



Language
Technologies
Institute

Carnegie
Mellon
University

Advanced Multimodal Machine Learning

Lecture 9.1: Attention models and alignment

Louis-Philippe Morency
Tadas Baltrušaitis

Upcoming Schedule

- First project assignment:
 - Proposal presentation (2/21 and 2/22)
 - First project report (3/5)
- Midterm project assignment
 - Midterm presentations (Tuesday 4/4 & Thursday 4/6)
 - Midterm report (Sunday 4/9)
- Final project assignment
 - Final presentation (5/2 & 5/4)
 - Final report (5/7)

Midterm Presentation Instructions

- 8-9 minute presentations (12-18 slides)
 - +2-3 minutes for written feedback and notes
- All team members should be involved during the presentation.
- The ordering of the presentations (Tuesday vs. Thursday) will be inverted based on proposal presentations.
- The presentations will be from 4:30pm – 6:15 to give everyone more time
 - Let us know if you need to leave before then

Midterm Presentation Instructions

- Motivation and general definition of your research problem (2-3 slides)
- Mathematical formalization of the research problem, including definition of the main variables and overall objective function (2-4 slides)
- Explain at least one multimodal baseline model for your research problem (2-4 slides)
- Present current results of this baseline model on your dataset. You should study the failure cases of the baseline model (2-4 slides)
- Describe the research directions you are planning to explore. Discuss how they will address some of the shortcomings of your baseline model. (2-3 slides)

Midterm Project Report Instructions

- PART 1
 - **Research problem:** describe and motivate the research problem you are planning to work on. Explain why this problem is important for the research community and, if possible, the society in general. Define in generic terms the main computational challenges involved in this research problem.
 - **Related Work:** Present an overview of the work happening in this research area. This section should include about 12-15 citations of prior work, grouped in similar topics. Also, you should present in more details the 3-4 research papers most related to your proposed work. The related work section should end emphasizing how your proposed approach differ from previous work.
 - **Dataset and Input Modalities:** Describe the dataset(s) you are planning to use for this project. If many options exist, please motivate your choice of dataset for this research problem. Describe the input modalities and annotations available in this dataset. Specify which subset of these modalities and annotations you are planning to use.

Midterm Project Report Instructions

- PART 2
 - **Problem statement:** formalize mathematically your research problem. This should include the mathematical definition of the variables involved in your problem.
 - **Multimodal baseline models:** Describe mathematically at least one multimodal baseline model for your research problem.
 - **Experimental methodology:** Describe your experimental methodology for evaluating the multimodal baseline model(s).
 - **Results and Discussion:** Present in tables and/or figures your experimental results. This section should include more than re-running existing baseline models.
 - **Proposed approach:** Describe what models you are planning to test for the final report experiments. Whenever possible, you should write down the loss function of these models, following the same mathematical formulation previously used.

Multi-modal alignment

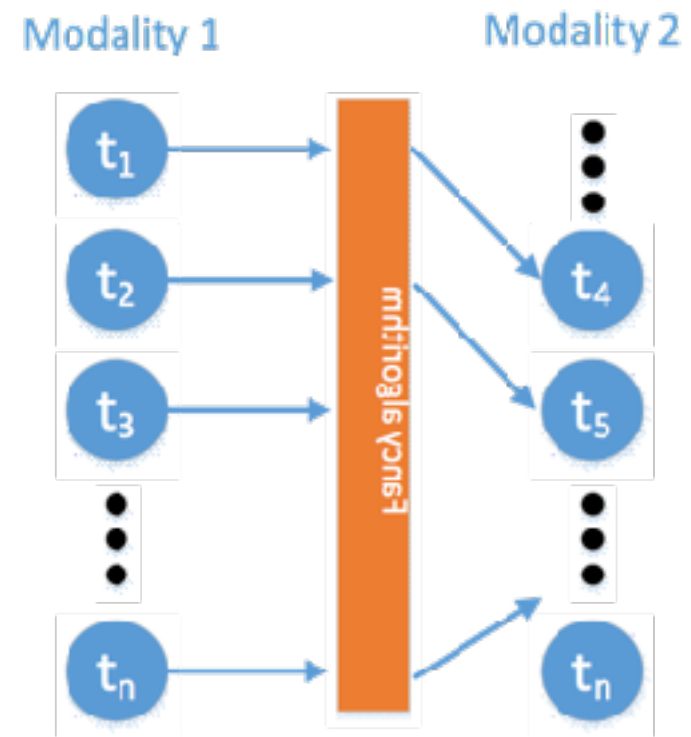


Lecture objectives

- Multimodal alignment
 - Implicit
 - Explicit
- Explicit signal alignment
 - Dynamic Time Warping
 - Canonical Time Warping
- Attention models in deep learning (implicit and explicit alignment)
 - Soft attention
 - Hard attention
 - Spatial Transformer Networks

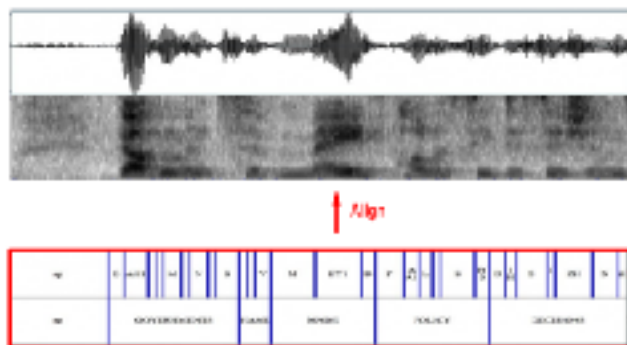
Multimodal-alignment

- Multimodal alignment – finding relationships and correspondences between two or more modalities
- Examples
 - Images with captions
 - Recipe steps with a how-to video
 - Phrases/words of translated sentences
- Two types
 - Explicit – alignment is the task in itself
 - Latent – alignment helps when solving a different task (for example “Attention” models)



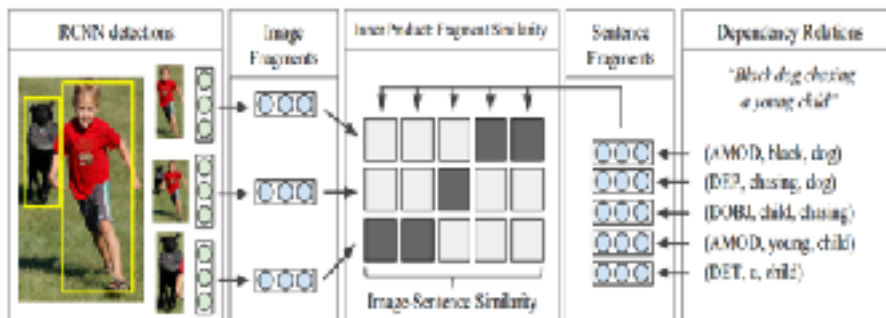
Multimodal-alignment

- Explicit alignment - goal is to find correspondences between modalities
 - Aligning speech signal to a transcript
 - Aligning two out-of sync sequences
 - Co-referring expressions



Multimodal-alignment

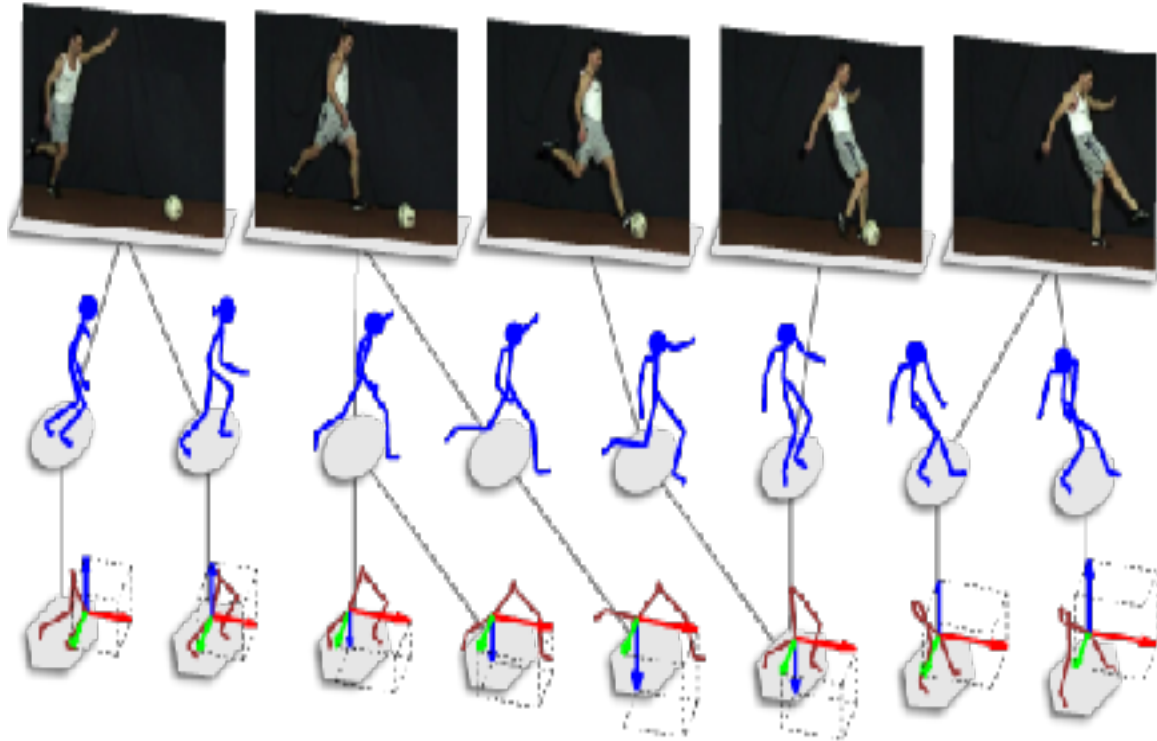
- Implicit alignment - uses internal latent alignment of modalities in order to better solve various problems
 - Machine Translation
 - Cross-modal retrieval
 - Image & Video Captioning
 - Visual Question Answering



Explicit alignment



Temporal sequence alignment



Applications:

- Re-aligning asynchronous data
- Finding similar data across modalities (we can estimate the aligned cost)
- Event reconstruction from multiple sources

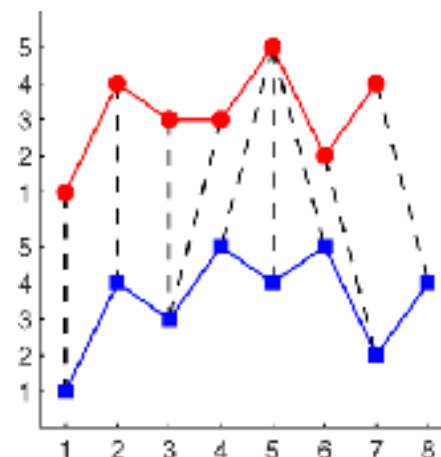
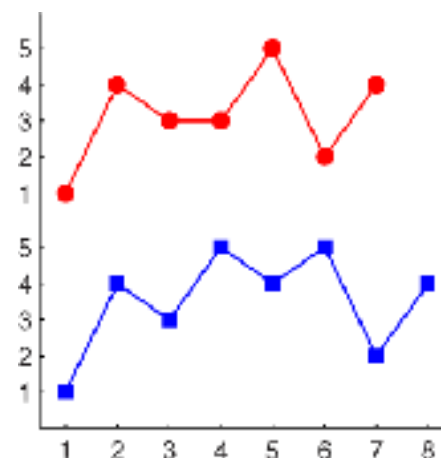


Let's start unimodal – Dynamic Time Warping

- We have two unaligned temporal unimodal signals
 - $\mathbf{X} = [x_1, x_2, \dots, x_{n_x}] \in \mathbb{R}^{d \times n_x}$
 - $\mathbf{Y} = [y_1, y_2, \dots, y_{n_y}] \in \mathbb{R}^{d \times n_y}$
- Find set of indices to minimize the alignment difference:

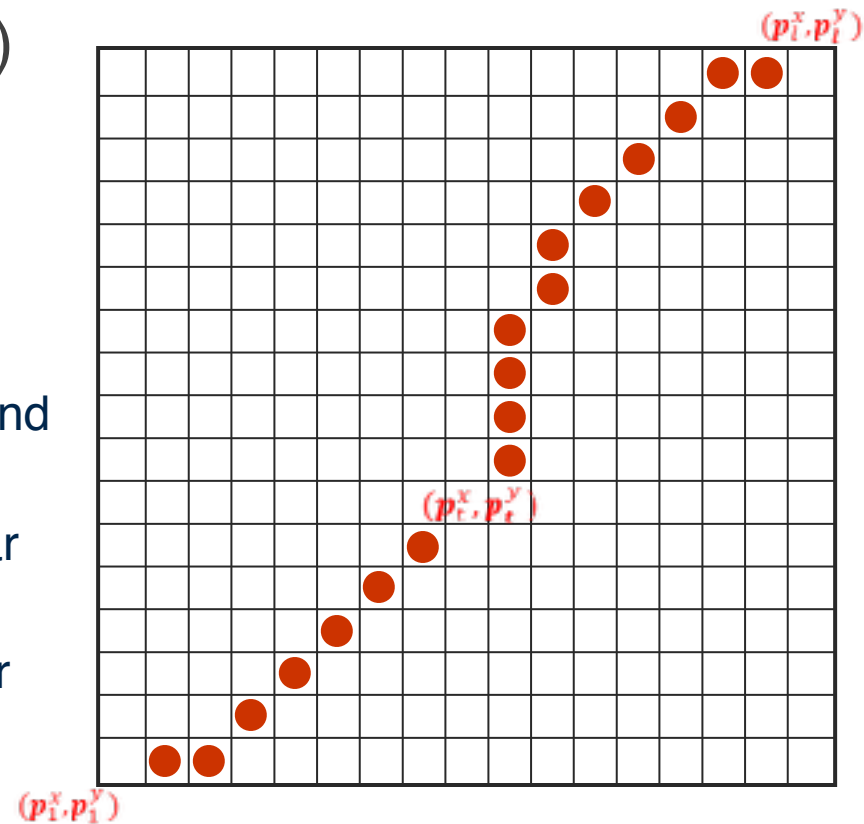
$$L(\mathbf{p}^x, \mathbf{p}^y) = \sum_{t=1}^l \|x_{p_t^x} - y_{p_t^y}\|_2^2$$

- Where $\mathbf{p}^x$ and $\mathbf{p}^y$ are index vectors of same length
- Finding these indices is called Dynamic Time Warping



