneous



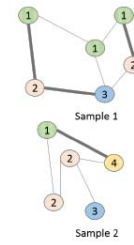
Gridded



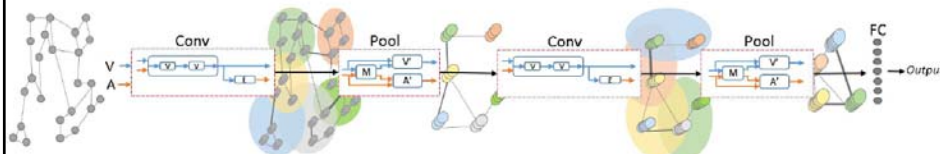
Non-Gridded



Homogeneous



Heterogeneous



Graph-CNN, Petroski Such '17

Ptucha '17

89

More Info: Upcoming Talks

- AI Seminar Series: “Connecting Vision & Language”
 - Mon, Sept 11, 12:20-1:15pm, Bamboo Room (2nd floor off of SAU)
 - Overlap with this talk, but goes into details of latest MIL research
- Brick City Homecoming- “What Is And How Will Machine Learning and Artificial Intelligence Change Our Lives”
 - Sat, Oct 14, 10:30-11:30am, 12:30-1:30pm
 - Intro to machine and deep learning, discuss impact of AI on our lives

Ptucha '17

90



For More Information: <http://www.rit.edu/mil>



Raymond W. Ptucha
Assistant Professor, Computer Engineering
Director, Machine Intelligence Laboratory
Rochester Institute of Technology

Email: rwpeec@rit.edu
Phone: +1 (585) 797-5561
Office: GLE (09) 3441

Ptucha '17

91