

CONVOLUTIONAL NEURAL NETS FOR SPEECH

YU / HUGHES

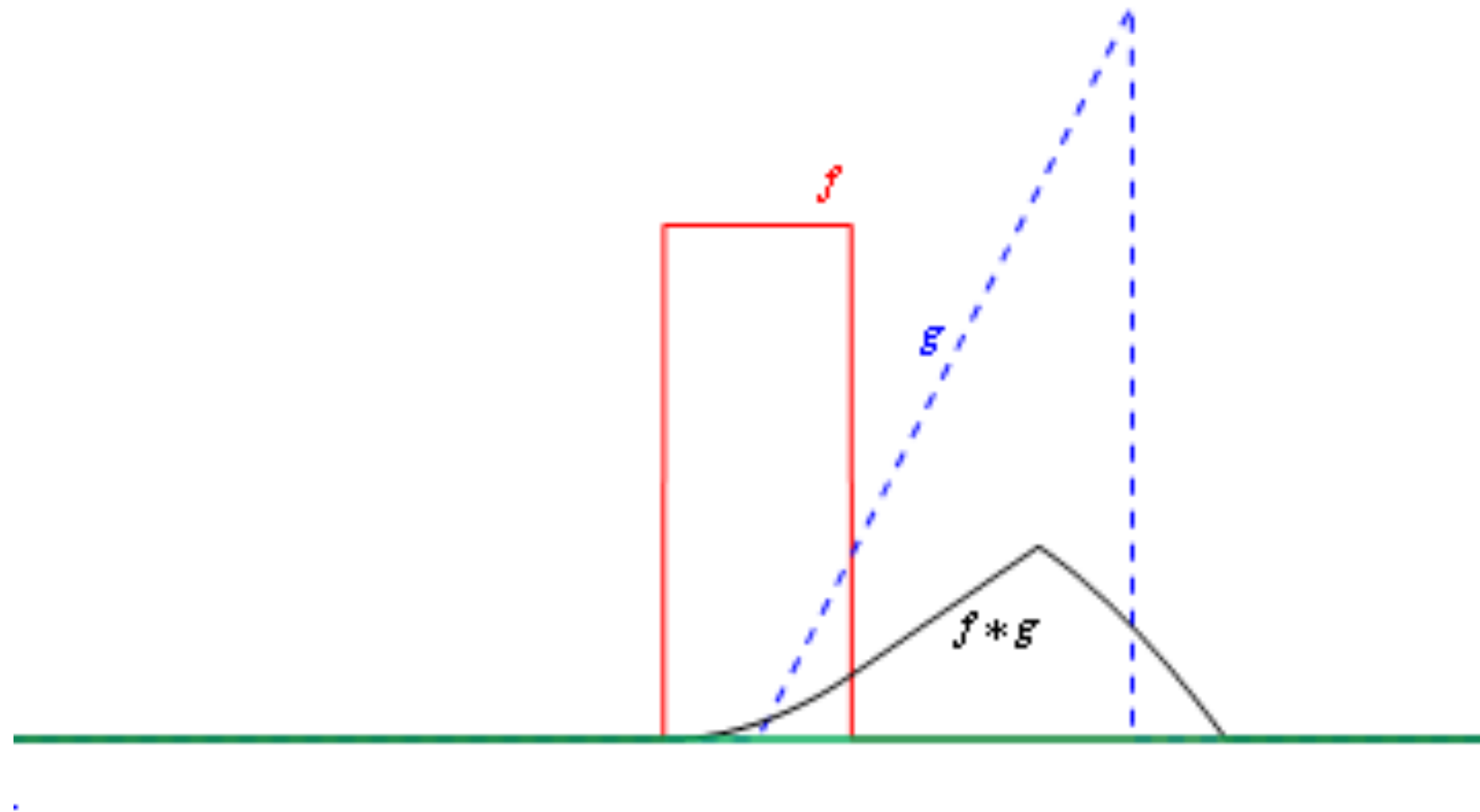
NOVEMBER 9, 2021

UMassAmherst

CONVOLUTION REVIEW

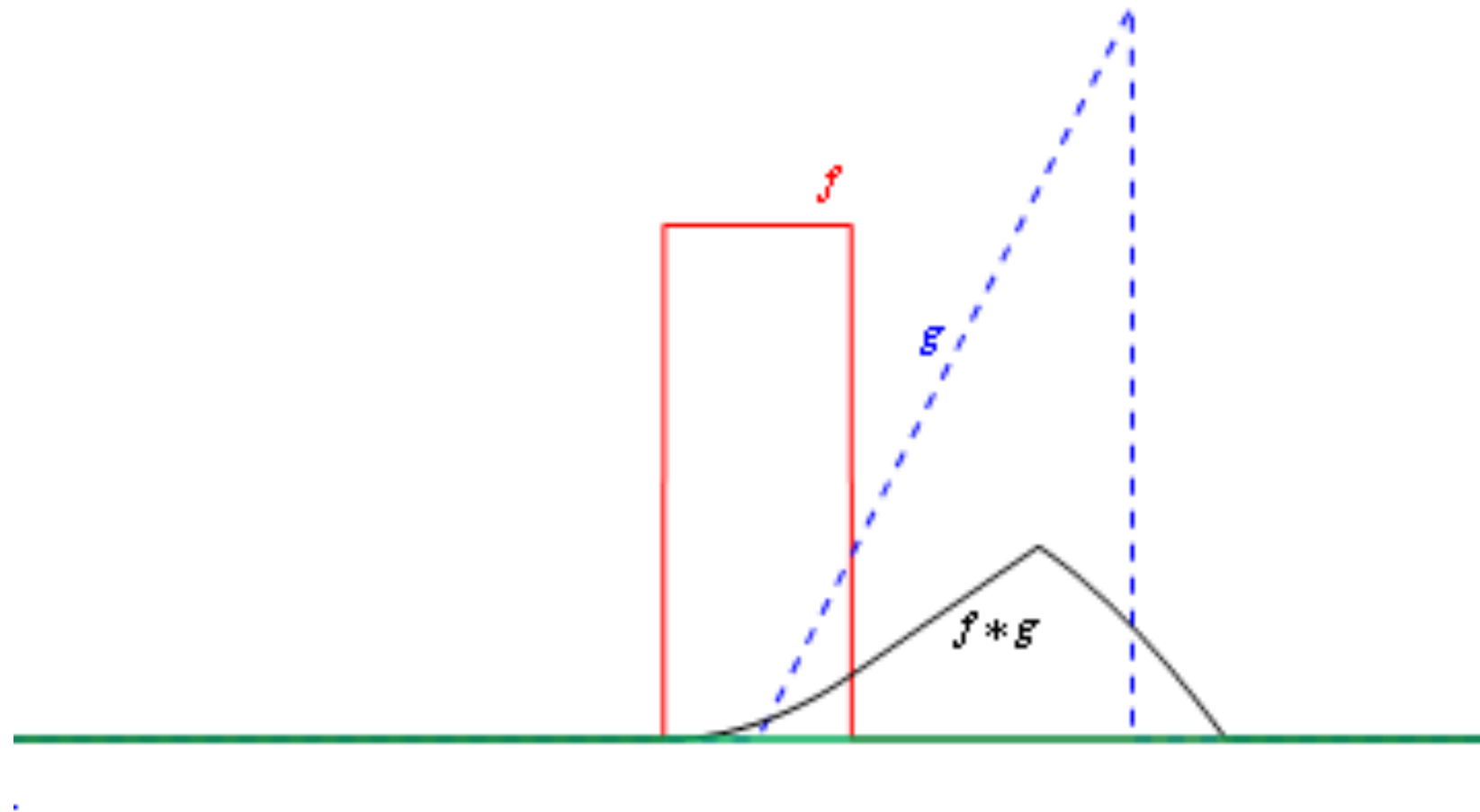
FLIP AND SHIFT! RECT AND RAMP

https://upload.wikimedia.org/wikipedia/commons/f/f8/Convolution_Animation_%28Boxcar_and_Ramp%29.gif



FLIP AND SHIFT! RECT AND RAMP

https://upload.wikimedia.org/wikipedia/commons/f/f8/Convolution_Animation_%28Boxcar_and_Ramp%29.gif

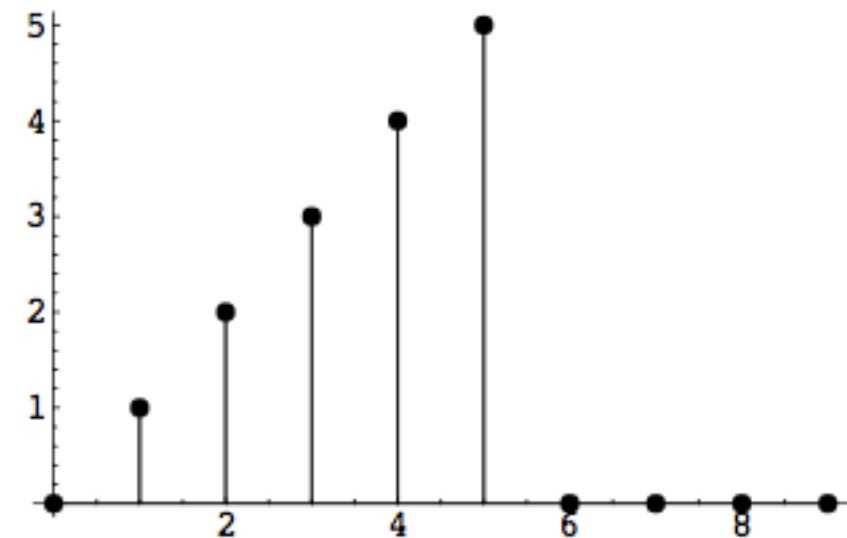


DIGITAL CONVOLUTION

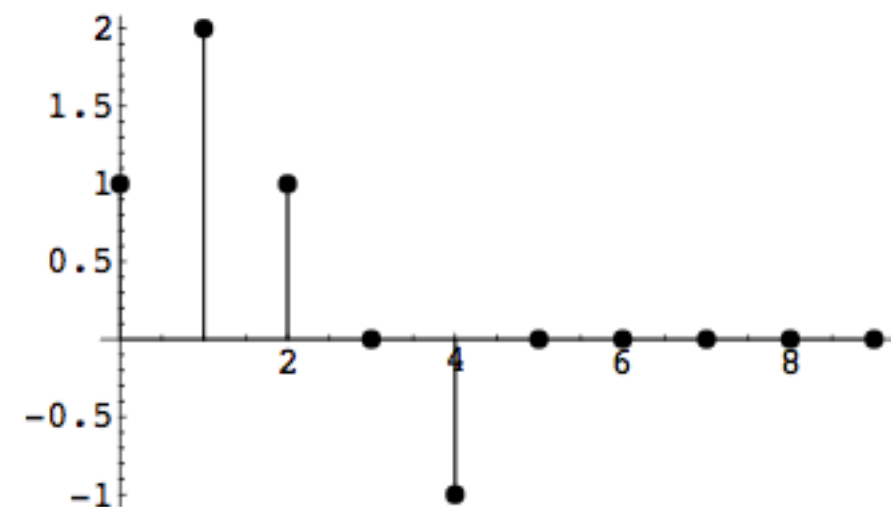
Shown at the right are two functions that we would like to process by convolution.

$f(n)$ has length $L = 6$ and $h(n)$ has length $K = 5$.

The functions have been extended by zero-padding to length $P = L + K - 1 = 10$.



Function $f(n)$.



Function $h(n)$.