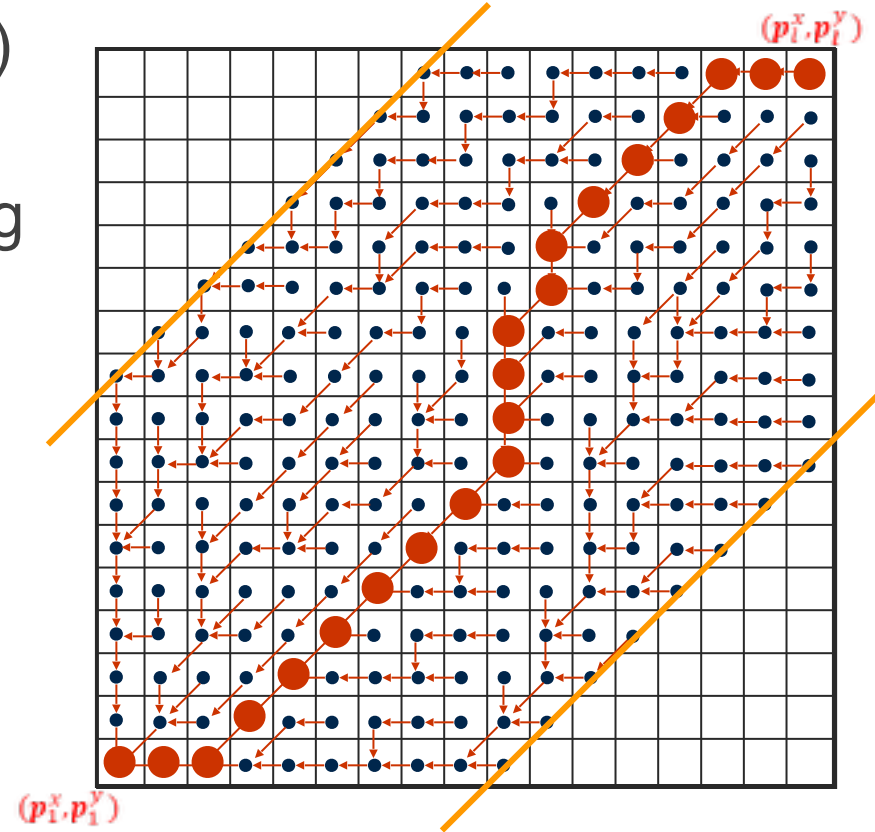
Dynamic Time Warping continued

- Lowest cost path in a cost matrix (expensive to compute)
- Restrictions
 - Monotonicity – no going back in time
 - Continuity - no gaps
 - Boundary conditions - start and end at the same points
 - Warping window - don't get too far from diagonal
 - Slope constraint – do not insert or skip too much



Dynamic Time Warping continued

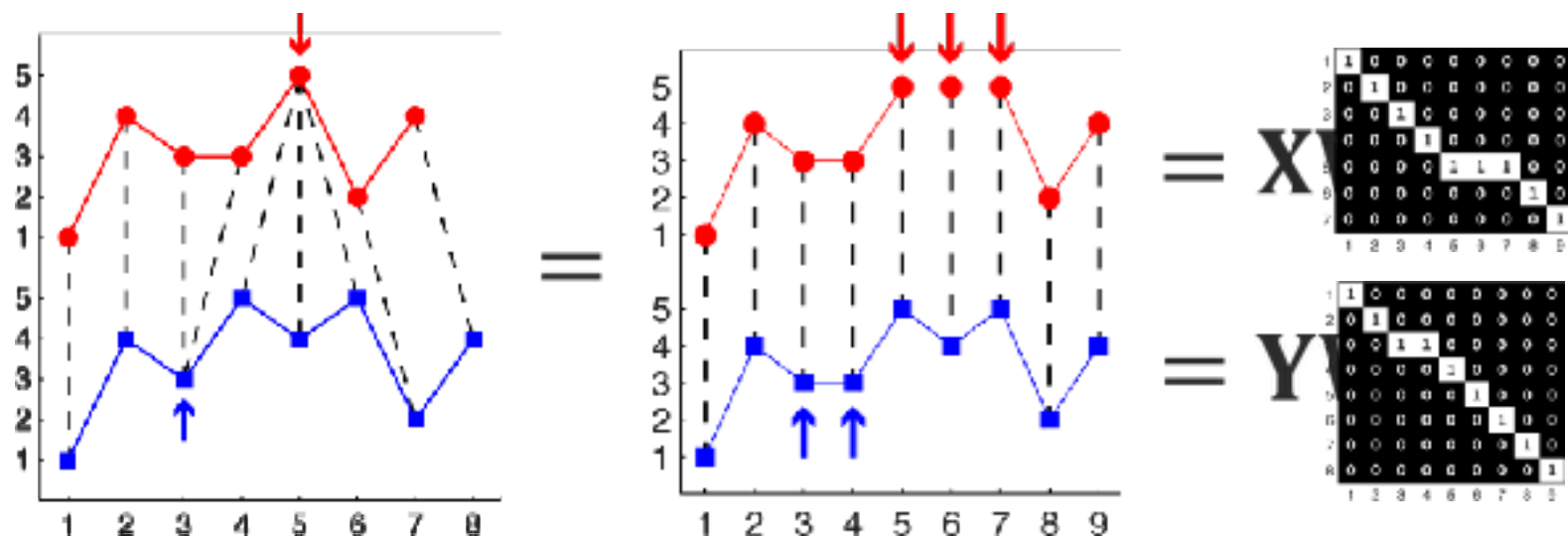
- Lowest cost path in a cost matrix (expensive to compute)
- Solved using dynamic programming whilst respecting the restrictions



DTW alternative formulation

$$L(\mathbf{p}^x, \mathbf{p}^y) = \sum_{t=1}^l \|x_{p_t^x} - y_{p_t^y}\|_2^2$$

Replication doesn't change the objective!



Alternative objective:

$$L(\mathbf{W}_x, \mathbf{W}_y) = \|\mathbf{X}\mathbf{W}_x - \mathbf{Y}\mathbf{W}_y\|_F^2$$

$\mathbf{X}, \mathbf{Y}$ — original signals (same #rows, possibly different #columns)

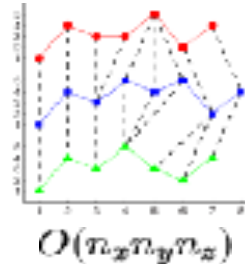
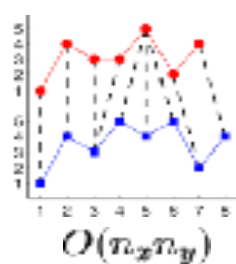
$\mathbf{W}_x, \mathbf{W}_y$ — alignment matrices

Frobenius norm $\|\mathbf{A}\|_F^2 = \sum_i \sum_j \mathbf{A}_{ij}^2$



DTW - limitations

- Computationally complex

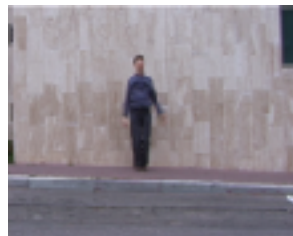


m sequences

$$O\left(\prod_{i=1}^m n_i\right)$$

- Sensitive to outliers

- Unimodal!

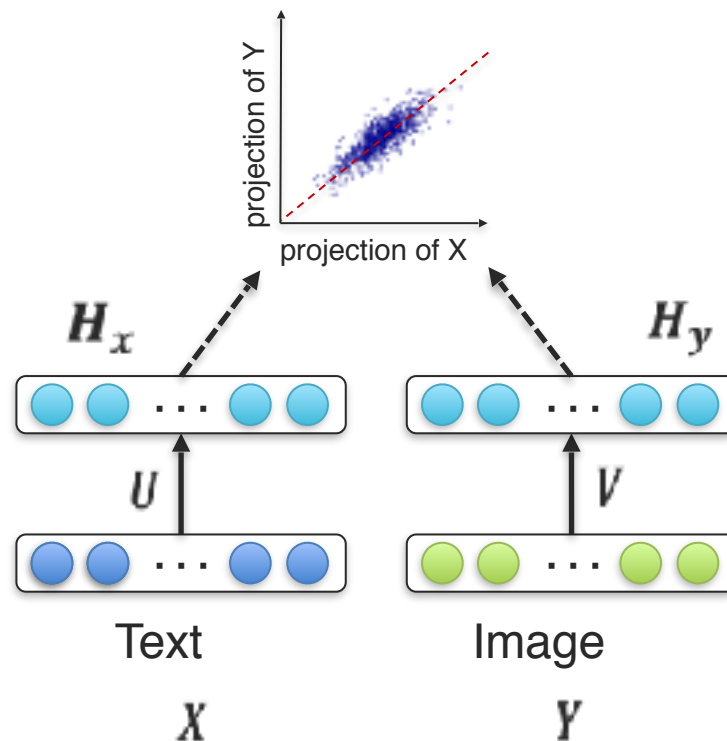


Canonical Correlation Analysis reminder

maximize: $tr(U^T \Sigma_{XY} V)$

subject to: $U^T \Sigma_{YY} U = V^T \Sigma_{YY} V = I$

- 1 Linear projections maximizing correlation
- 2 Orthogonal projections
- 3 Unit variance of the projection vectors

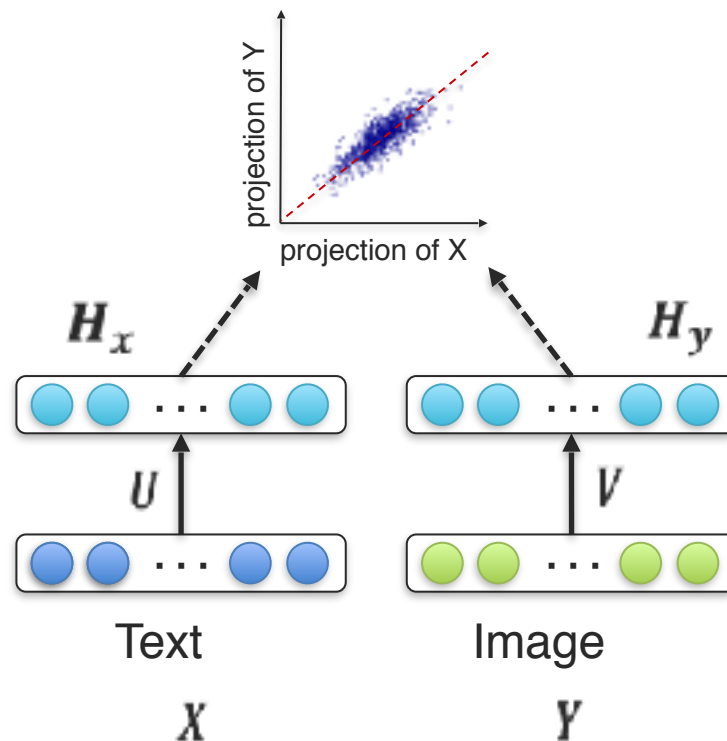


Canonical Correlation Analysis reminder

- When data is normalized it is actually equivalent to smallest RMSE reconstruction
- CCA loss can also be re-written as:

$$L(U, V) = \|U^T X - V^T Y\|_F^2$$

subject to: $U^T \Sigma_{YY} U = V^T \Sigma_{YY} V = I$



Canonical Time Warping

- Dynamic Time Warping + Canonical Correlation Analysis = Canonical Time Warping

$$L(\mathbf{U}, \mathbf{V}, \mathbf{W}_x, \mathbf{W}_y) = \|\mathbf{U}^T \mathbf{X} \mathbf{W}_x - \mathbf{V}^T \mathbf{Y} \mathbf{W}_y\|_F^2$$

- Allows to align multi-modal or multi-view (same modality but from a different point of view)
- $\mathbf{W}_x, \mathbf{W}_y$ – temporal alignment
- $\mathbf{U}, \mathbf{V}$ – cross-modal (spatial) alignment

[Canonical Time Warping for Alignment of Human Behavior, Zhou and De la Tore, 2009]

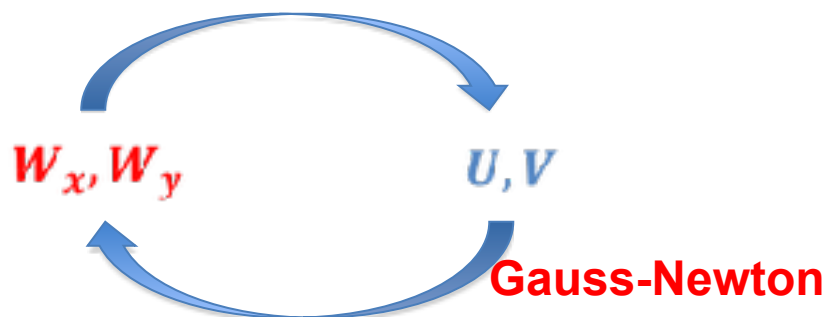


Canonical Time Warping

$$L(\mathbf{U}, \mathbf{V}, \mathbf{W}_x, \mathbf{W}_y) = \|\mathbf{U}^T \mathbf{X} \mathbf{W}_x - \mathbf{V}^T \mathbf{Y} \mathbf{W}_y\|_F^2$$

Optimized by Coordinate-descent – fix one set of parameters, optimize another

Generalized Eigen-decomposition



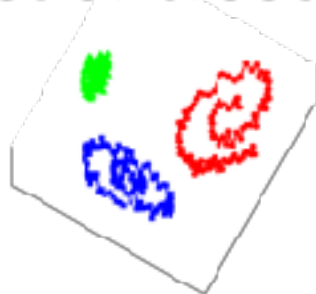
[Canonical Time Warping for Alignment of Human Behavior, Zhou and De la Tore, 2009, NIPS]

Generalized Time warping

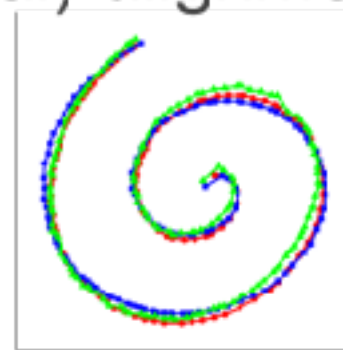
- Generalize to multiple sequences all of different modality

$$L(\mathbf{U}_i, \mathbf{W}_i) = \sum_{i=1} \sum_{j=1} \|\mathbf{U}_i^T \mathbf{x}_i \mathbf{W}_i - \mathbf{U}_j^T \mathbf{x}_j \mathbf{W}_j\|_F^2$$