DIGITAL CONVOLUTION

One Way to Think of Convolution

$$x(t) * h(t) = \int_{-\infty}^{\infty} x(\tau) h(t - \tau) d\tau$$

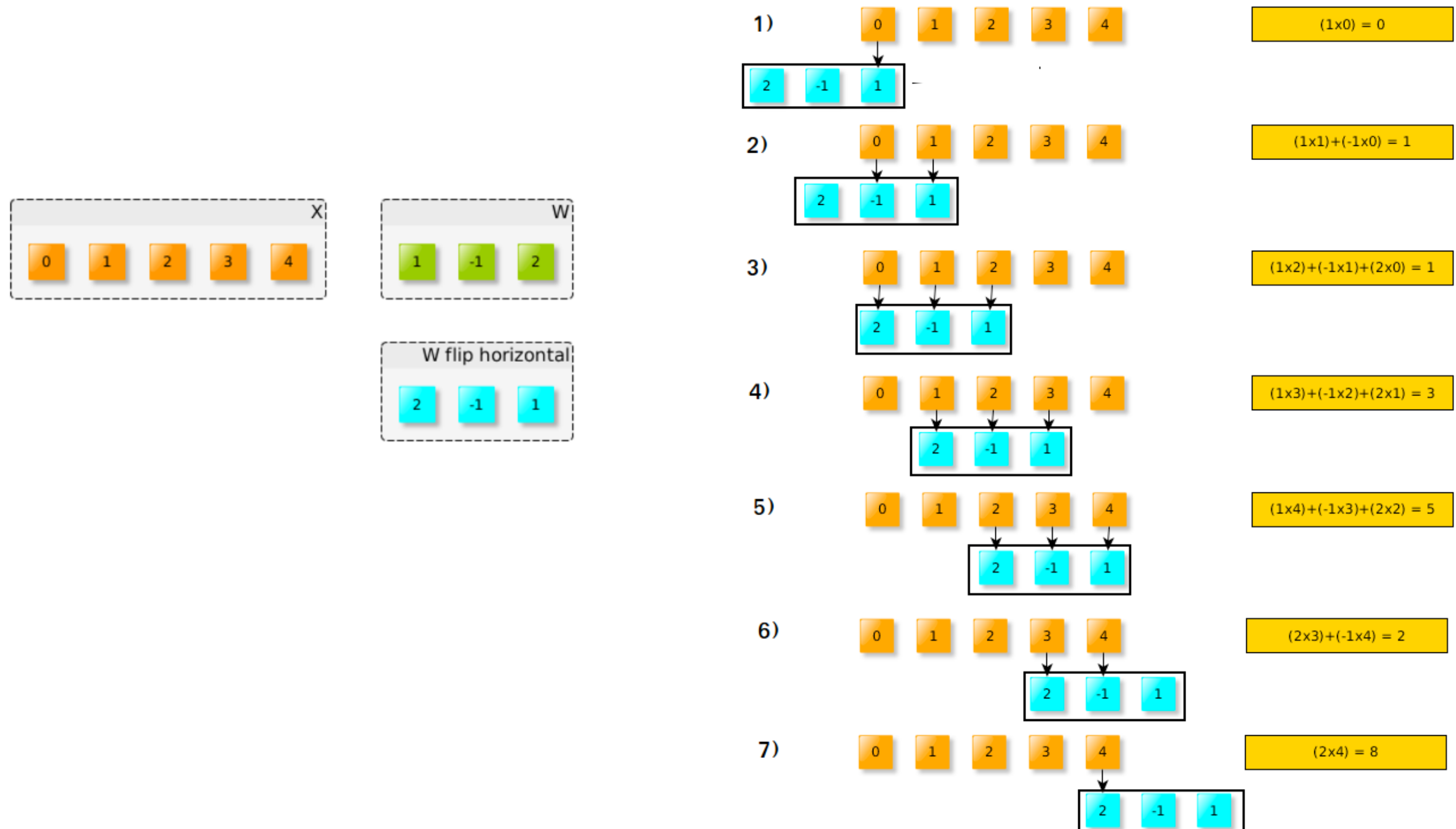
$$x[j] * h[j] = \sum_k x[k] \cdot h[j - k]$$

Think of it this way:

- Shift a copy of h to each position t (or discrete position k)
- Multiply by the value at that position $x(t)$ (or discrete sample $x[k]$)
- Add shifted, multiplied copies for all t (or discrete k)

FOR A WORKED OUT NUMERICAL EXAMPLE...

https://leonardoaraujosantos.gitbook.io/artificial-intelligence/machine_learning/deep_learning/convolution



FOR A WORKED OUT NUMERICAL EXAMPLE...

https://leonardoaraujosantos.gitbook.io/artificial-intelligence/machine_learning/deep_learning/convolution

```
In [1]: import numpy as np  
  
In [2]: x = np.array([0,1,2,3,4])  
  
In [3]: w = np.array([1,-1,2])  
  
In [4]: res = np.convolve(x,w)  
  
In [5]: print res  
[0 1 1 3 5 2 8]
```

REVIEW QUESTION: DIGITAL CONVOLUTION

https://leonardoaraujosantos.gitbook.io/artificial-intelligence/machine_learning/deep_learning/convolution

Work through the convolution example at
https://leonardoaraujosantos.gitbook.io/artificial-intelligence/machine_learning/deep_learning/convolution#doing-by-hand

NEURAL NETS IN SPEECH RECOGNITION ARCHITECTURE

FUNDAMENTAL EQUATION OF SPEECH RECOGNITION

Most likely
word sequence

Word
sequence

The diagram features a central rectangular box with a black border. Inside the box is the equation $\mathbf{W}^* = \arg \max_{\mathbf{W}} P(\mathbf{W} | \mathbf{X})$. Three red arrows point towards the equation: one from the text 'Most likely word sequence' above the left side, one from the text 'Word sequence' above the right side, and one from the text 'Observations (Acoustic feature vectors)' below the right side.

$$\mathbf{W}^* = \arg \max_{\mathbf{W}} P(\mathbf{W} | \mathbf{X})$$

Observations
(Acoustic feature vectors)

FUNDAMENTAL EQUATION OF SPEECH RECOG

$$\mathbf{W}^* = \arg \max_{\mathbf{W}} P(\mathbf{W} | \mathbf{X})$$

$$P(\mathbf{W} | \mathbf{X}) = \frac{P(\mathbf{X} | \mathbf{W})P(\mathbf{W})}{P(\mathbf{X})}$$
$$\propto P(\mathbf{X} | \mathbf{W})P(\mathbf{W})$$

$$\mathbf{W}^* = \arg \max_{\mathbf{W}} P(\mathbf{X} | \mathbf{W}) P(\mathbf{W})$$

FUNDAMENTAL EQUATION OF SPEECH RECOG

$$\mathbf{W}^* = \arg \max_{\mathbf{W}} P(\mathbf{W} | \mathbf{X})$$

$$P(\mathbf{W} | \mathbf{X}) = \frac{P(\mathbf{X} | \mathbf{W})P(\mathbf{W})}{P(\mathbf{X})}$$
$$\propto P(\mathbf{X} | \mathbf{W})P(\mathbf{W})$$

$$\mathbf{W}^* = \arg \max_{\mathbf{W}} P(\mathbf{X} | \mathbf{W}) P(\mathbf{W})$$

Acoustic
Model

FUNDAMENTAL EQUATION OF SPEECH RECOG

$$\mathbf{W}^* = \arg \max_{\mathbf{W}} P(\mathbf{W} | \mathbf{X})$$

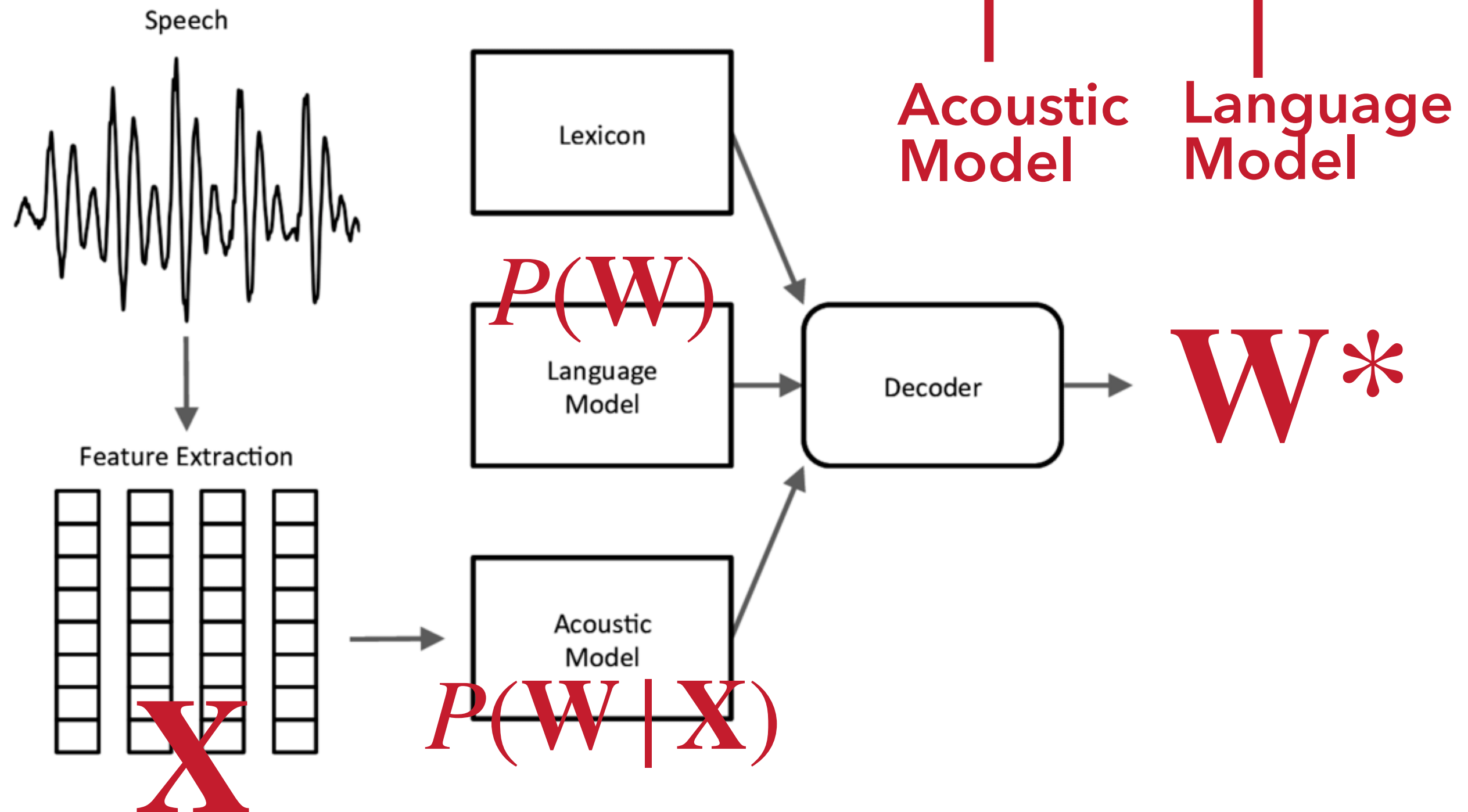
$$P(\mathbf{W} | \mathbf{X}) = \frac{P(\mathbf{X} | \mathbf{W})P(\mathbf{W})}{P(\mathbf{X})}$$
$$\propto P(\mathbf{X} | \mathbf{W})P(\mathbf{W})$$

$$\mathbf{W}^* = \arg \max_{\mathbf{W}} P(\mathbf{X} | \mathbf{W}) P(\mathbf{W})$$

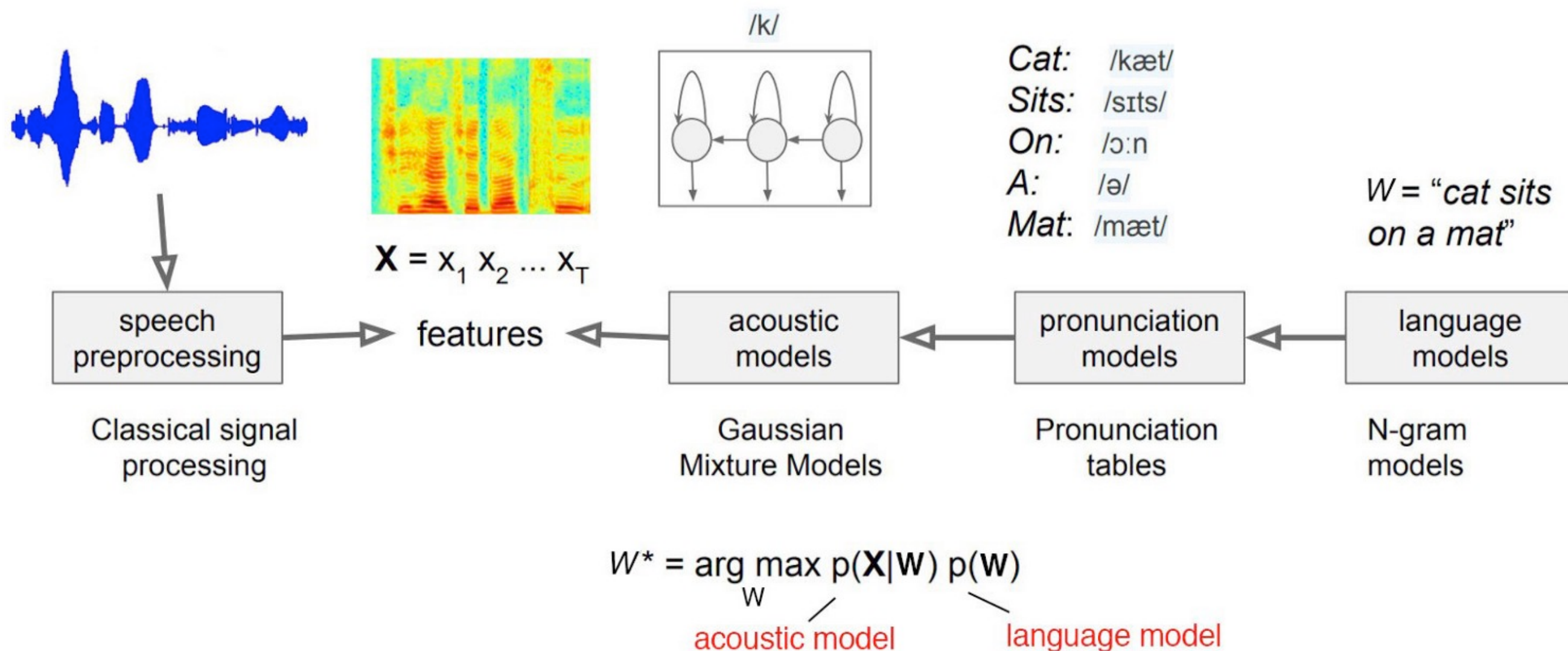
Acoustic
Model

Language
Model

$$\mathbf{W}^* = \arg \max_{\mathbf{W}} \boxed{P(\mathbf{X} | \mathbf{W})} \boxed{P(\mathbf{W})} \quad 13$$