- $\mathbf{W}_i$ – set of temporal alignments
- $\mathbf{U}_i$ – set of cross-modal (spatial) alignments



(1) Time warping
(2) Spatial embedding



[Generalized Canonical Time Warping, Zhou and De la Tore, 2016, TPAMI]

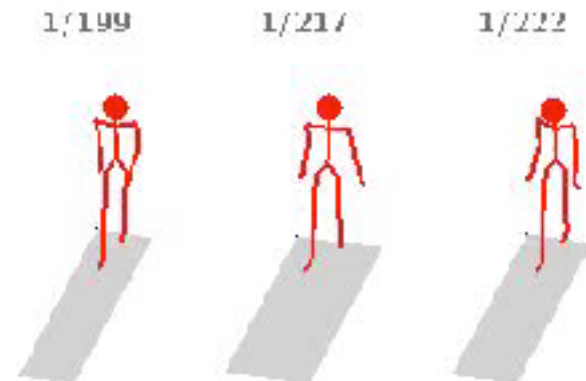
Alignment examples (unimodal)

CMU Motion Capture

Subject 1: 199 frames

Subject 2: 217 frames

Subject 3: 222 frames

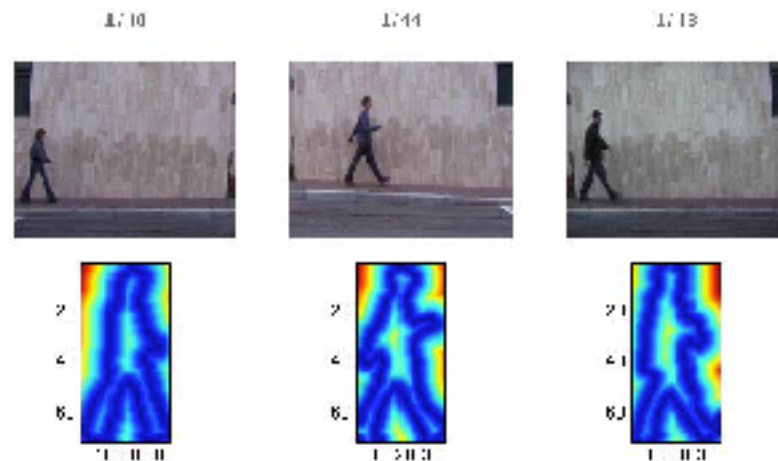


Weizmann

Subject 1: 40 frames

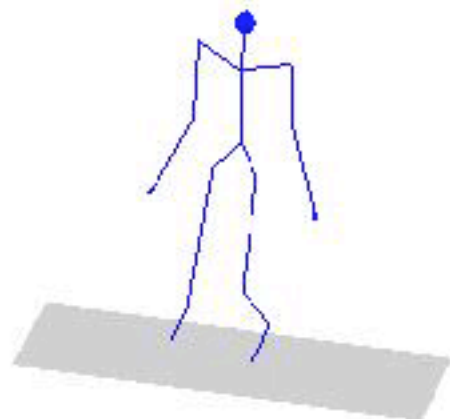
Subject 2: 44 frames

Subject 3: 43 frames



Alignment examples (multimodal)

1/273



1/51



1/127



Canonical time warping - limitations

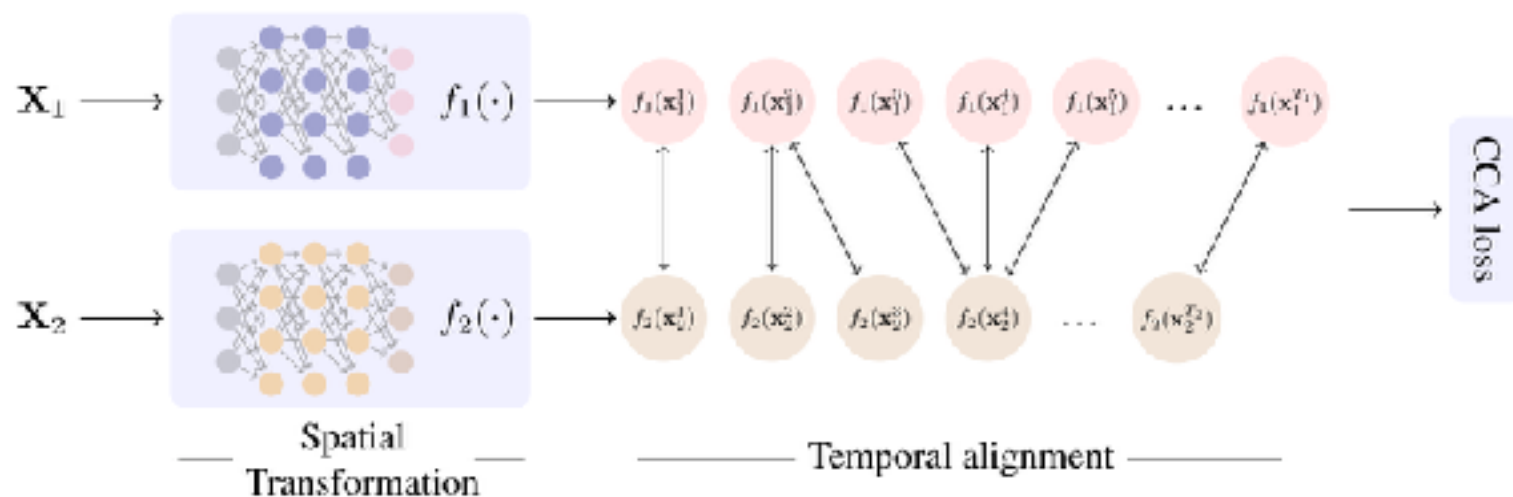
- Linear transform between modalities
- How to address this?



Deep Canonical Time Warping

$$L(\boldsymbol{\theta}_1, \boldsymbol{\theta}_2, \mathbf{W}_x, \mathbf{W}_y) = \|f_{\boldsymbol{\theta}_1}(\mathbf{X})\mathbf{W}_x - f_{\boldsymbol{\theta}_1}(\mathbf{Y})\mathbf{W}_y\|_F^2$$

- Could be seen as generalization of DCCA and GTW



[Deep Canonical Time Warping, Trigeorgis et al., 2016, CVPR]

Deep Canonical Time Warping

- $L(\theta_1, \theta_2, W_x, W_y) = \|f_{\theta_1}(\mathbf{X})W_x - f_{\theta_1}(\mathbf{Y})W_y\|_F^2$
- Optimization is again iterative:
 - Solve for alignment (W_x, W_y) with fixed projections (θ_1, θ_2)
 - Eigen decomposition
 - Solve for projections (θ_1, θ_2) with fixed alignment (W_x, W_y)
 - Gradient descent
 - Repeat till convergence

[Deep Canonical Time Warping, Trigeorgis et al., 2016, CVPR]

Implicit alignment

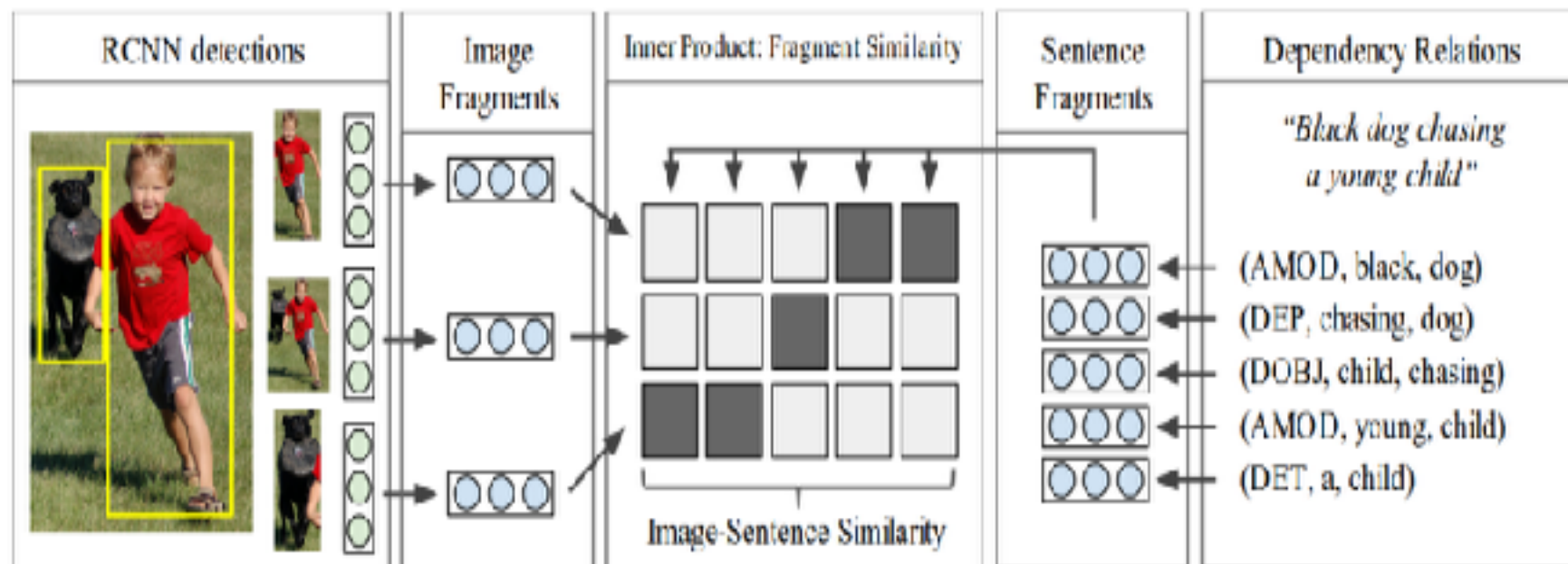


Implicit alignment

- We looked how to explicitly align data
- Could use that as a pre-processing step in our pipelines
- Can we instead allow/encourage the model to align data when solving a particular problem?



Deep fragment embeddings



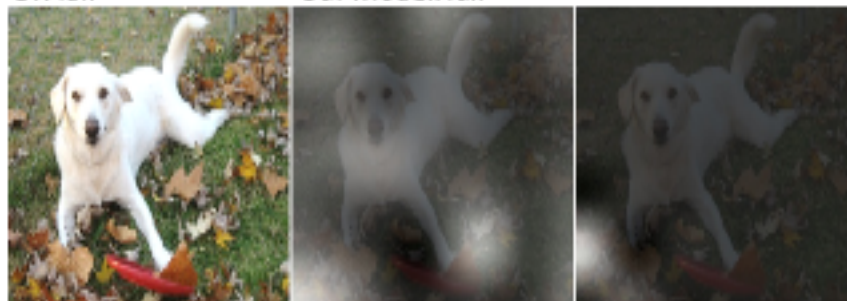
Visual Question Answering

- Xu and Saenko, 2015

What season does this appear to be?

GT: fall

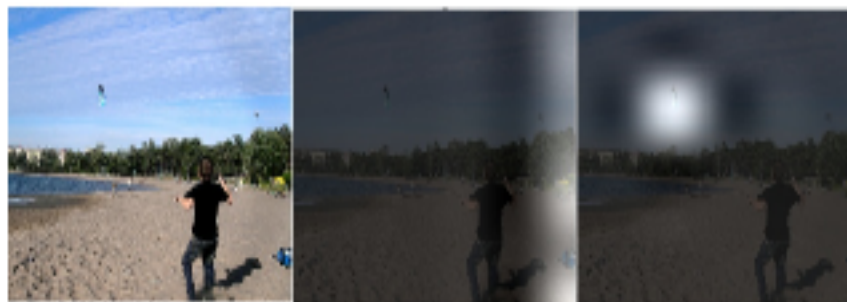
Our Model: fall



What is soaring in the sky?

GT: kite

Our Model: kite

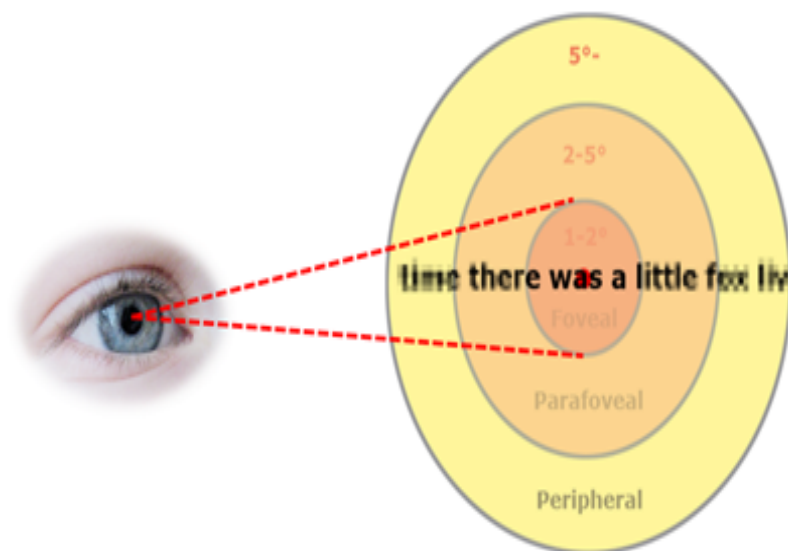


Attention models



Attention in humans

- Foveal vision – we only see in “high resolution” in 2 degrees of vision
- We focus our attention selectively to certain words (for example our names)
- We attend to relevant speech in a noisy room



Attention models in deep learning

- Lots of attention
- Why:
 - Allows for implicit data alignment
 - Good results empirically
 - In some cases faster (don't need to focus on all the image)
 - Better Interpretability



Types of models

- Recent attention models can be roughly split into two major categories
- Soft attention
 - Acts more like a gate function
 - Deterministic
- Hard attention
 - Is more reminiscent of reinforcement learning
 - Stochastic