CLASSICAL SPEECH RECOGNITION



Pronunciation model = Lexicon

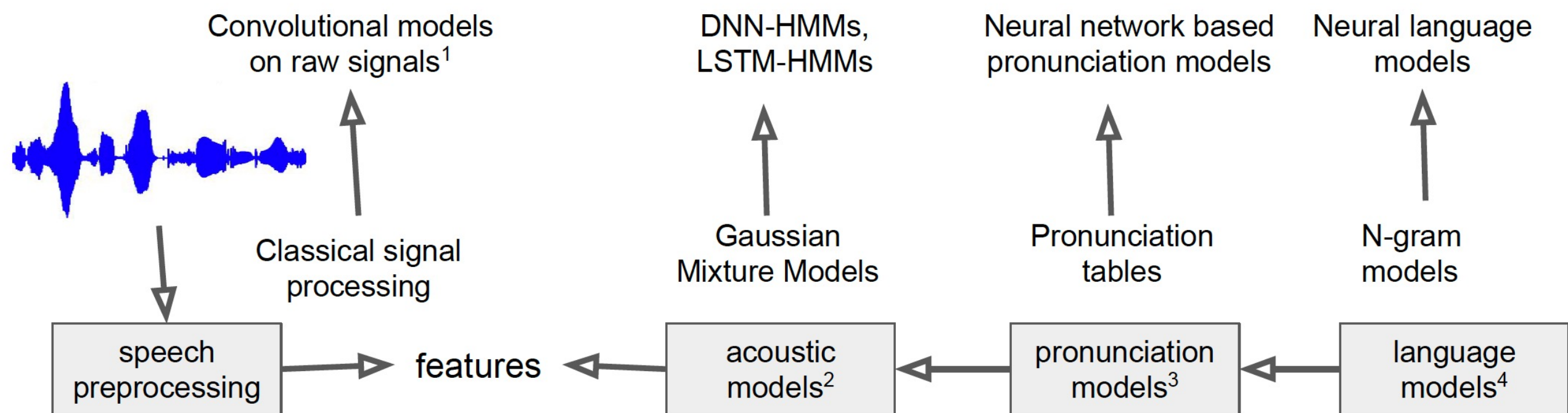
<https://jonathan-hui.medium.com/speech-recognition-gmm-hmm-8bb5eff8b196>

<https://web.stanford.edu/class/archive/cs/cs224n/cs224n.1174/lectures/cs224n-2017-lecture12.pdf>

NEURAL NETWORK “INVASION”

Speech Recognition -- the neural network invasion

- Each of the components seems to be better off with a neural network



1. Jaitly, Navdeep, and Geoffrey Hinton. "Learning a better representation of speech soundwaves using restricted boltzmann machines." *Acoustics, Speech and Signal Processing (ICASSP), 2011 IEEE International Conference on*. IEEE, 2011.

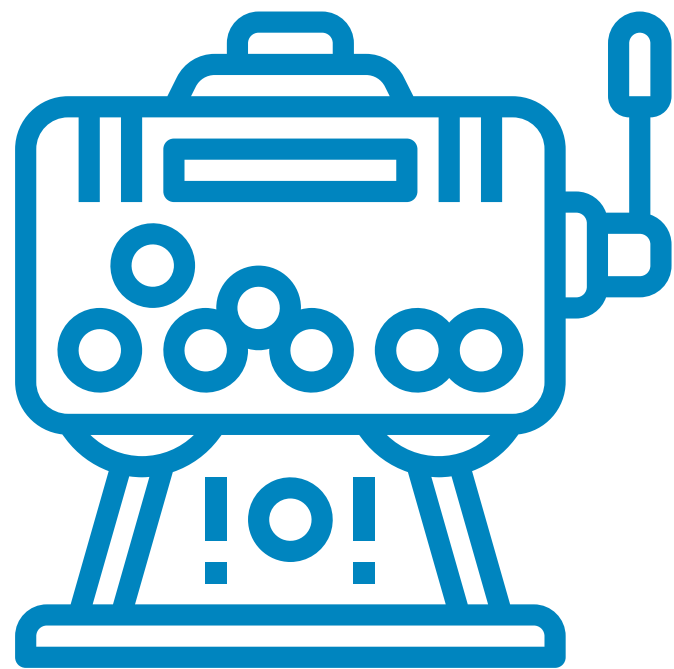
2. Hinton, Geoffrey, et al. "Deep neural networks for acoustic modeling in speech recognition: The shared views of four research groups." *IEEE Signal Processing Magazine* 29.6 (2012): 82-97.

3. Rao, Kanishka, et al. "Grapheme-to-phoneme conversion using long short-term memory recurrent neural networks." *Acoustics, Speech and Signal Processing (ICASSP), 2015 IEEE International Conference on*. IEEE, 2015.

4. Mikolov, Tomas, et al. "Recurrent neural network based language model." *Interspeech*. Vol. 2. 2010.

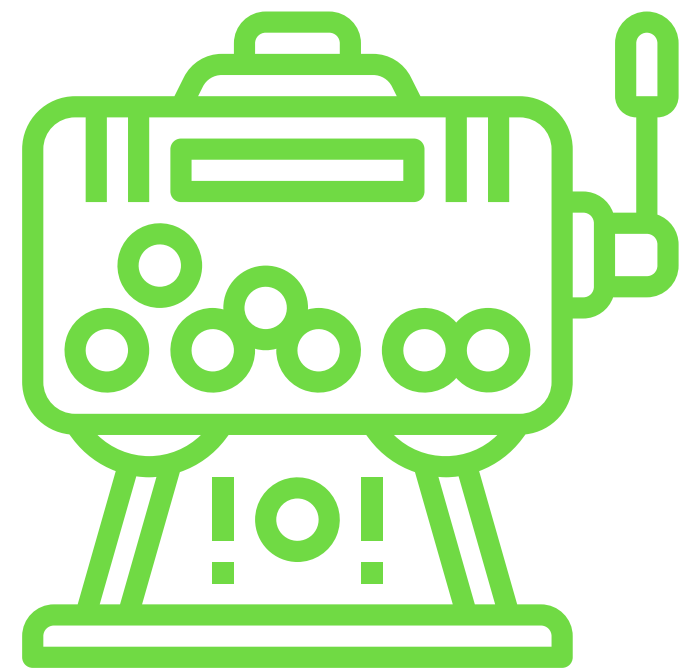
OBSERVATION PROBABILITIES

Machine m_1



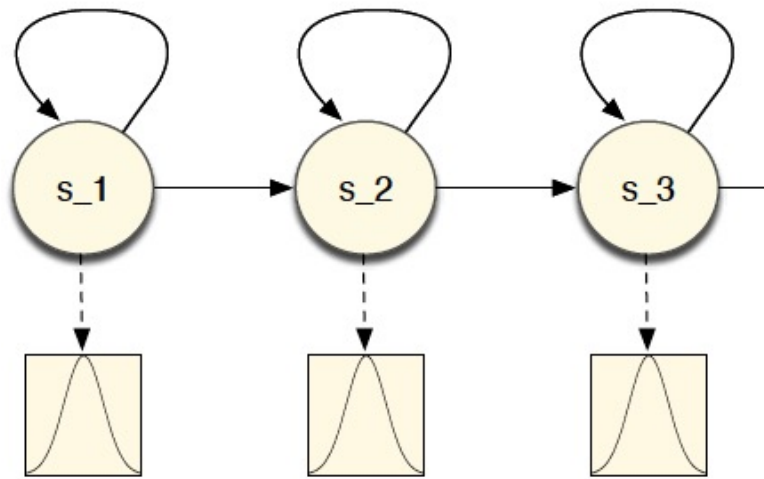
60% grapefruit
40% apple

Machine m_2

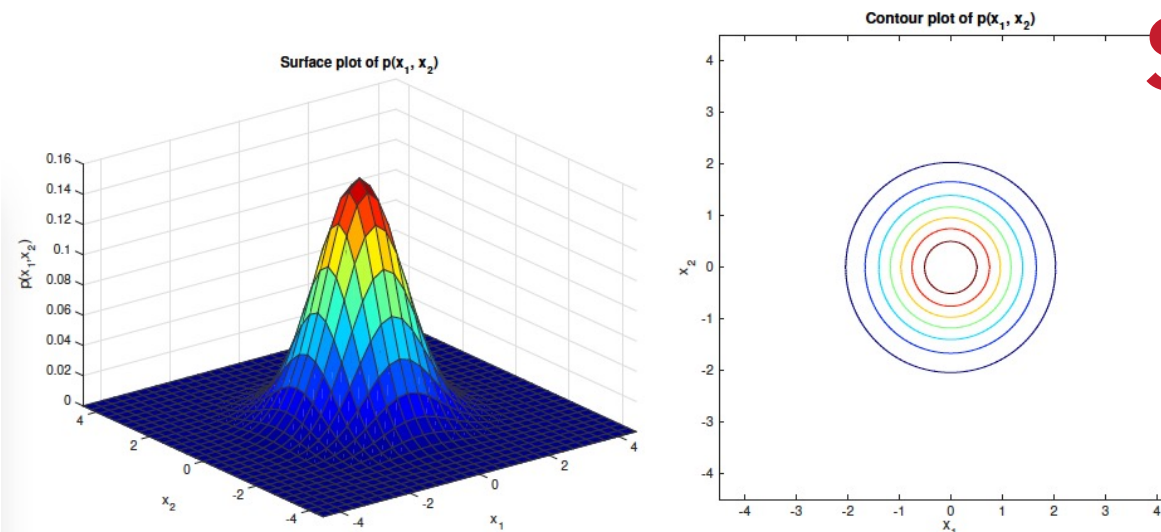


10% grapefruit
90% apple

GAUSSIAN MODEL OF OBSERVATION PROBS



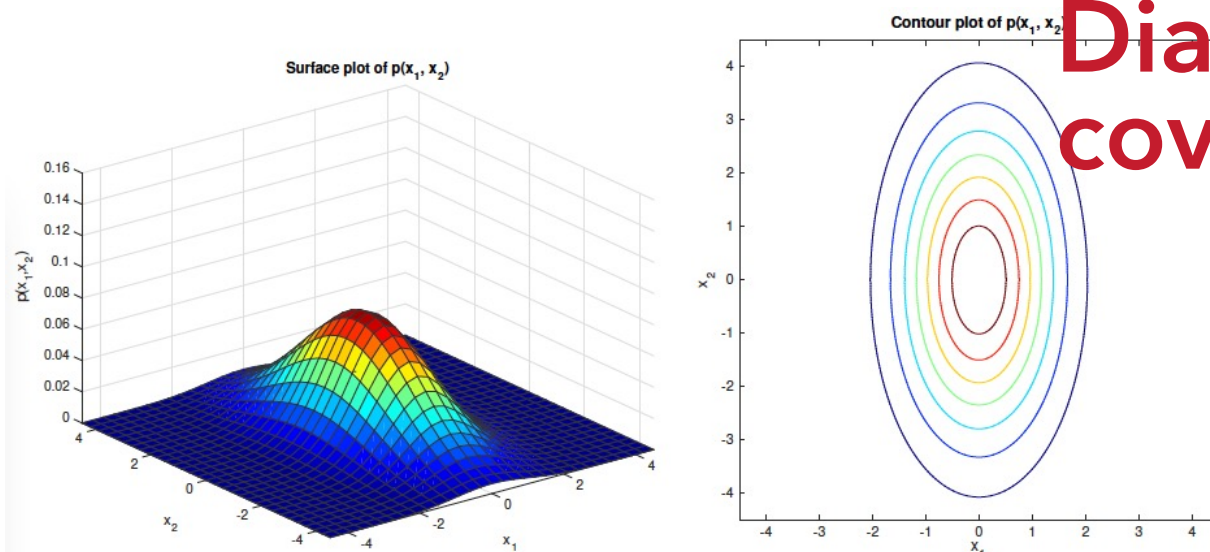
/s/



Spherical

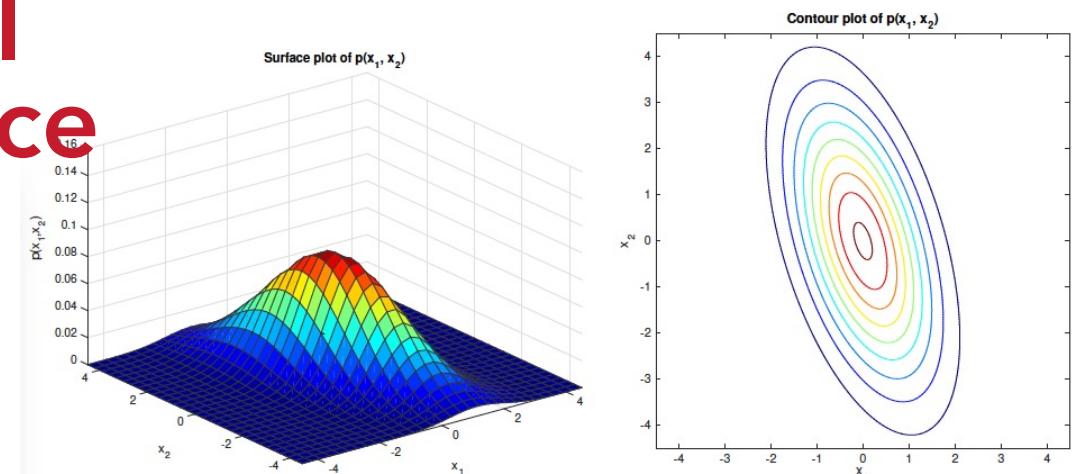
$$\mu = \begin{pmatrix} 0 \\ 0 \end{pmatrix} \quad \Sigma = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} \quad \rho_{12} = 0$$