Soft attention



Machine Translation

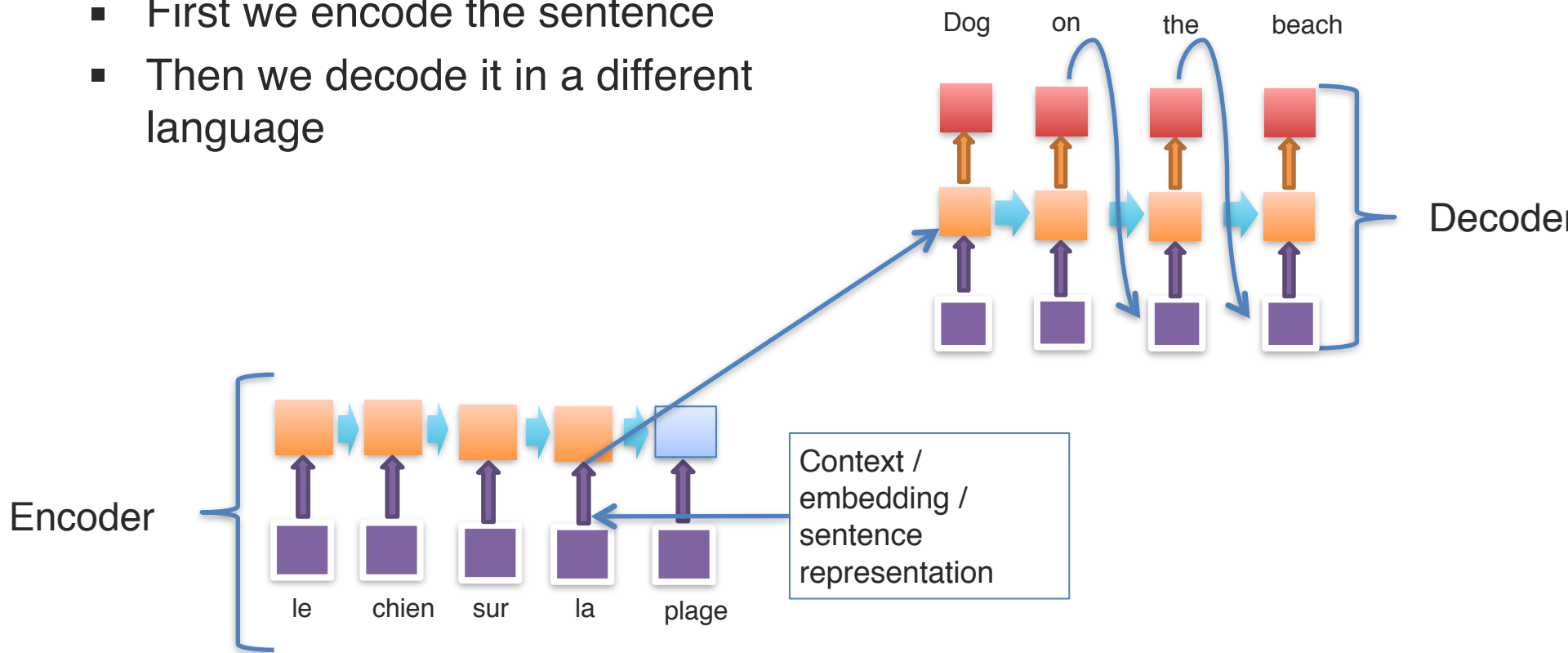
- Given a sentence in one language translate it to another

Dog on the beach → le chien sur la plage

- Not exactly multimodal task – but a good start! Each language can be seen almost as a modality.

Machine Translation with RNNs

- A quick reminder about encoder decoder frameworks
- First we encode the sentence
- Then we decode it in a different language



Machine Translation with RNNs

- What is the problem with this?
- What happens when the sentences are very long?
- We expect the encoders hidden state to capture everything in a sentence, a very complex state in a single vector, such as

The agreement on the European Economic Area was signed in August 1992.

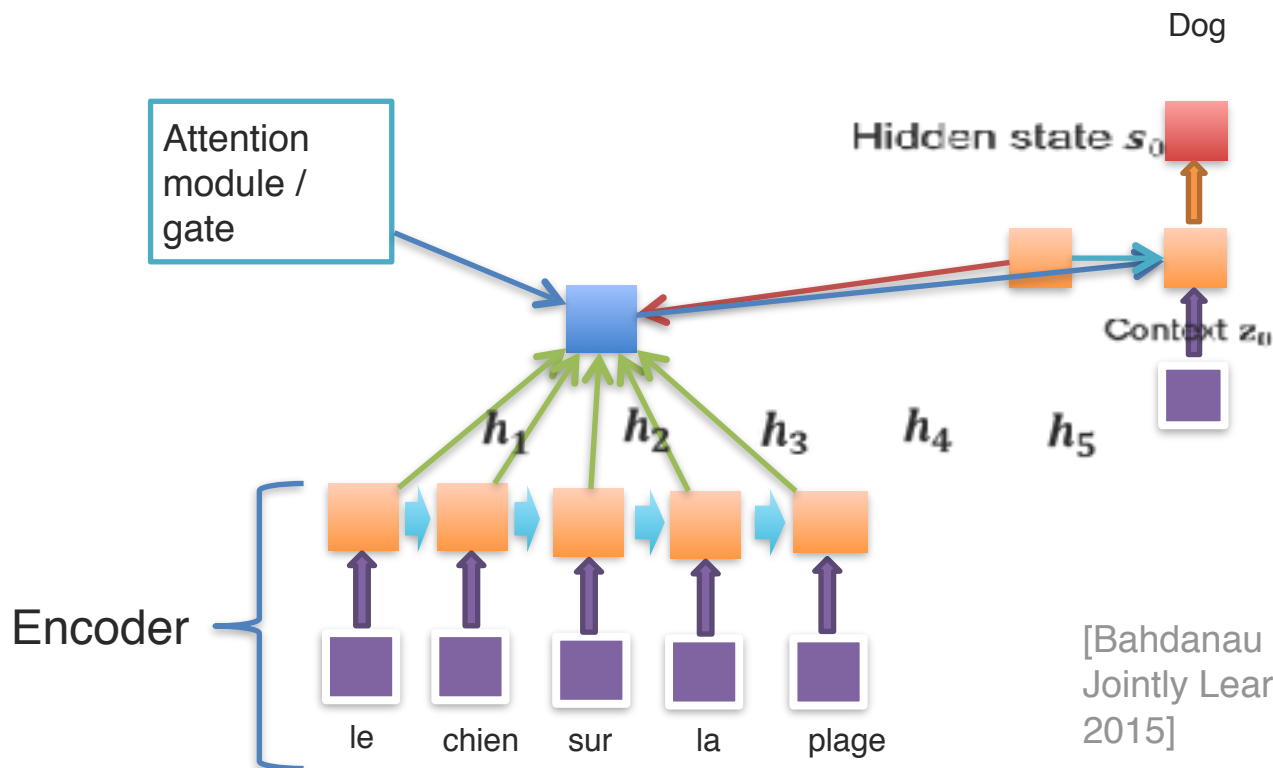


L' accord sur la zone économique européenne a été signé en août 1992.



Decoder – attention model

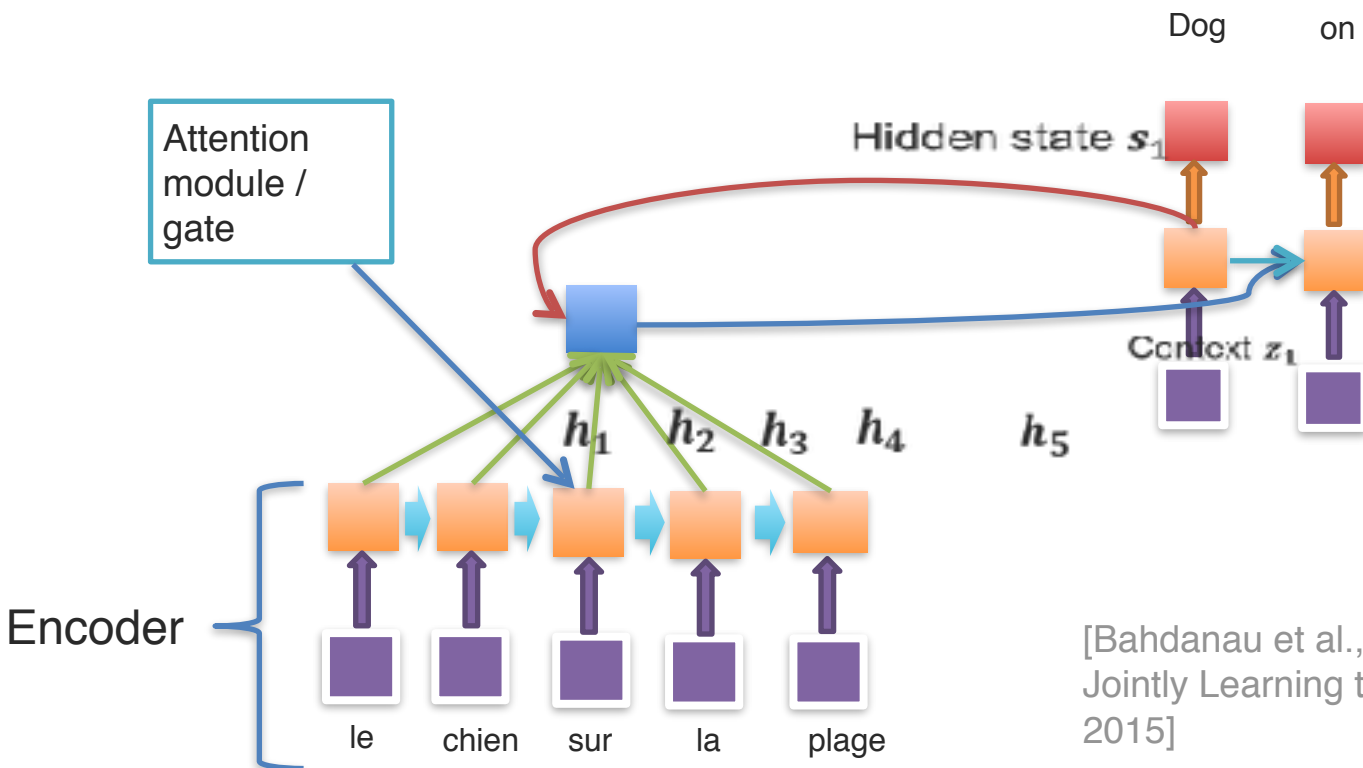
- Before encoder would just take the final hidden state, now we actually care about the intermediate hidden states



[Bahdanau et al., "Neural Machine Translation by Jointly Learning to Align and Translate", ICLR 2015]

Decoder – attention model

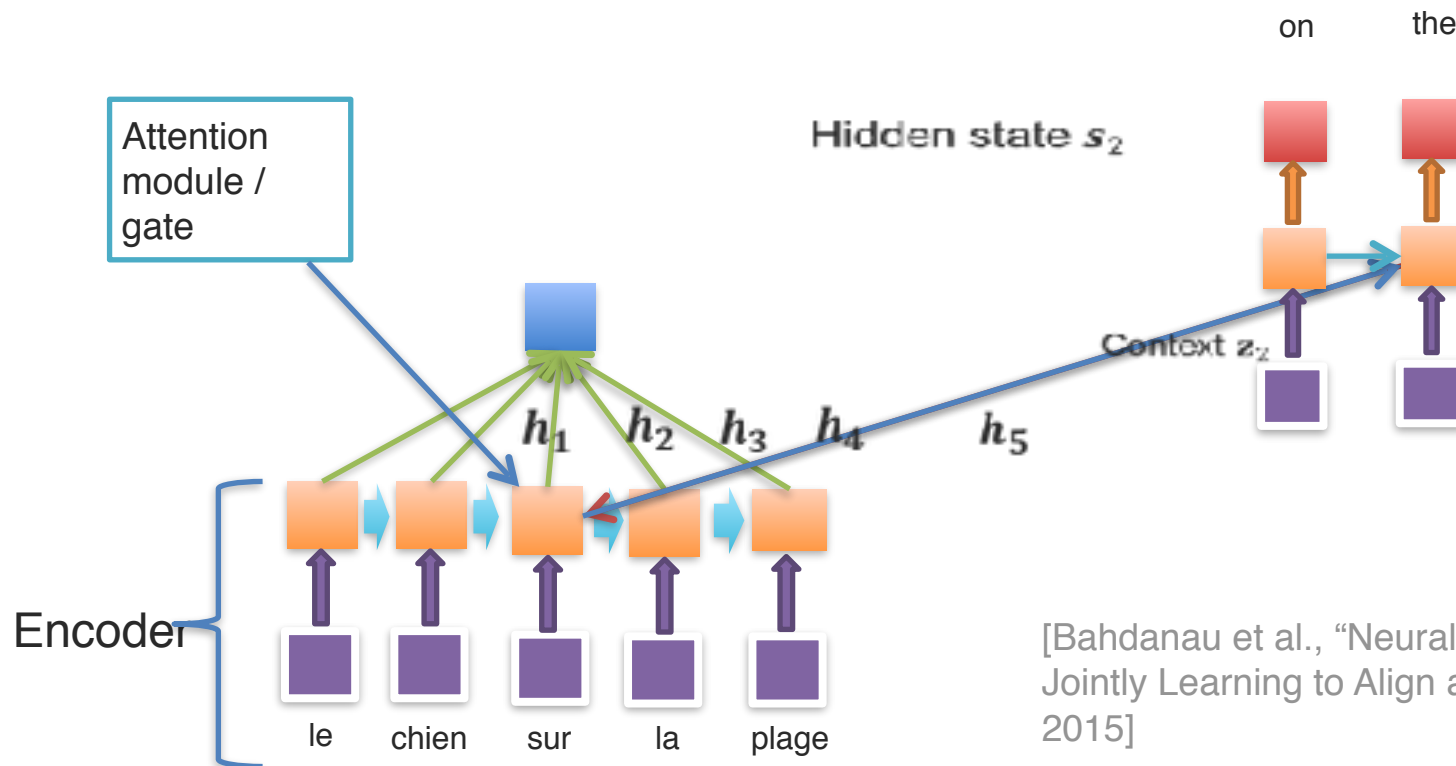
- Before encoder would just take the final hidden state, now we actually care about the intermediate hidden states



[Bahdanau et al., "Neural Machine Translation by Jointly Learning to Align and Translate", ICLR 2015]

Decoder – attention model

- Before encoder would just take the final hidden state, now we actually care about the intermediate hidden states



[Bahdanau et al., "Neural Machine Translation by Jointly Learning to Align and Translate", ICLR 2015]

How do we encode attention

■ Before:

- $p(y_i | y_1, \dots, y_{i-1}, \mathbf{x}) = g(y_{i-1}, \mathbf{s}_i, \mathbf{z})$, where $\mathbf{z} = \mathbf{h}_T$, and $\mathbf{s}_i$ - the current state of the decoder