Full
covariance



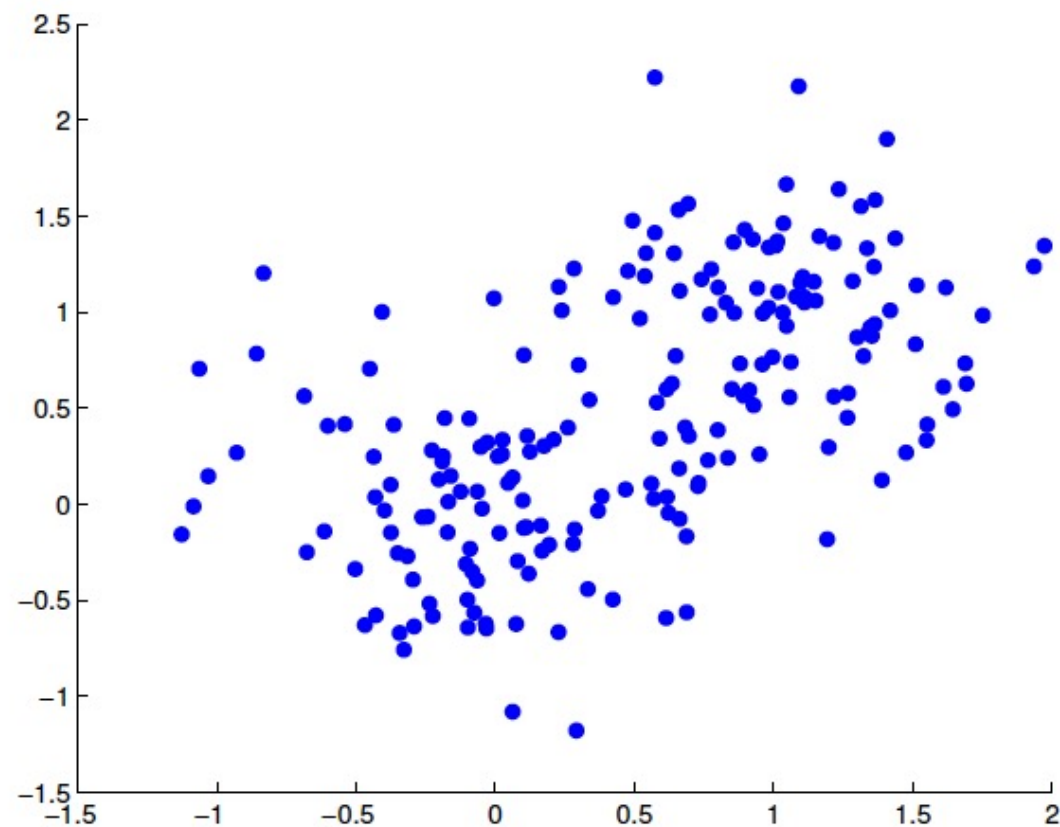
Diagonal
covariance

$$\mu = \begin{pmatrix} 0 \\ 0 \end{pmatrix} \quad \Sigma = \begin{pmatrix} 1 & 0 \\ 0 & 4 \end{pmatrix} \quad \rho_{12} = 0$$



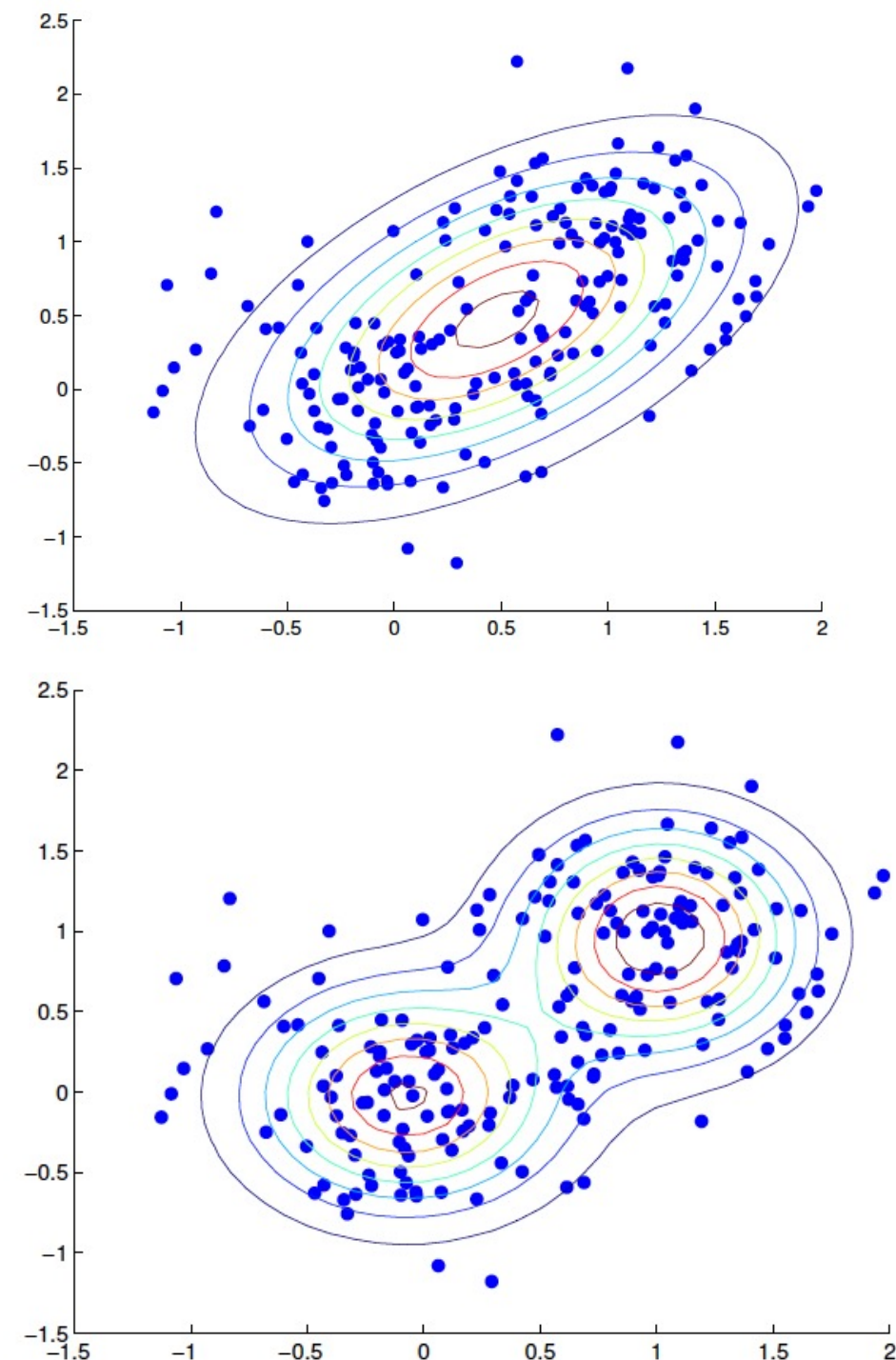
$$\mu = \begin{pmatrix} 0 \\ 0 \end{pmatrix} \quad \Sigma = \begin{pmatrix} 1 & -1 \\ -1 & 4 \end{pmatrix} \quad \rho_{12} = -0.5$$

GAUSSIANS TO GAUSSIAN MIXTURE MODELS



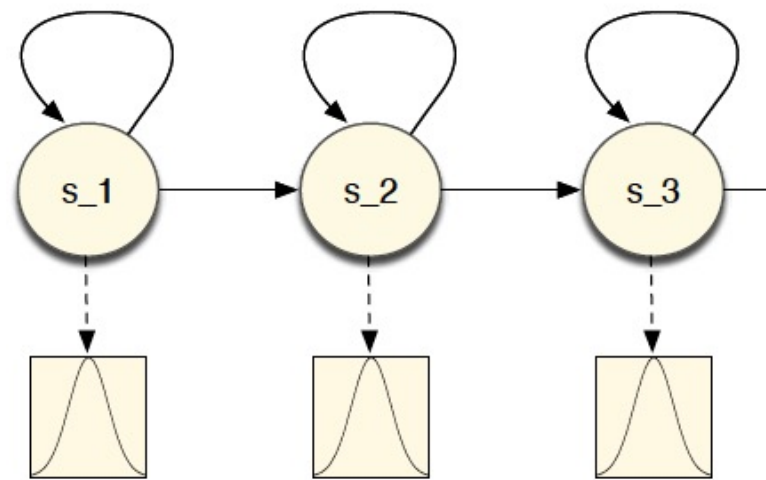
$$\mu_1 = (0, 0)^T \quad \mu_2 = (1, 1)^T$$

$$\Sigma_1 = \Sigma_2 = 0.2I$$

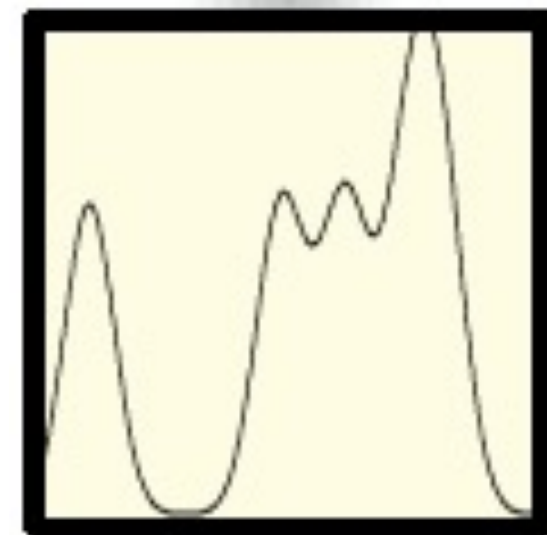
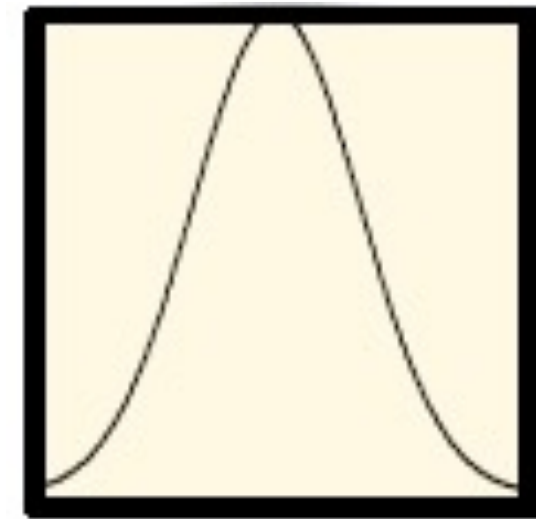


Fitted with a two component GMM using EM

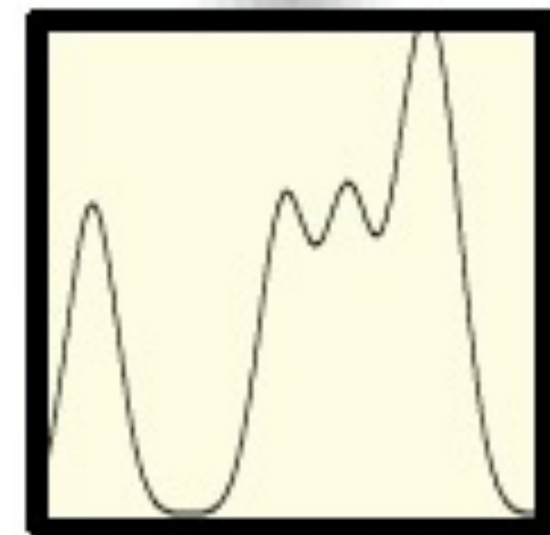
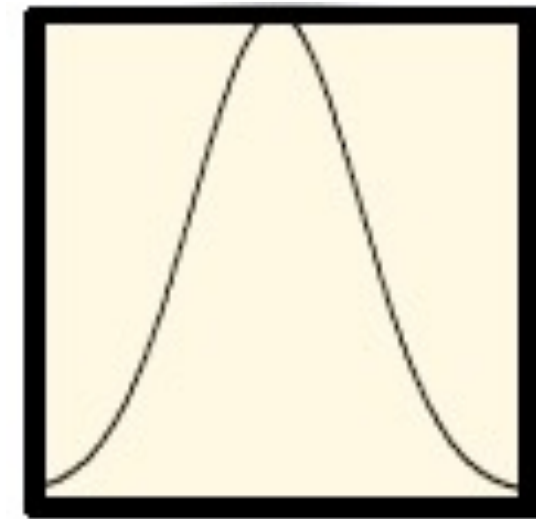
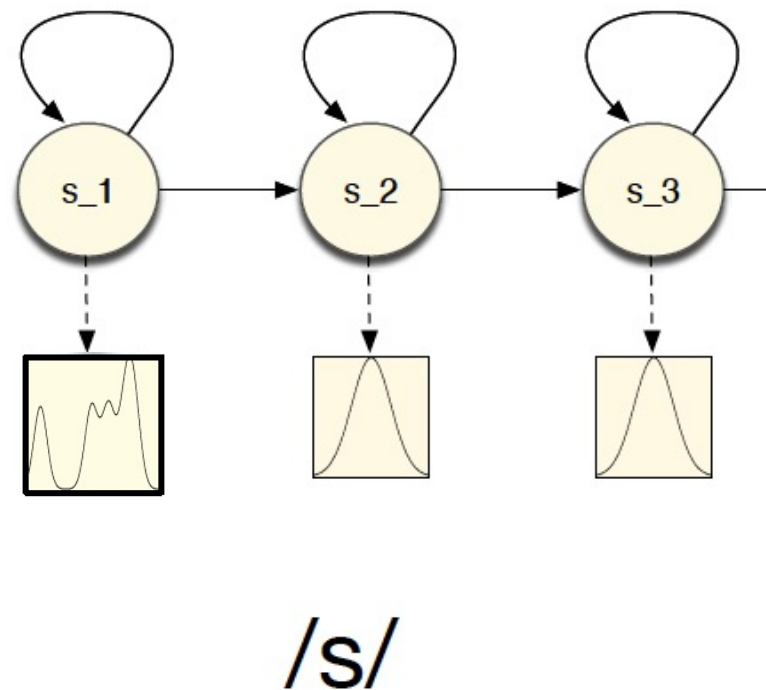
GAUSSIANS TO GAUSSIAN MIXTURE MODELS



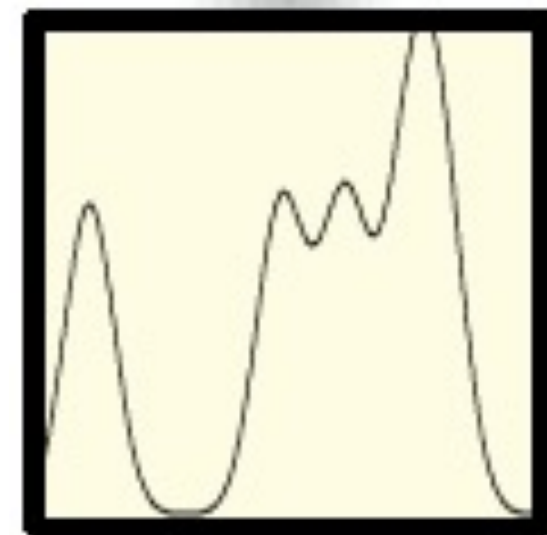
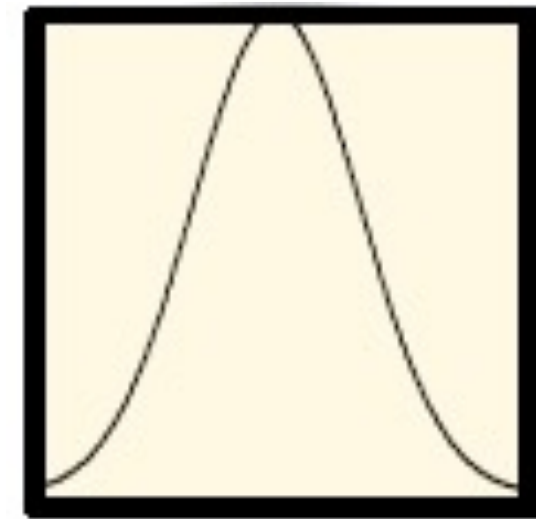
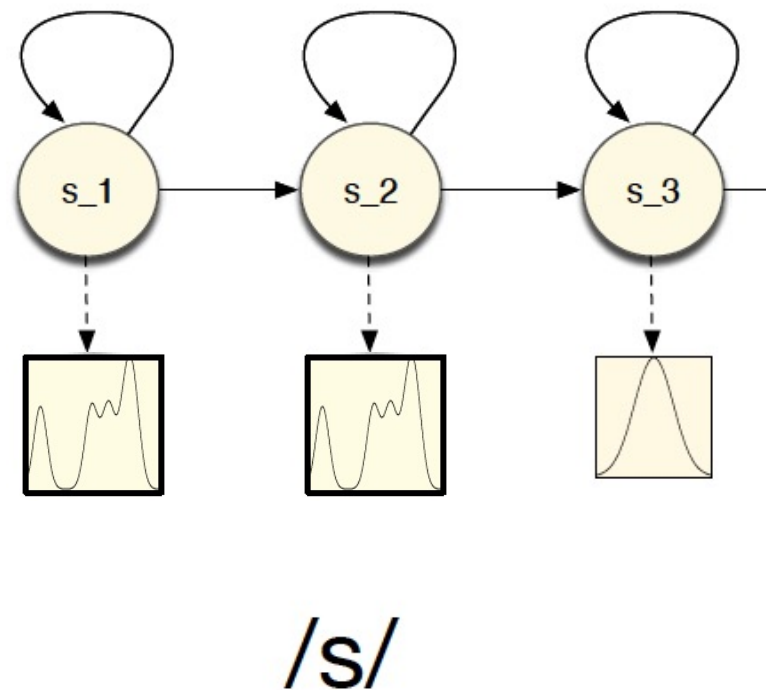
/s/



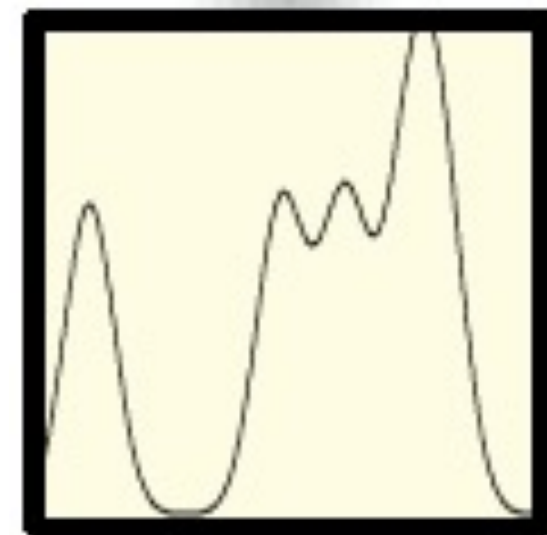
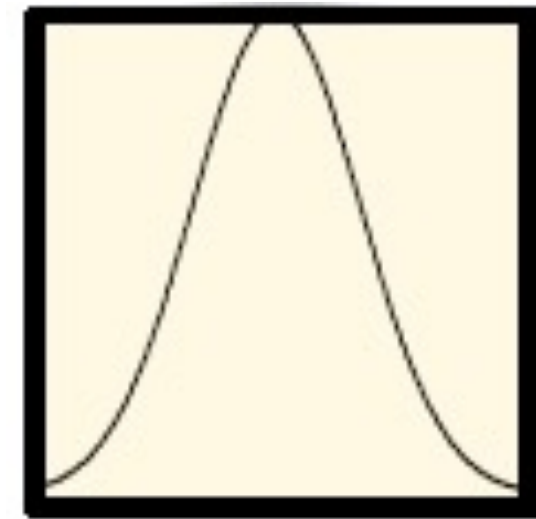
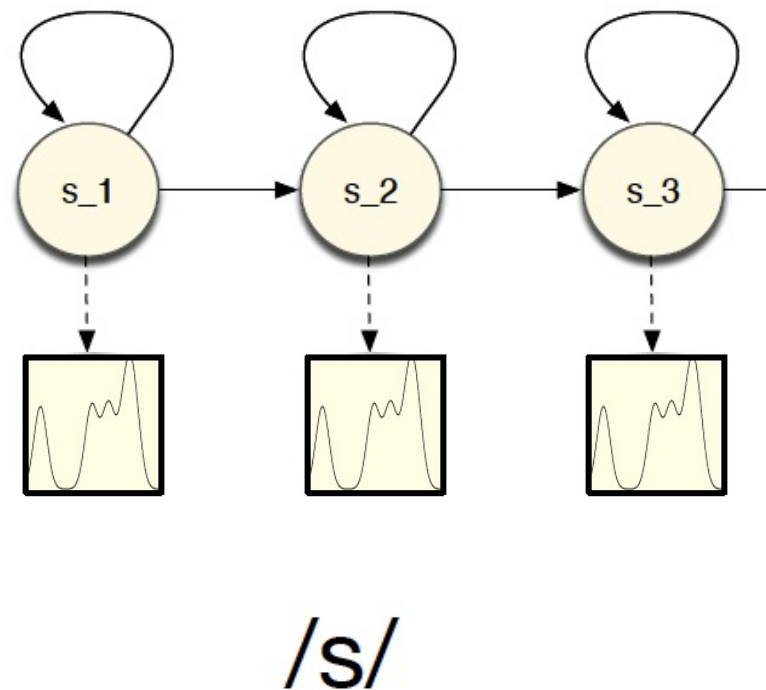
GAUSSIANS TO GAUSSIAN MIXTURE MODELS