■ Now:

- $p(y_i | y_1, \dots, y_{i-1}, \mathbf{x}) = g(y_{i-1}, \mathbf{s}_i, \mathbf{z}_i)$
- Have an attention “gate”
 - A different context $\mathbf{z}_i$ used at each time step!
 - $\mathbf{z}_i = \sum_{j=1}^{T_x} \alpha_{ij} \mathbf{h}_j$

α_{ij} - the (scalar) attention for word j at generation step i



MT with attention

- So how do we determine α_{ij} ,

- $\alpha_{ij} = \frac{\exp(e_{ij})}{\sum_{k=1}^{T_x} \exp(e_{ik})}$ - softmax, making sure they sum to 1

- Where:

- $e_{ij} = \mathbf{v}^T \sigma(W \mathbf{s}_{i-1} + U \mathbf{h}_j)$
 - a feedforward network that can tell us given the current state of decoder how important the current encoding is now
 - $\mathbf{v}, W, U$ - learnable weights

- $\mathbf{z}_i = \sum_{j=1}^{T_x} \alpha_{ij} \mathbf{h}_j$

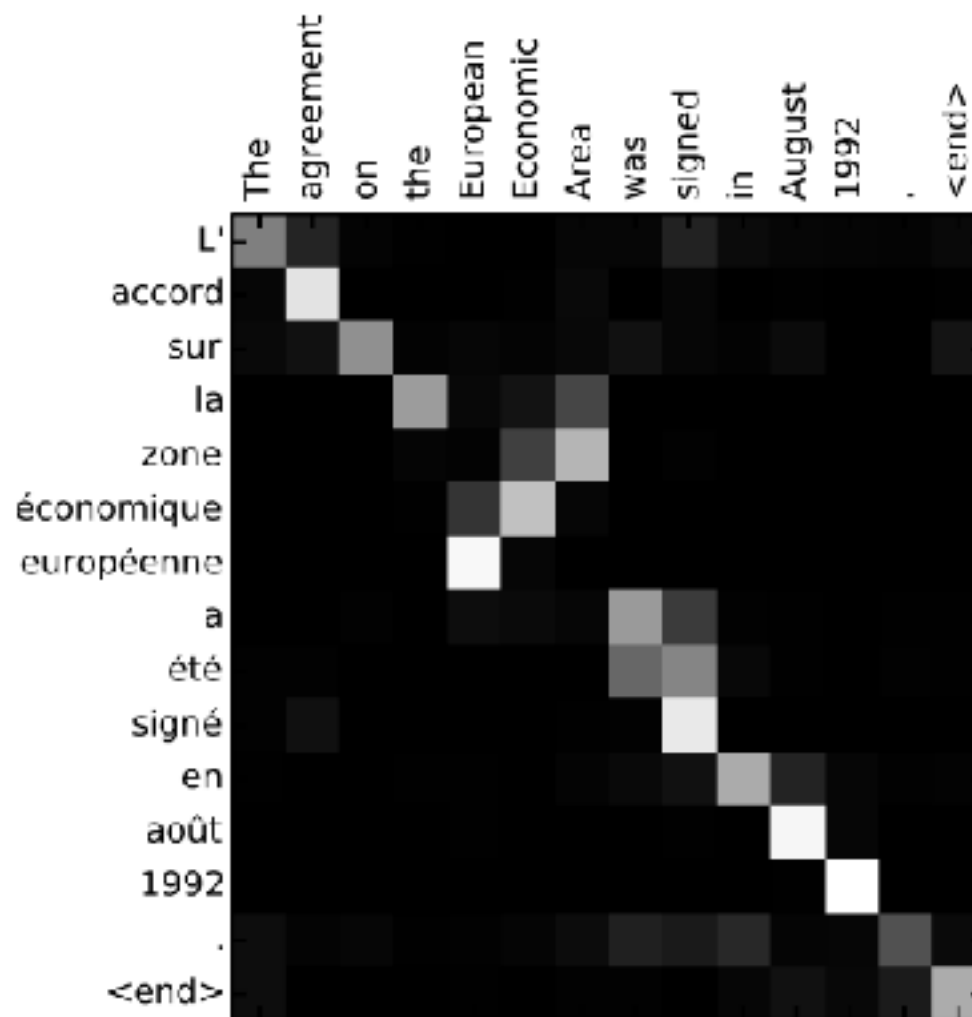
expectation of the context (a fancy way to say it's a weighted average)



MT with attention

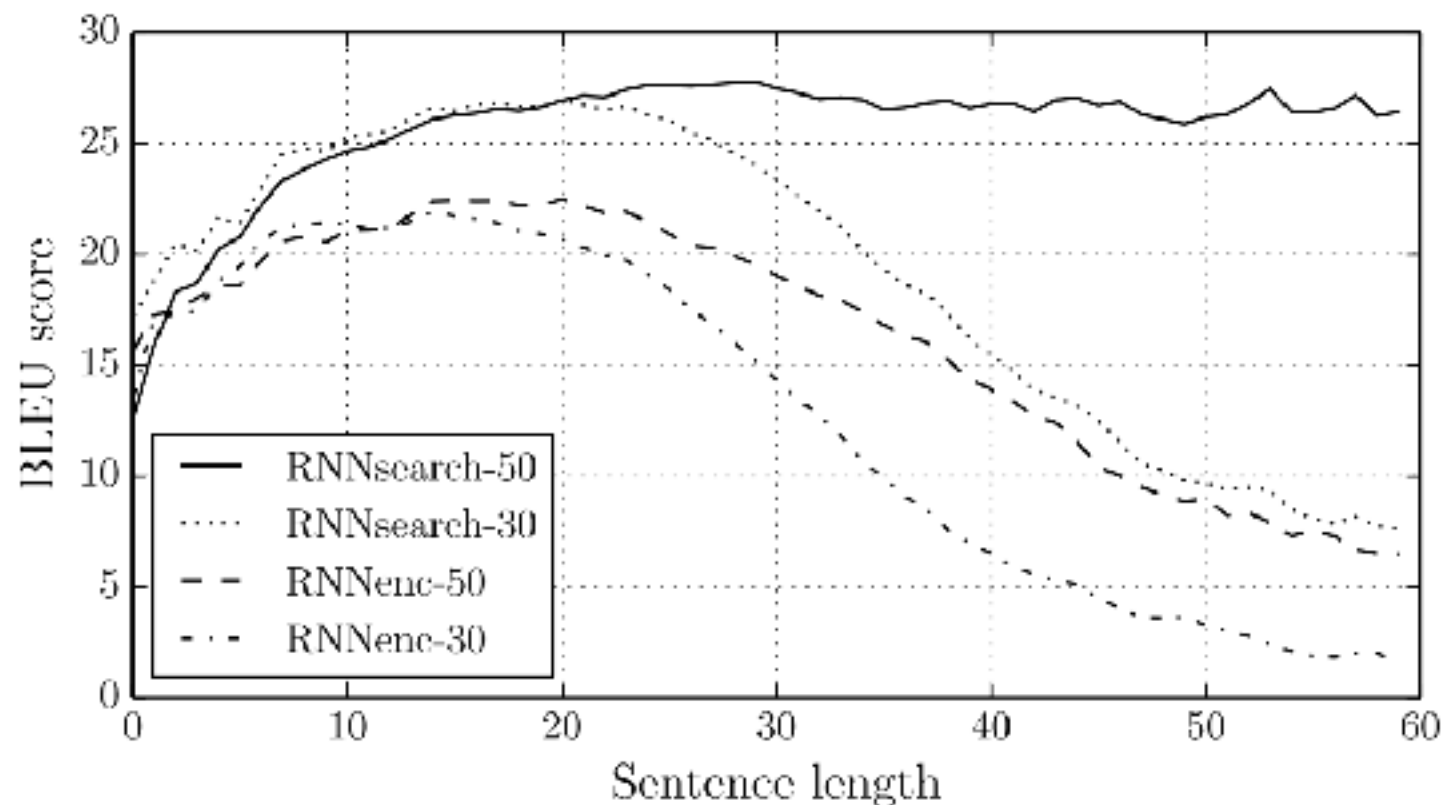
- Basically we are using a neural network to tell us where a neural network should be looking!
- Can use with RNN, LSTM or GRU
- Encoder being used is the same structure as before
 - Can use uni-directional
 - Can use bi-directional
- Model can be trained using our regular back-propagation through time, all of the modules are differentiable

Does it work?



Does it work

- Especially good for long sentences



MT with attention recap

- Get good translation results (especially for long sentences)
- Also get a (soft) alignment of sentences in different languages
 - Extra interpretability of method functioning
- How do we move to multimodal?

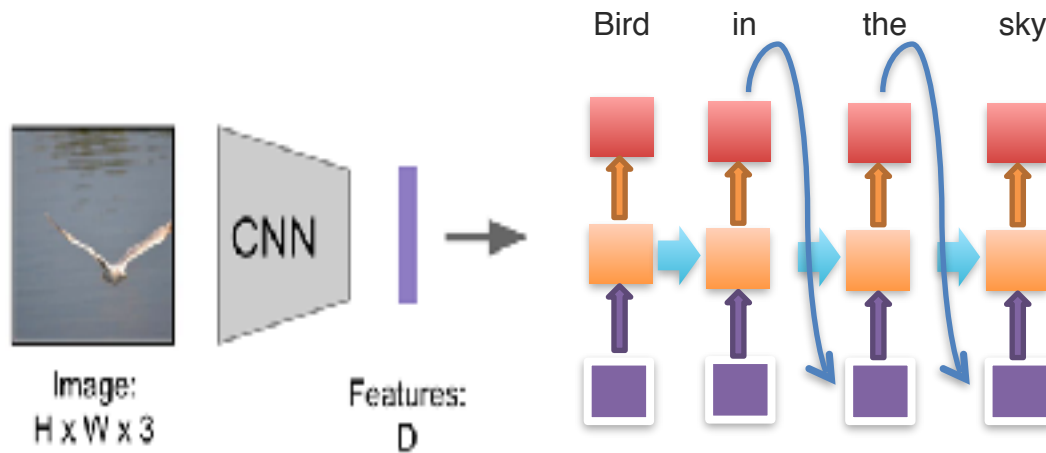
Visual captioning with soft attention

- A similar model but with a visual modality over which we pay attention
- This week's reading group paper



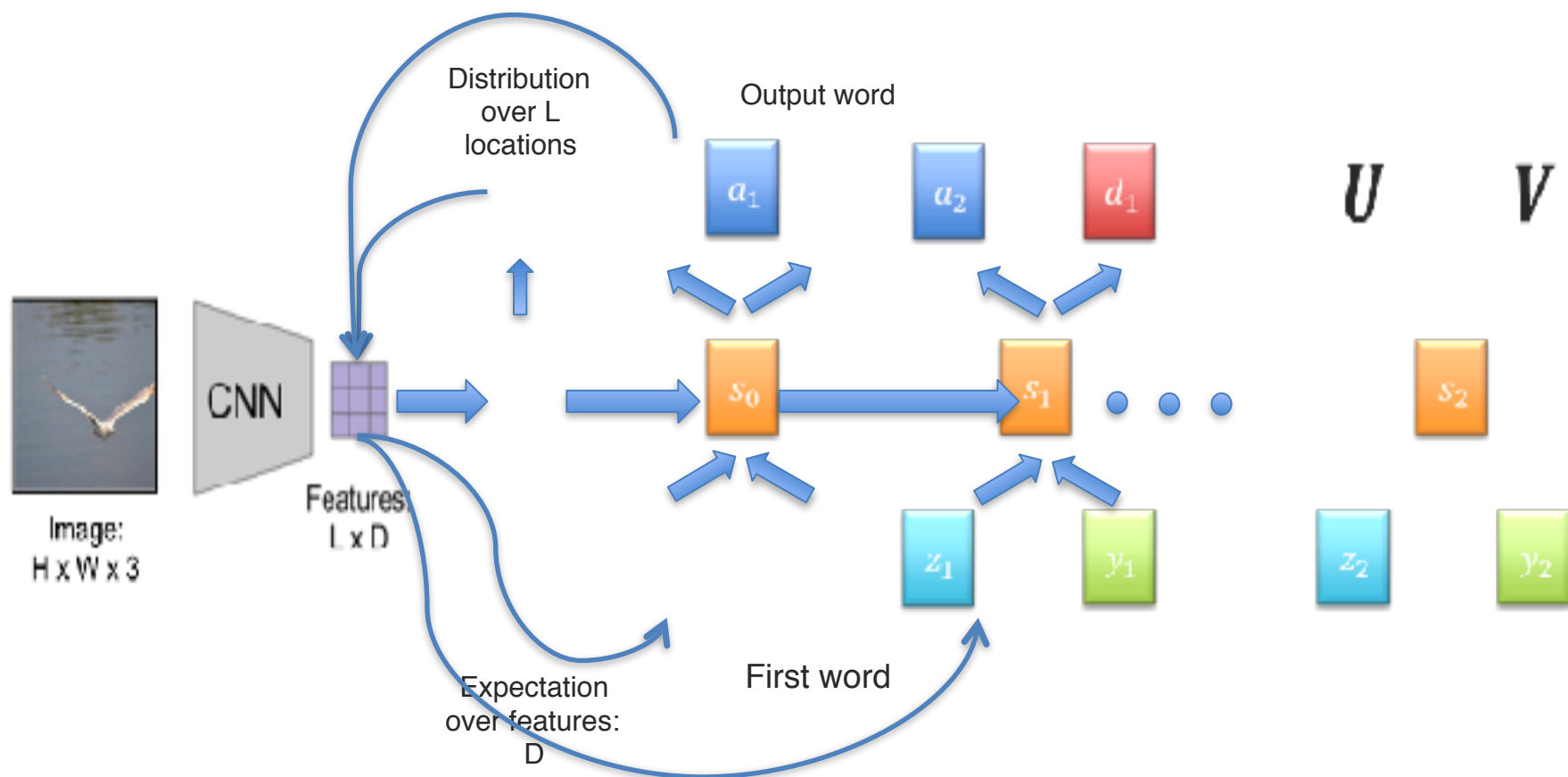
[Show, Attend and Tell: Neural Image Caption Generation with Visual Attention, Xu et al., 2015]

Recap RNN for Captioning



Why might we not want to focus on the final layer?