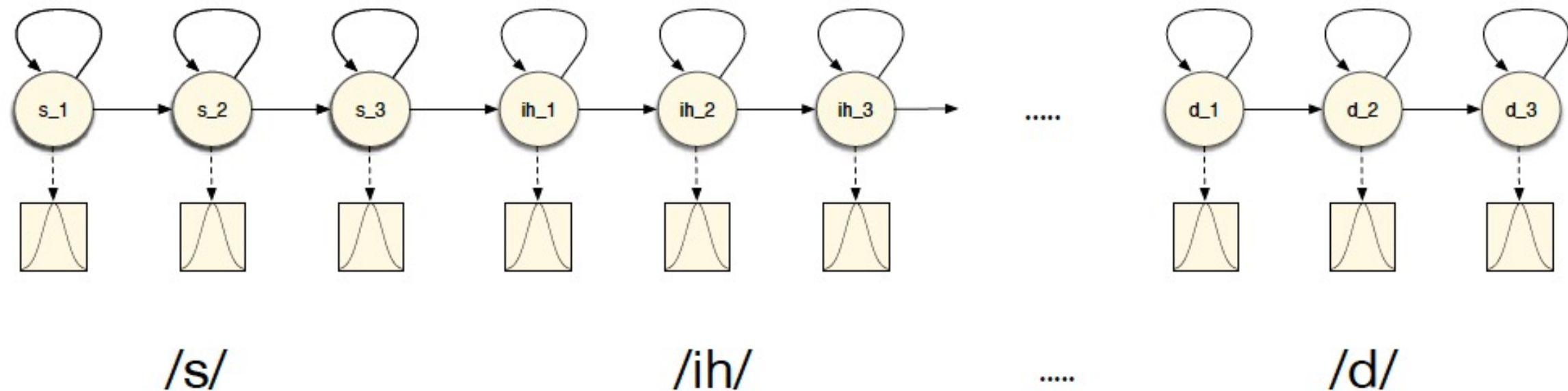
GAUSSIANS TO GAUSSIAN MIXTURE MODELS



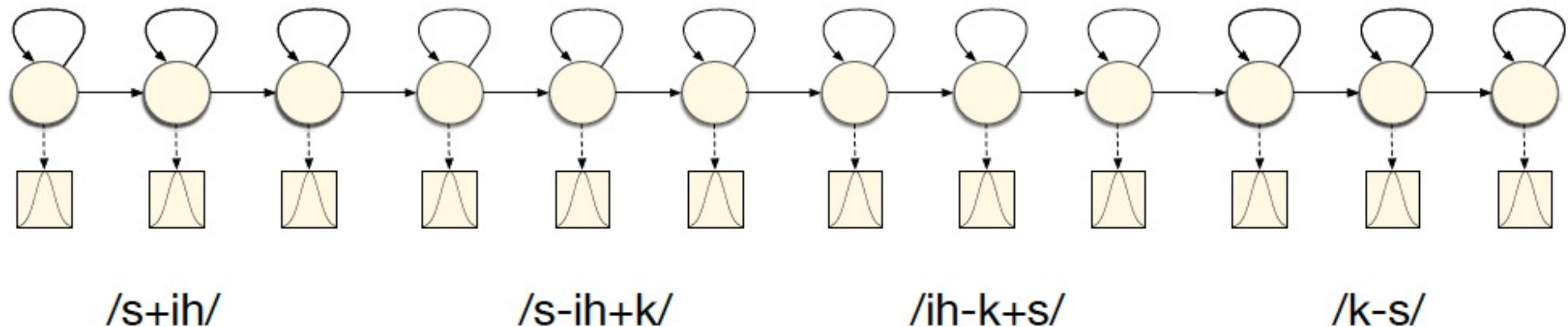
GAUSSIANS TO GAUSSIAN MIXTURE MODELS



CONTEXT DEPENDENT PHONE MODELS

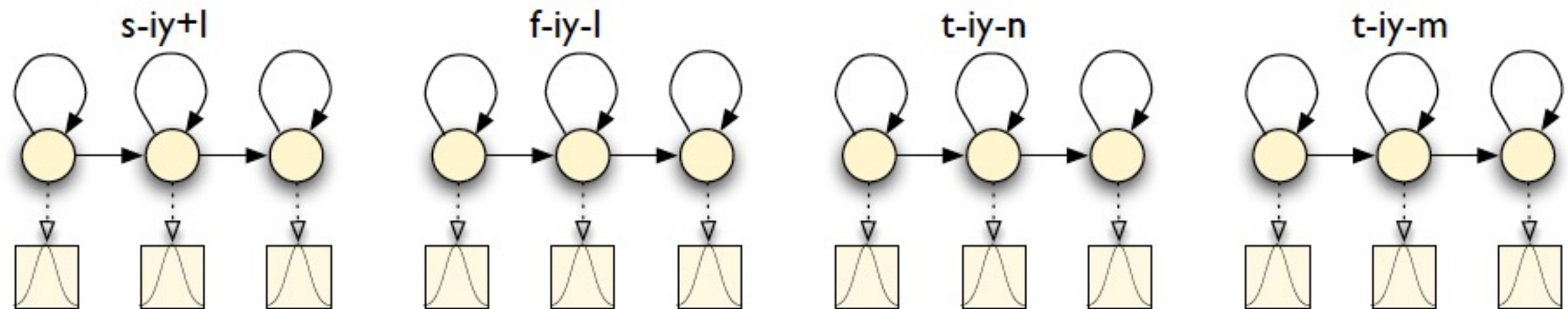


“six quid” **Triphones: left -x - right**

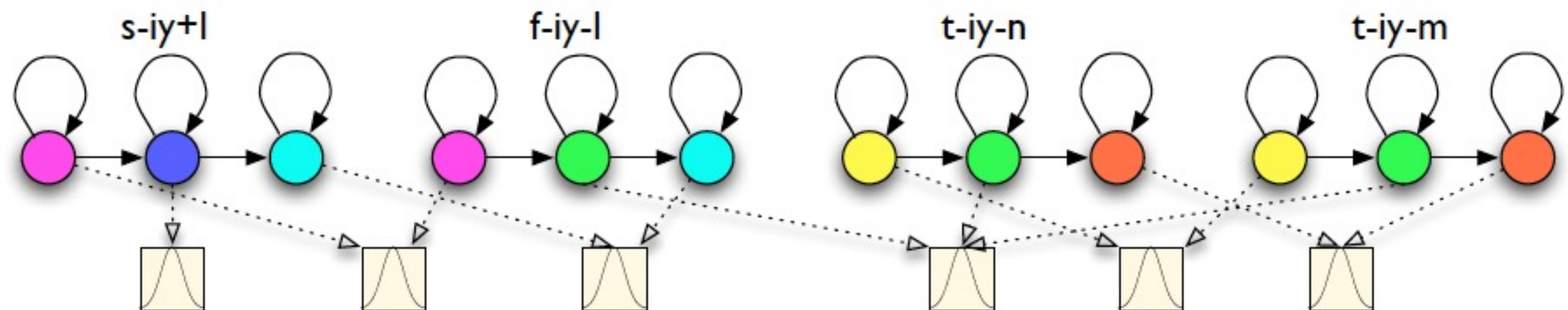


“six”

TOO MANY PARAMETERS: SHARE STATES

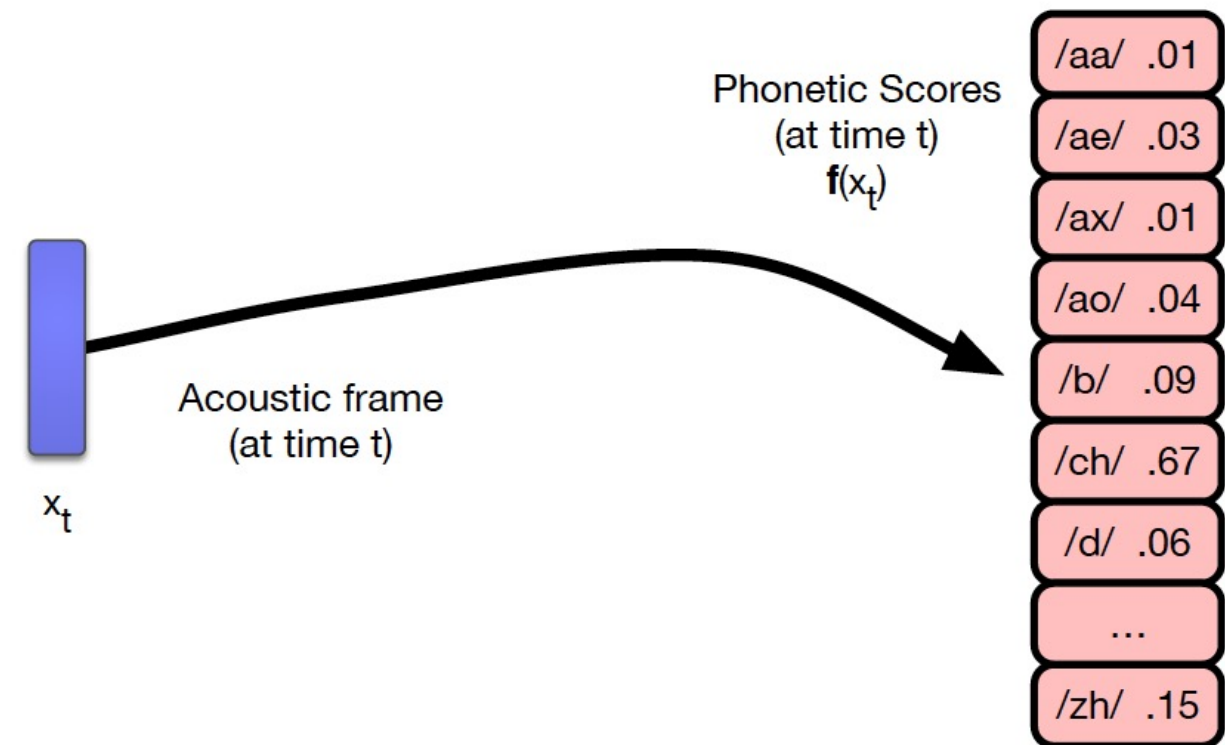
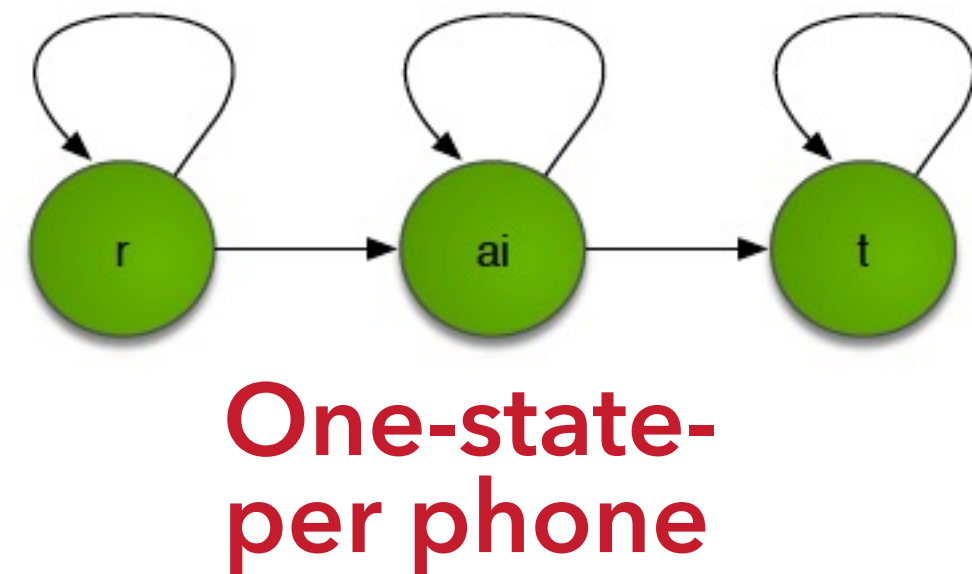


Simple triphones (no sharing)

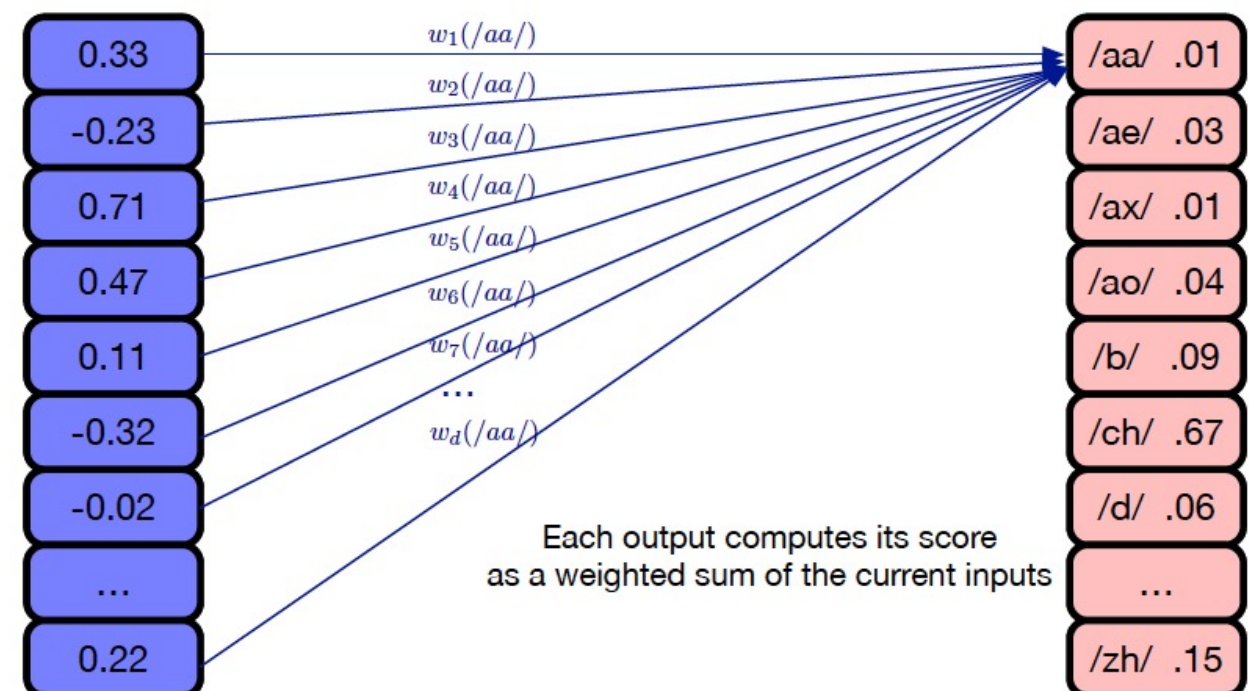


State-clustered triphones (state sharing)