Looking at more fine grained features



Visual captioning with soft attention

- Works pretty well – outperforms a number of baselines
- Allows us to get an idea of what the network “sees”
- A very similar model to that used for translation
- As well can be optimized using back propagation
 - Code is available <https://github.com/kelvinxu/arctic-captions>
- Also explored hard attention (our next topic)

Other attention work

- Good at paper naming!
 - Show, Attend and Tell (extension of Show and Tell)
 - Listen, Attend and Walk
 - Listen, Attend and Spell
 - Ask, Attend and Answer



Video Descriptions

- Yao et al. 2015



+Local+Global: A **man** and a **woman** are **talking** on the **road**



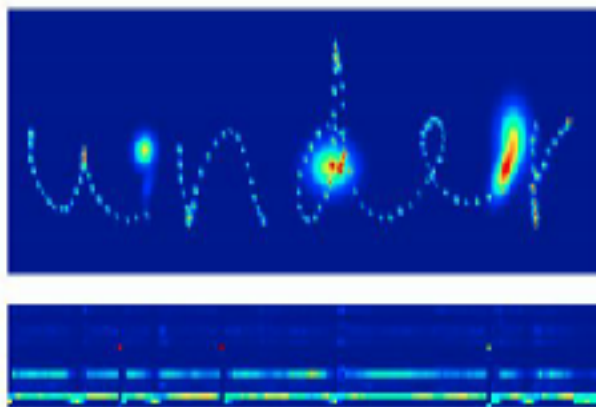
+Local+Global: the **girl** **grins** at **him**

- Soft attention model

Speech

- Speech transcription – Listen, Attend and Spell [Chan et al.]
- Speech recognition - Attention-Based Models for Speech Recognition [Chorowski et al.]

Generative models with attention



Groves, "Generating Sequences with Recurrent Neural Networks", arXiv 2013

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

- DRAW paper from DeepMind

Hard attention



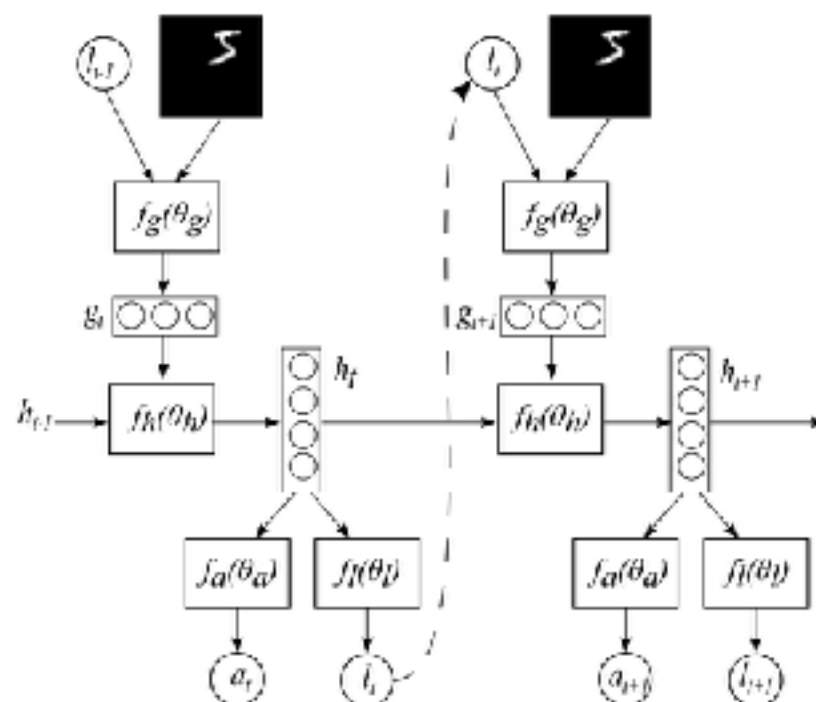
Hard attention

- Soft attention requires computing a representation for the whole image or sentence
- Hard attention on the other hand forces looking only at one part
- Main motivation was reduced computational cost rather than improved accuracy (although that happens a bit as well)
- **Saccade followed by a glimpse – how human visual system works**

[Recurrent Models of Visual Attention, Mnih, 2014]
[Multiple Object Recognition with Visual Attention, Ba, 2015]

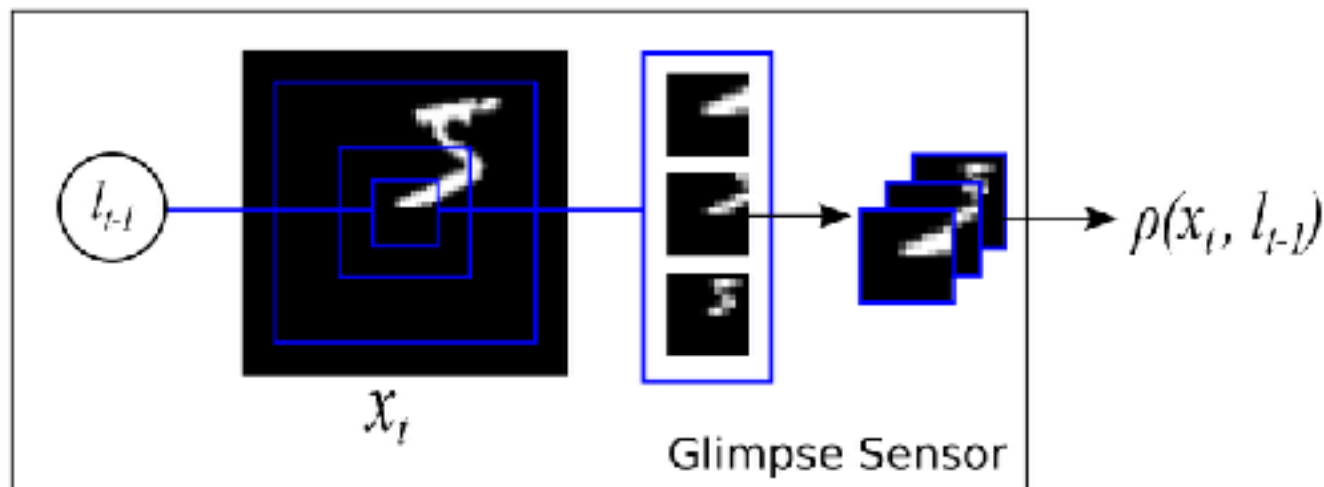
Emission network

- Given an image and a glimpse/attention location the model produces next glimpse location l_t , and optionally an action a_t
- Action can be:
 - Some action in a dynamic system – press a button etc.
 - Classification of an object
 - Word output
- This is an RNN with two output gates and a slightly more complex input gate!



Glimpse

- Looking at a part of an image at different scales

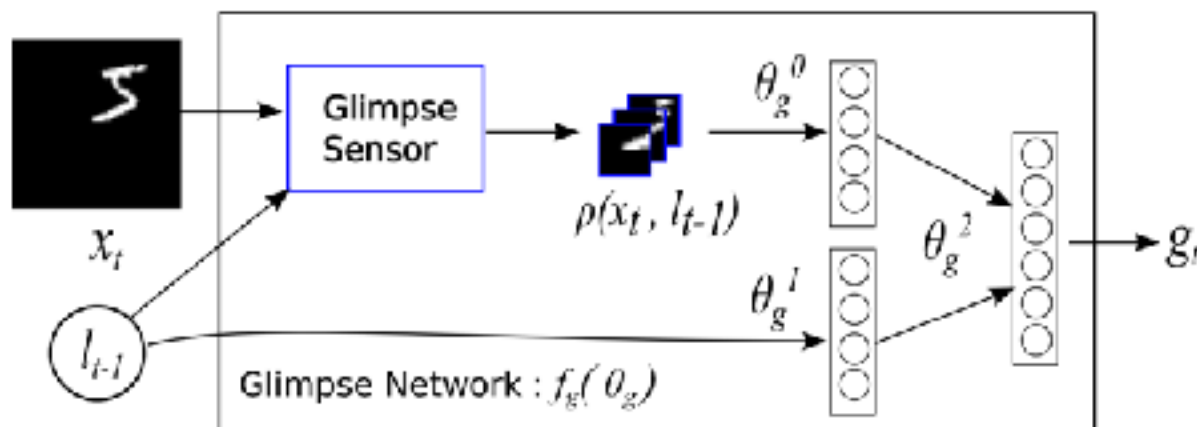


- At a number of different scales combined to a single multichannel image (human retina like representation)
- Given a location l_t output an image summary at that location

[Recurrent Models of Visual Attention, Mnih, 2014]

Glimpse network

- Combining the Glimpse and the location of the glimpse into a joint network



- The glimpse is followed by a feedforward network (CNN or a DNN)
- The exact formulation of how the location and appearance are combined varies, the important thing is combining **what** and **where**
- Differentiable with respect to glimpse parameters but not the location

Recurrent model of Visual Attention (RAM)

■ Model definition

$$L = \log[p(y|X, W)] = \log \sum_l p(l|X, W) p(y|l, X, W)$$

W – set of parameters of the RNN ($\theta_g, \theta_l, \theta_a$)

X the input (image, frame etc.)

l – the set of actions and locations

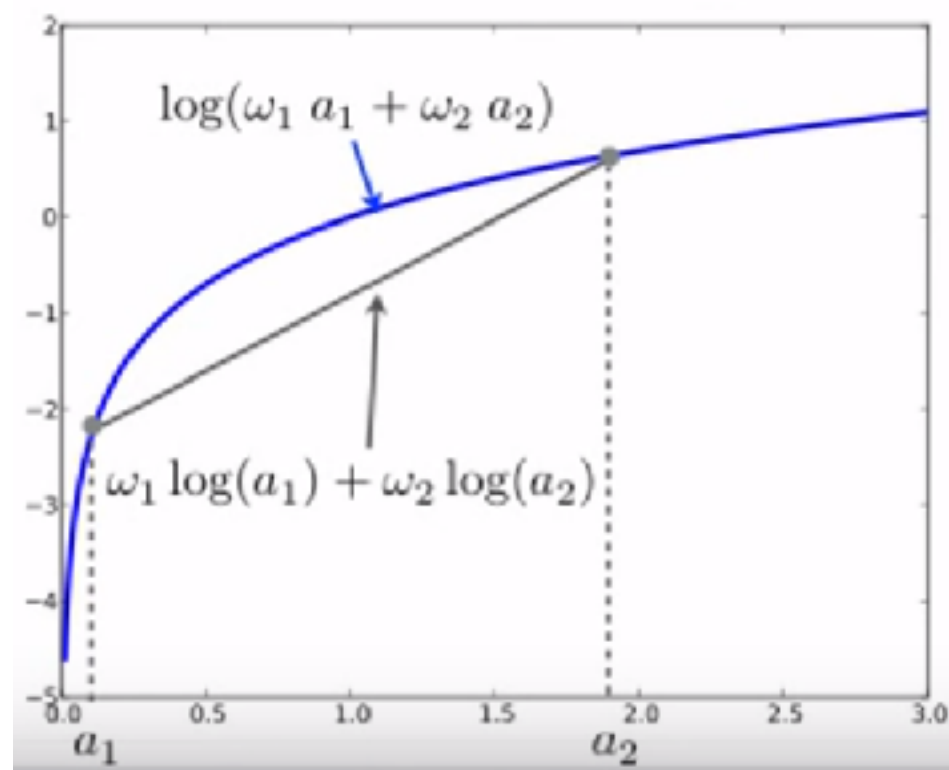
y – correct output (digit classification), correct word prediction etc.

- How do we sum across all of them?
 - It seems summing is the biggest problem in ML
 - Inability to sum was also the main issue in RBM and DBN training



Jensen's Inequality

- For further steps we need and aside
 - Weighted average of concave functions
- $\log(\sum_i w_i a_i) \geq \sum_i w_i \log(a_i)$
- If $\sum_i w_i = 1$ and $w_i \geq 0$



Variational bound

- Going back to our formulation of our loss:

$$L = \log[p(y|X, W)] = \log \sum_l p(l|X, W) p(y|l, X, W)$$

$$\log \sum_l p(l|X, W) p(y|l, X, W) \geq \sum_l p(l|X, W) \log[p(l|X, W)] = F$$

- F is called variational free energy lower bound
- By maximizing it we are indirectly maximizing the true likelihood



Optimization

$$\frac{\partial F}{\partial W} = \sum_l p(l|X, W) \left[\frac{\partial \log p(y|l, X, W)}{\partial W} + \log p(y|l, X, W) \frac{\partial \log p(l|X, W)}{\partial W} \right]$$

Still have to sum across all possible glimpses, so stochastic version:

$$\frac{\partial F}{\partial W} = \frac{1}{M} \sum_{m=1}^M \left[\frac{\partial \log p(y|\tilde{l}^m, X, W)}{\partial W} + \log p(y|\tilde{l}^m, X, W) \frac{\partial \log p(\tilde{l}^m|X, W)}{\partial W} \right]$$

Drawing M samples from some prior (Gaussian, Bernoulli), for example taking a glimpse or by attending to a particular region for caption generation



Optimization

$$\frac{\partial F}{\partial W} = \sum_l p(l|X, W) \left[\frac{\partial \log p(y|l, X, W)}{\partial W} + \log p(y|l, X, W) \frac{\partial \log p(l|X, W)}{\partial W} \right]$$

Still have to sum across all possible glimpses, so use a stochastic version:

$$\frac{\partial F}{\partial W} = \frac{1}{M} \sum_{m=1}^M \left[\frac{\partial \log p(y|\tilde{l}^m, X, W)}{\partial W} + \underbrace{\log p(y|\tilde{l}^m, X, W)}_{\text{unbounded}} \frac{\partial \log p(\tilde{l}^m|X, W)}{\partial W} \right]$$

Unbounded as can have very low negative values

Often replace it with an indicator function instead (Deep RAM and RAM)

Or use a moving average (Thursday's reading)

Optimization

$$\frac{\partial F}{\partial W} = \sum_l p(l|X, W) \left[\frac{\partial \log p(y|l, X, W)}{\partial W} + \log p(y|l, X, W) \frac{\partial \log p(l|X, W)}{\partial W} \right]$$

Still have to sum across all possible glimpses, so stochastic version:

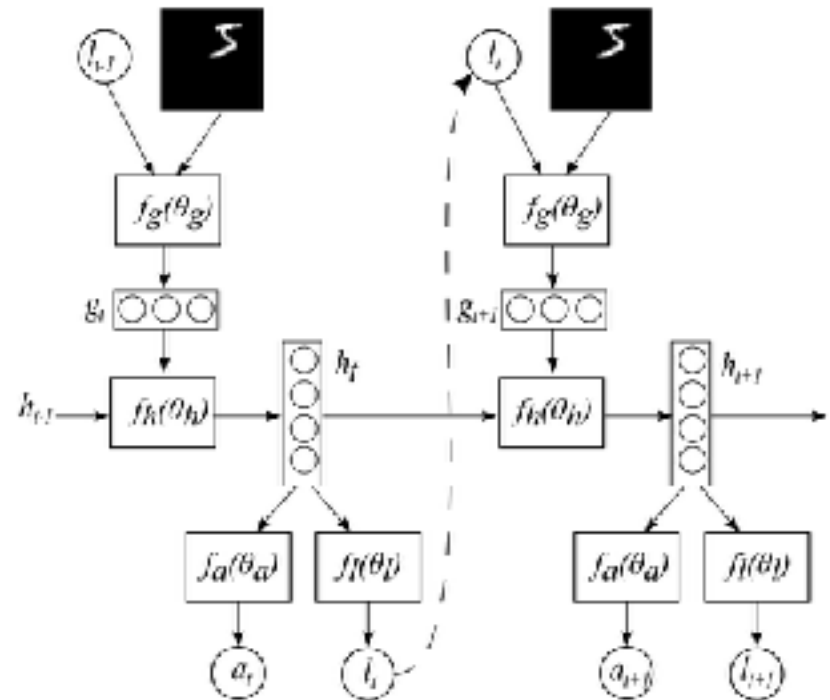
$$\frac{\partial F}{\partial W} = \frac{1}{M} \sum_{m=1}^M \left[\frac{\partial \log p(y|\tilde{l}^m, X, W)}{\partial W} + \log p(y|\tilde{l}^m, X, W) \frac{\partial \log p(\tilde{l}^m|X, W)}{\partial W} \right]$$

To optimize our model sample stochastically for each iteration, then update the weights. Rinse and repeat.



Putting everything together

- Sample locations of glimpses leading to updates in the network
- Use gradient descent to update the weights (the glimpse network weights are differentiable)
- The emission network is an RNN
- Not as simple as backprop but doable
- Some similarities to Restricted Boltzmann Machine training

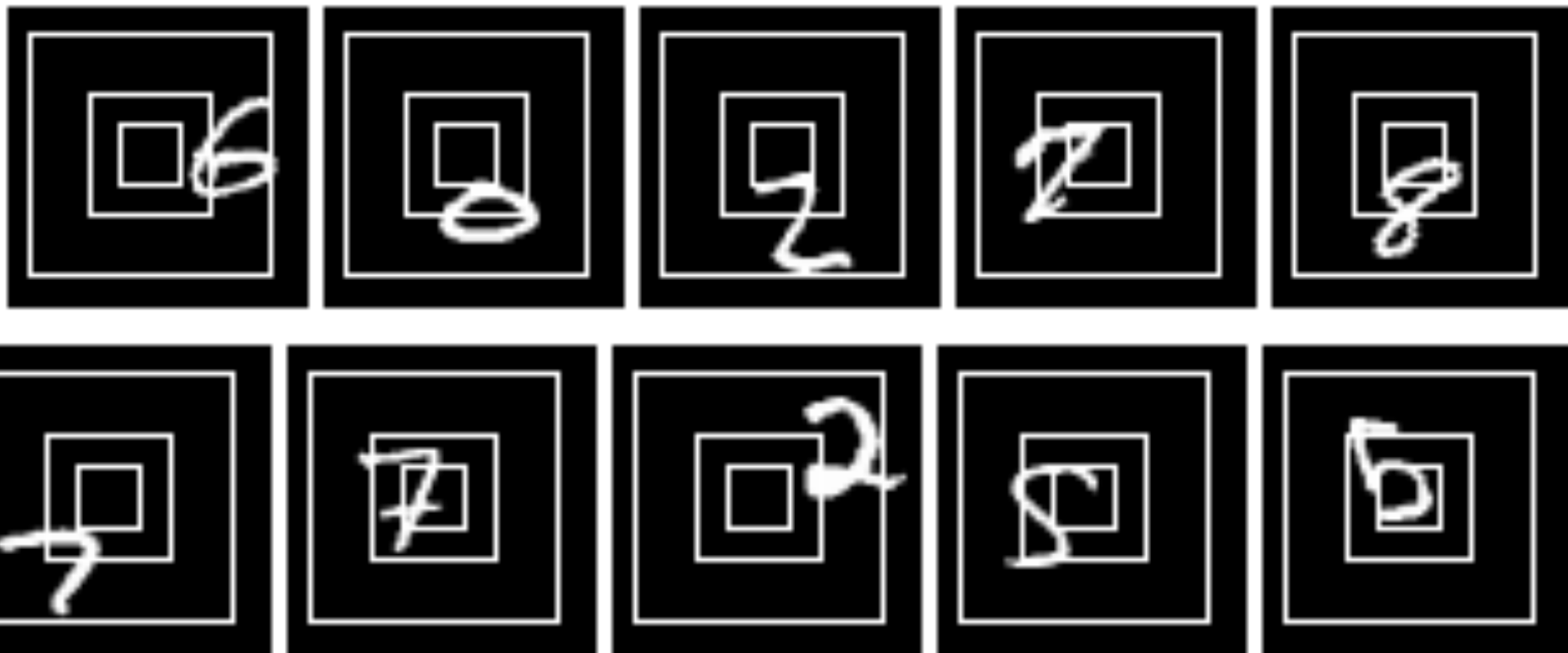


Reinforcement learning

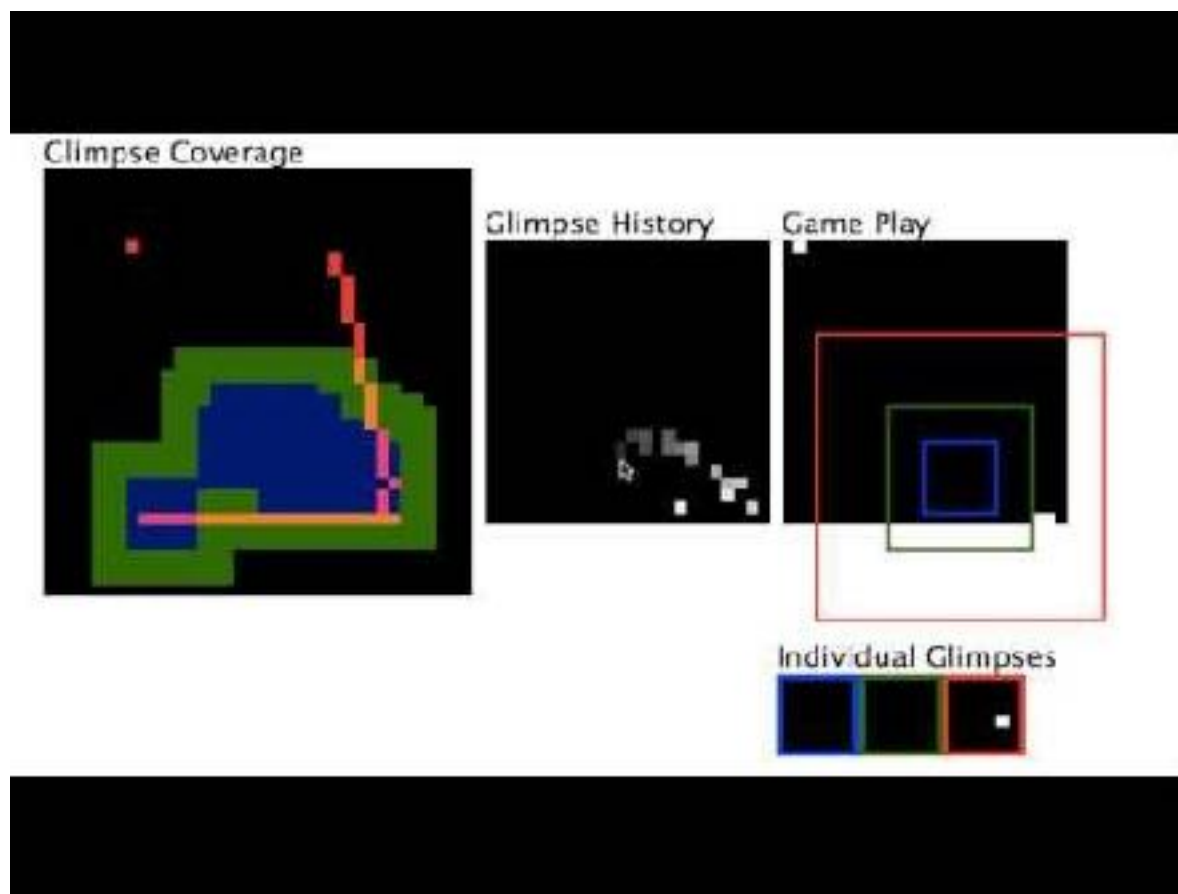
- Turns out this is very similar and in some cases equivalent to reinforcement learning using the REINFORCE learning rule [Williams, 1992]
- Intuition
 - We want to perform a set of actions leading to a reward
 - The reward is correct object classification
 - What actions do I need to take to get there
 - How do I summarize previous actions - RNN



Hard attention examples

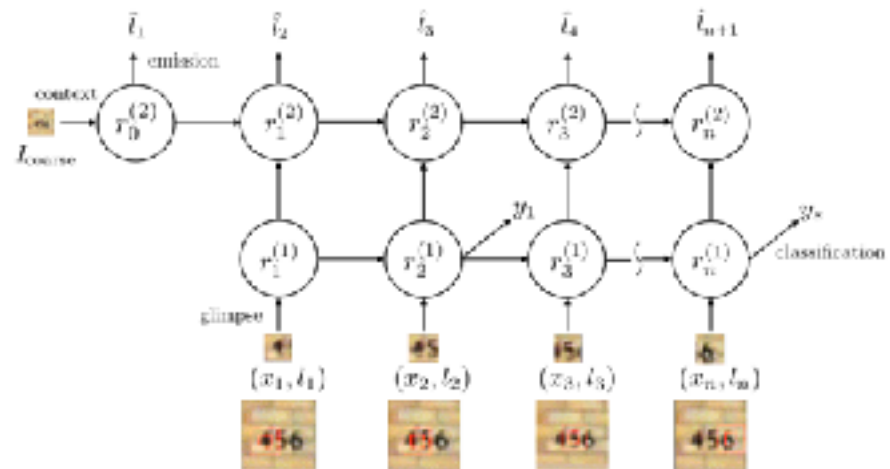


Some results



DRAM - a slight extension to RAM model

- Coarse image input initially
- Interestingly location output from top LSTM
- But the bottom one does classification
 - Why?
- This prevents the coarse image to be used for classification, thus shortcutting attention



Spatial Transformer networks

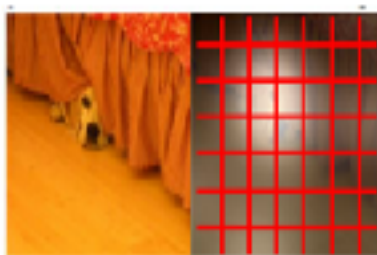


Some limitations of grid based attention

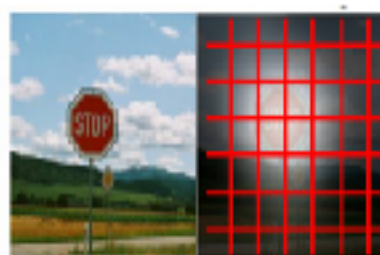
- Limited to a fixed grid analysis, glimpses solve this a bit but it's difficult to train, can we fixate on small parts of image but still have easy end-to-end training?



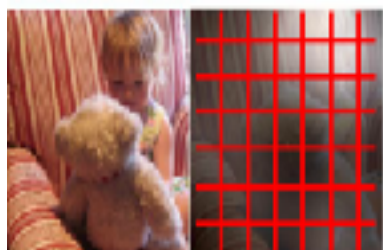
A woman is throwing a frisbee in a park.



A dog is standing on a hardwood floor.



A stop sign is on a road with a mountain in the background.



A little girl sitting on a bed with a teddy bear.

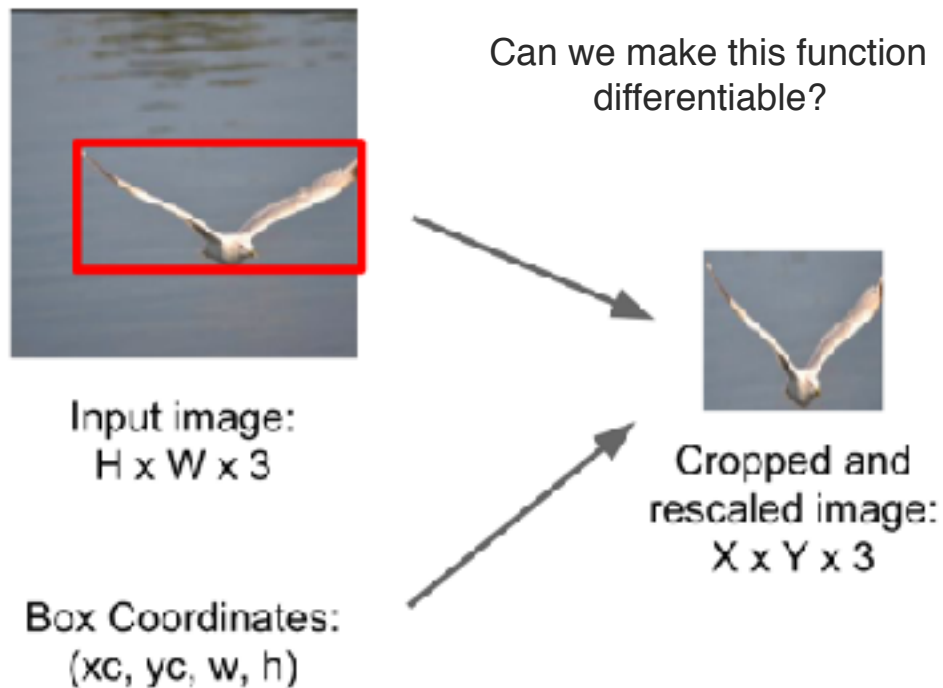


A group of people sitting on a boat in the water.