FROM GMMS TO NEURAL NETS

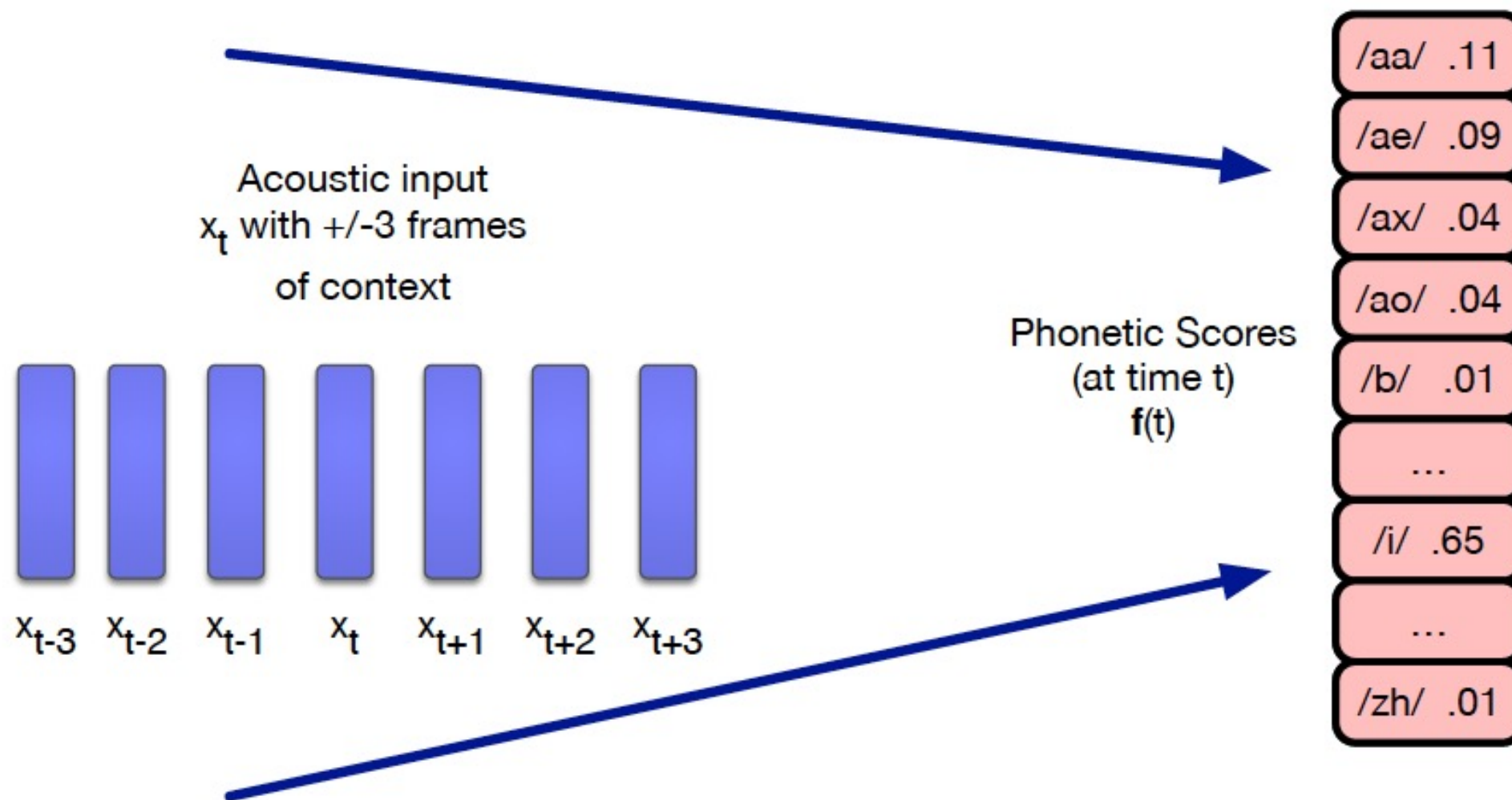


Compute the phonetic scores using a single layer neural network (linear regression!)

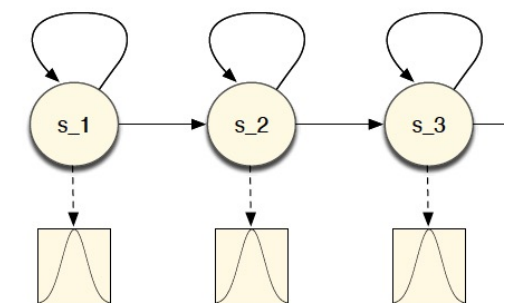


MODELING ACOUSTIC CONTEXT WITH NNS

Use multiple frames of acoustic context



Correlated features! No more conditional independence, cf.



/s/

NNS ALLOW UNCORRELATED FEATURES

- ▶ Can't use filter bank features because strongly correlated with one another, requiring full covariance and diagonal covariance Gaussians for GMMs
- ▶ NNs don't require uncorrelated feature vectors
 - ▶ Can use filter bank vectors
 - ▶ Can use multiple frames to model context-dependency

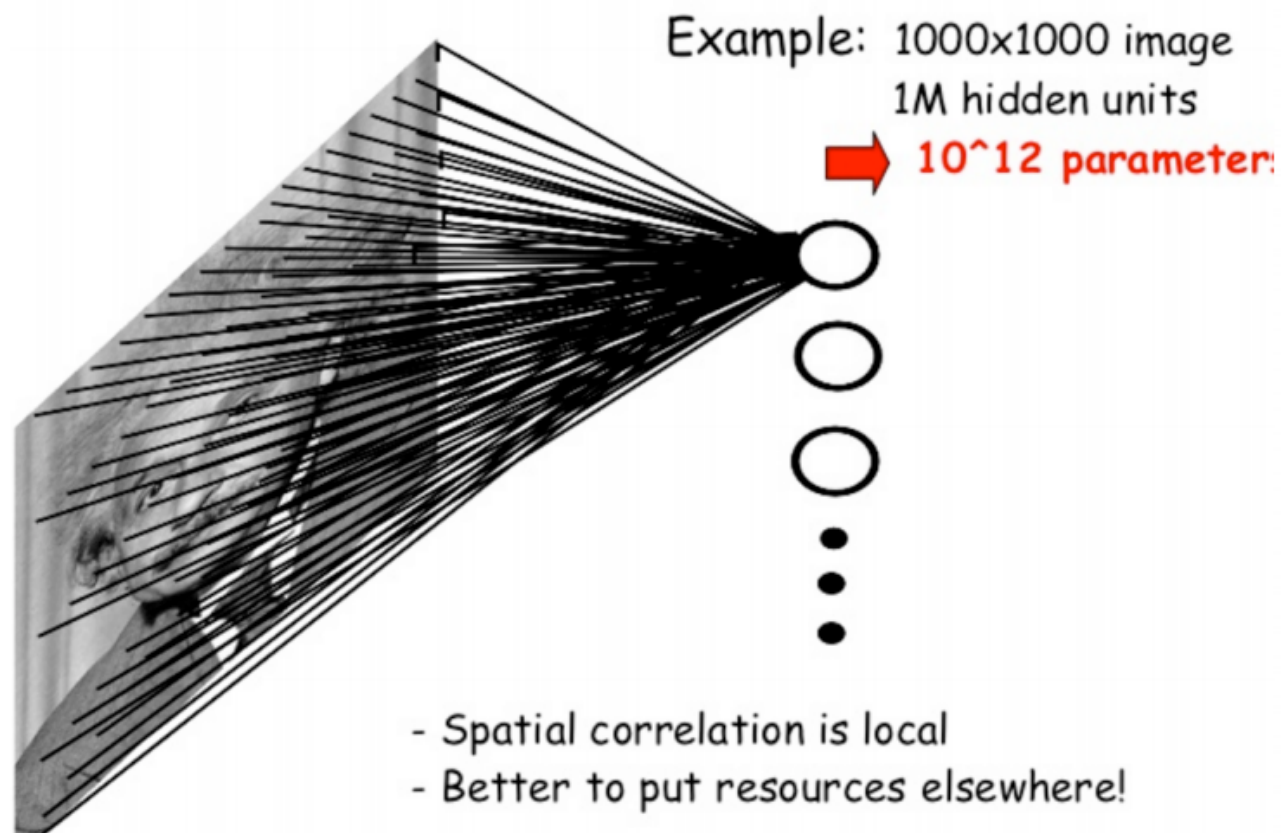
CNNS FOR SPEECH

TIME-DELAY NEURAL NETS

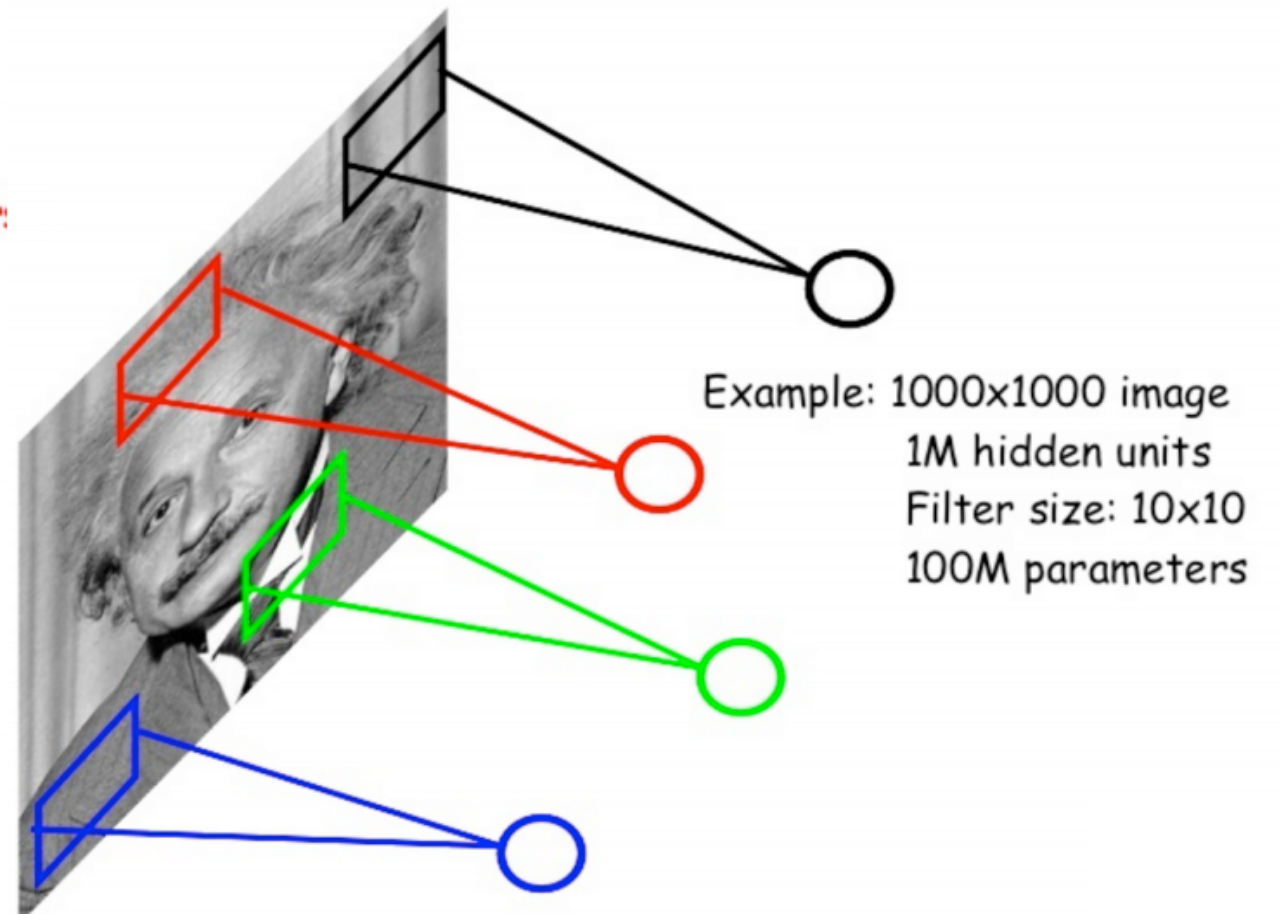
- ▶ In deep neural nets, we model acoustic context by including neighboring frames in input layer
- ▶ In TDNNs (time-delay neural nets), each layer processes a window of context from previous layer
- ▶ Each higher layer has a wider “receptive” field, wider acoustic context
- ▶ “TDNNs are 1D convolutional neural nets (without pooling and with dilations)”
 - ▶ (What part of them is 1D?)