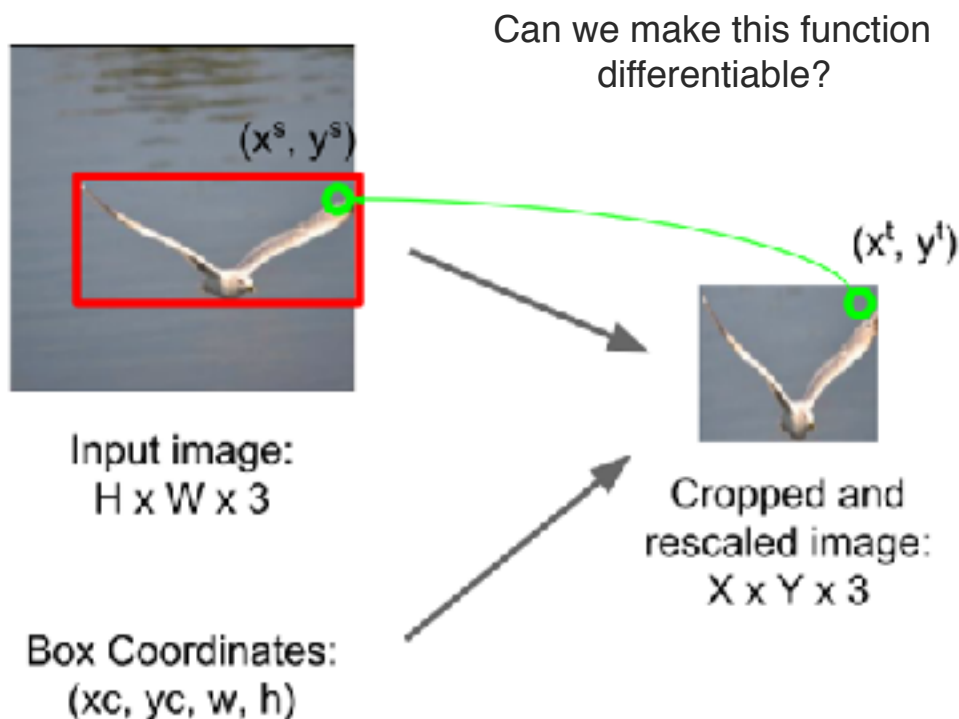
A giraffe standing in a forest with trees in the background.

Spatial Transformer Networks



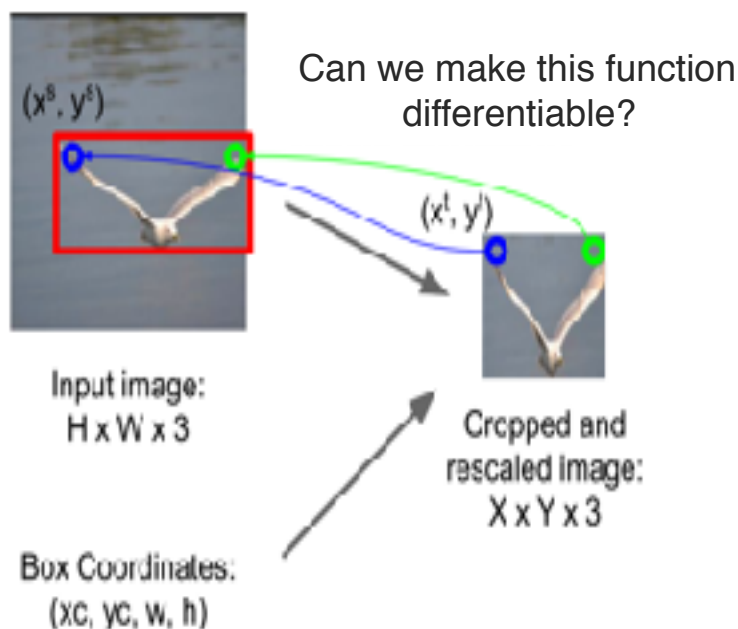
Spatial Transformer Networks

Idea: Function mapping pixel coordinates (x^t, y^t) of output to pixel coordinates (x^s, y^s) of input



$$\begin{pmatrix} x_i^s \\ y_i^s \end{pmatrix} = \begin{bmatrix} \theta_{1,1} & \theta_{1,2} & \theta_{1,3} \\ \theta_{2,1} & \theta_{2,2} & \theta_{2,3} \end{bmatrix} \begin{pmatrix} x_i^t \\ y_i^t \\ 1 \end{pmatrix}$$

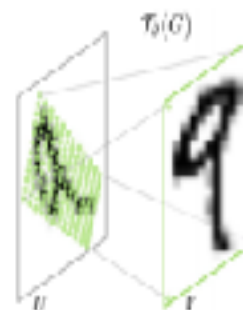
Spatial Transformer Networks



Idea: Function mapping pixel coordinates (x^t, y^t) of output to pixel coordinates (x^s, y^s) of input

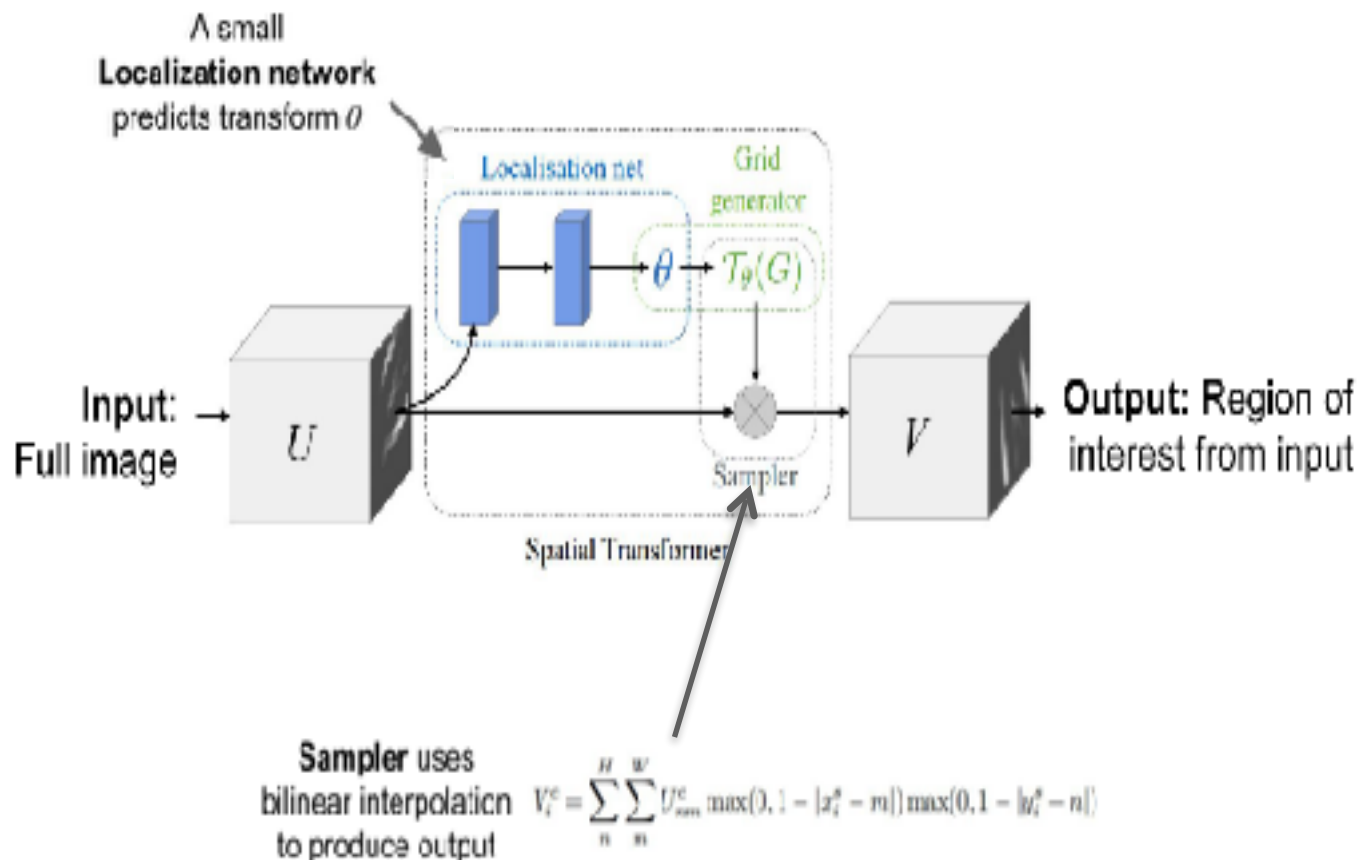
$$\begin{pmatrix} x_i^s \\ y_i^s \end{pmatrix} = \begin{bmatrix} \theta_{1,1} & \theta_{1,2} & \theta_{1,3} \\ \theta_{2,1} & \theta_{2,2} & \theta_{2,3} \end{bmatrix} \begin{pmatrix} x_i^t \\ y_i^t \\ 1 \end{pmatrix}$$

Network "attends" to input by predicting θ



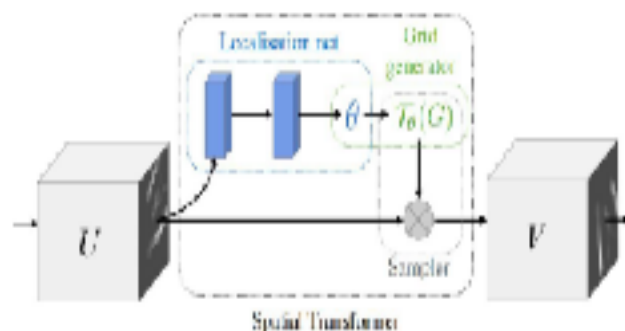
Repeat for all pixels in output to get a sampling grid

Spatial Transformer Networks

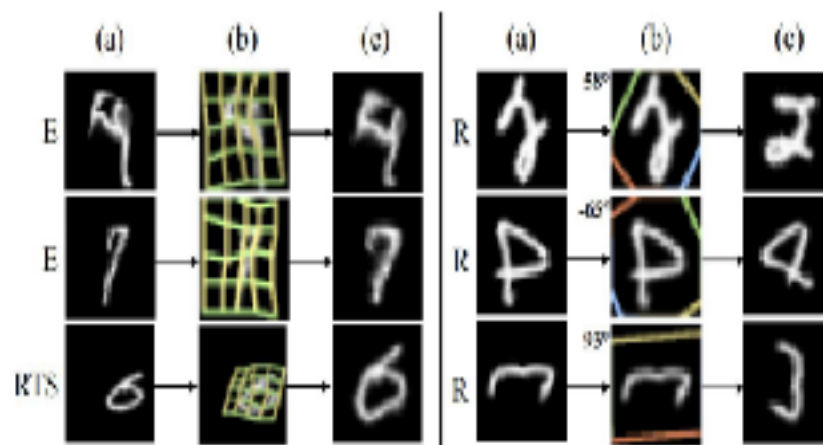


Spatial Transformer Networks

Differentiable "attention / transformation" module

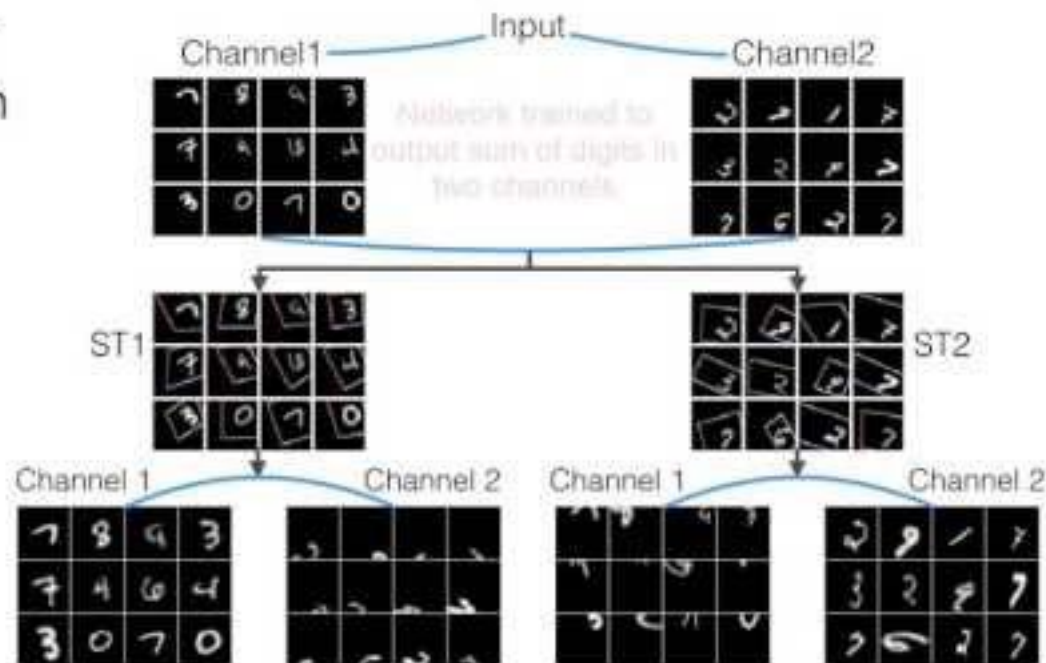


Insert spatial transformers into a classification network and it learns to attend and transform the input



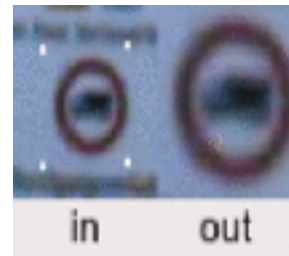
Spatial Transformer Networks

MNIST
Addition



Examples on real world data

- State-of-the-art results on traffic sign recognition



Code available http://torch.ch/blog/2015/09/07/spatial_transformers.html

Recap on Spatial Transformer Networks

- Differentiable so we can just use back-prop for training end-to-end
- Can use complex models for focusing on an image
 - Affine and Piece-Wise Affine
 - Perspective
 - Thin Plate Splines
- Can use to focus on certain parts of an image
- Seems to outperform soft and hard attention models on a number of visual tasks
- We can use it instead of (or in addition to) soft and hard attention for multi-modal tasks

Multi-modal alignment recap



Multimodal-alignment recap

- Explicit alignment - aligns two or more modalities (or views) as an actual task. The goal is to find correspondences between modalities
 - Dynamic Time Warping
 - Canonical Time Warping
 - Deep Canonical Time Warping
- Implicit alignment - uses internal latent alignment of modalities in order to better solve various problems
 - Attention models
 - Soft attention
 - Hard attention
 - Spatial transformer networks