CNNs AND LOCAL FEATURES

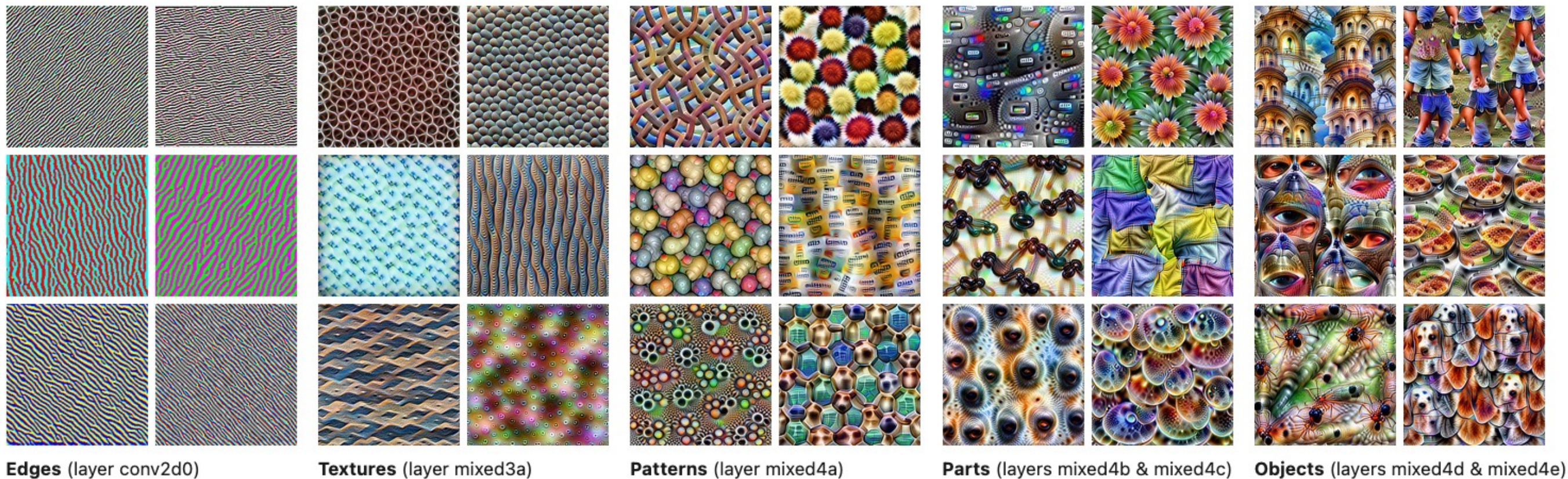
FULLY CONNECTED NEURAL NET



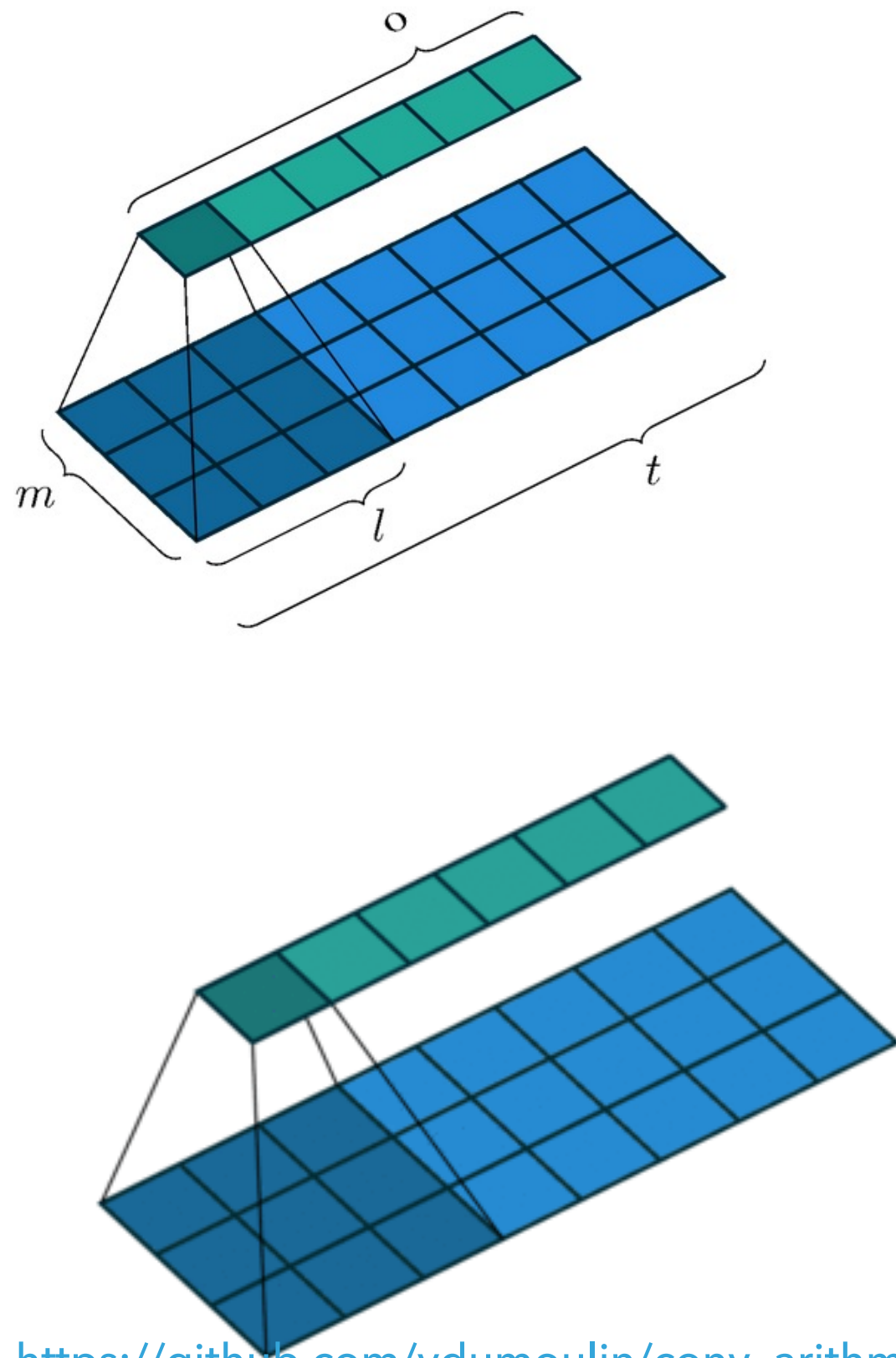
LOCALLY CONNECTED NEURAL NET



CNNs AND LOCAL FEATURES

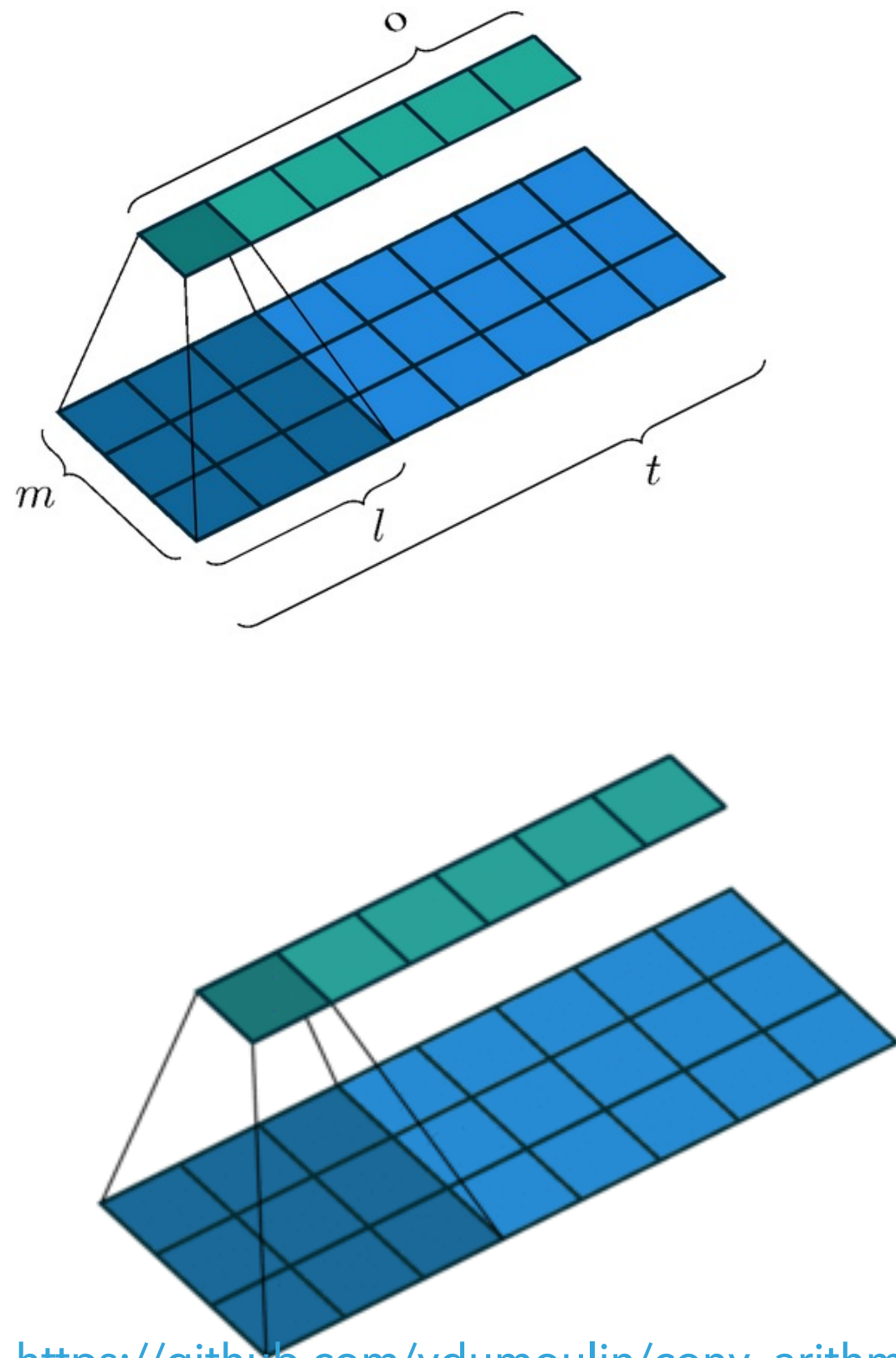


TIME-DELAY NEURAL NETS



- ▶ Input m -dimensional filter bank vector for each frame of speech signal, with t frames in total (light blue)
- ▶ Kernel/filter weight matrix $\mathbf{W}$ with height of m and width of l (dark blue)
- ▶ Output of length o , determined by number of times kernel can fit across the frames of input: (green)

TIME-DELAY NEURAL NETS



- ▶ Input m -dimensional filter bank vector for each frame of speech signal, with t frames in total (light blue)
- ▶ Kernel/filter weight matrix $\mathbf{W}$ with height of m and width of l (dark blue)
- ▶ Output of length o , determined by number of times kernel can fit across the frames of input: (green)

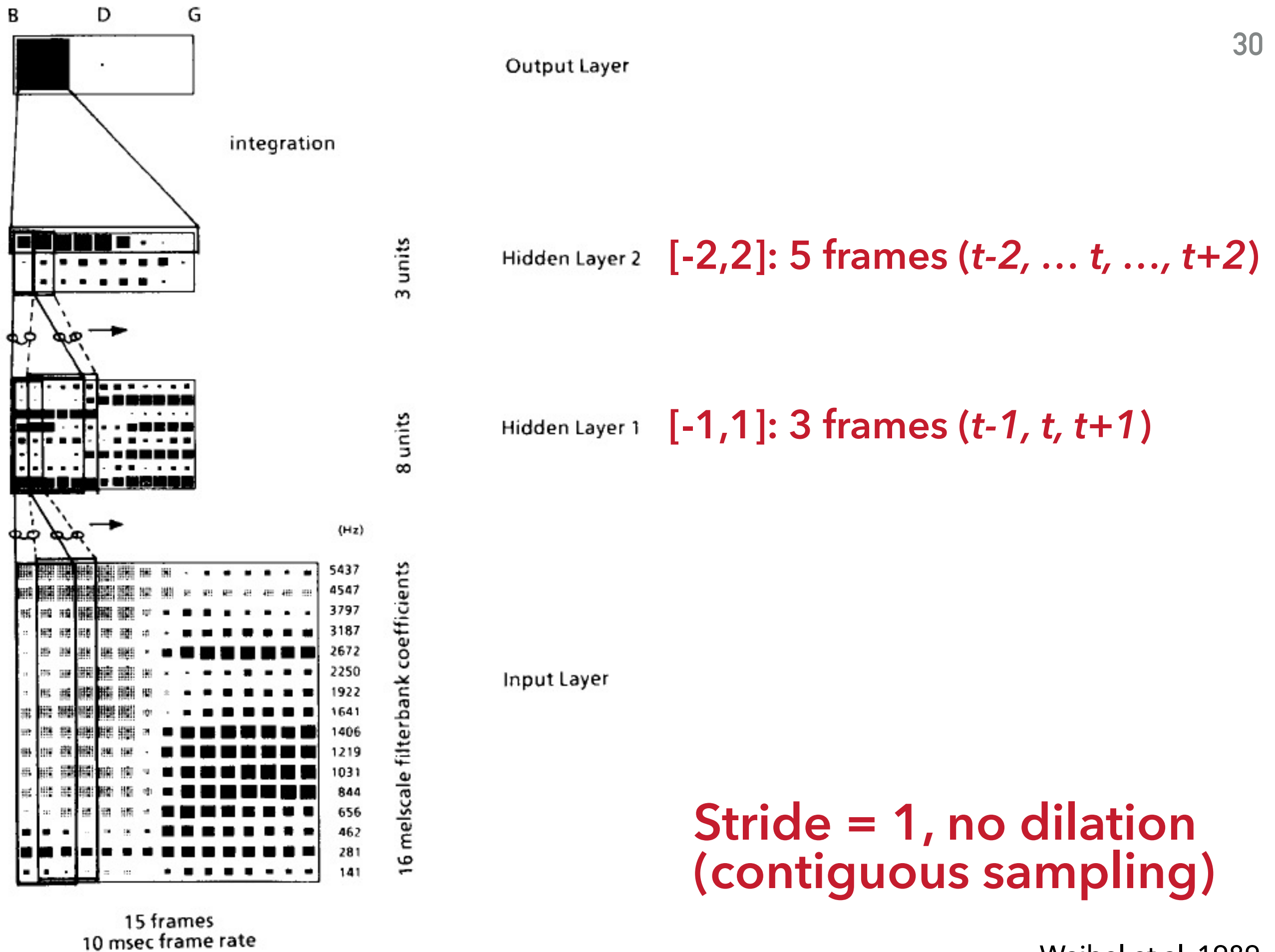


Fig. 2. The architecture of the TDNN.

SUB-SAMPLED TIME-DELAY NEURAL NETS

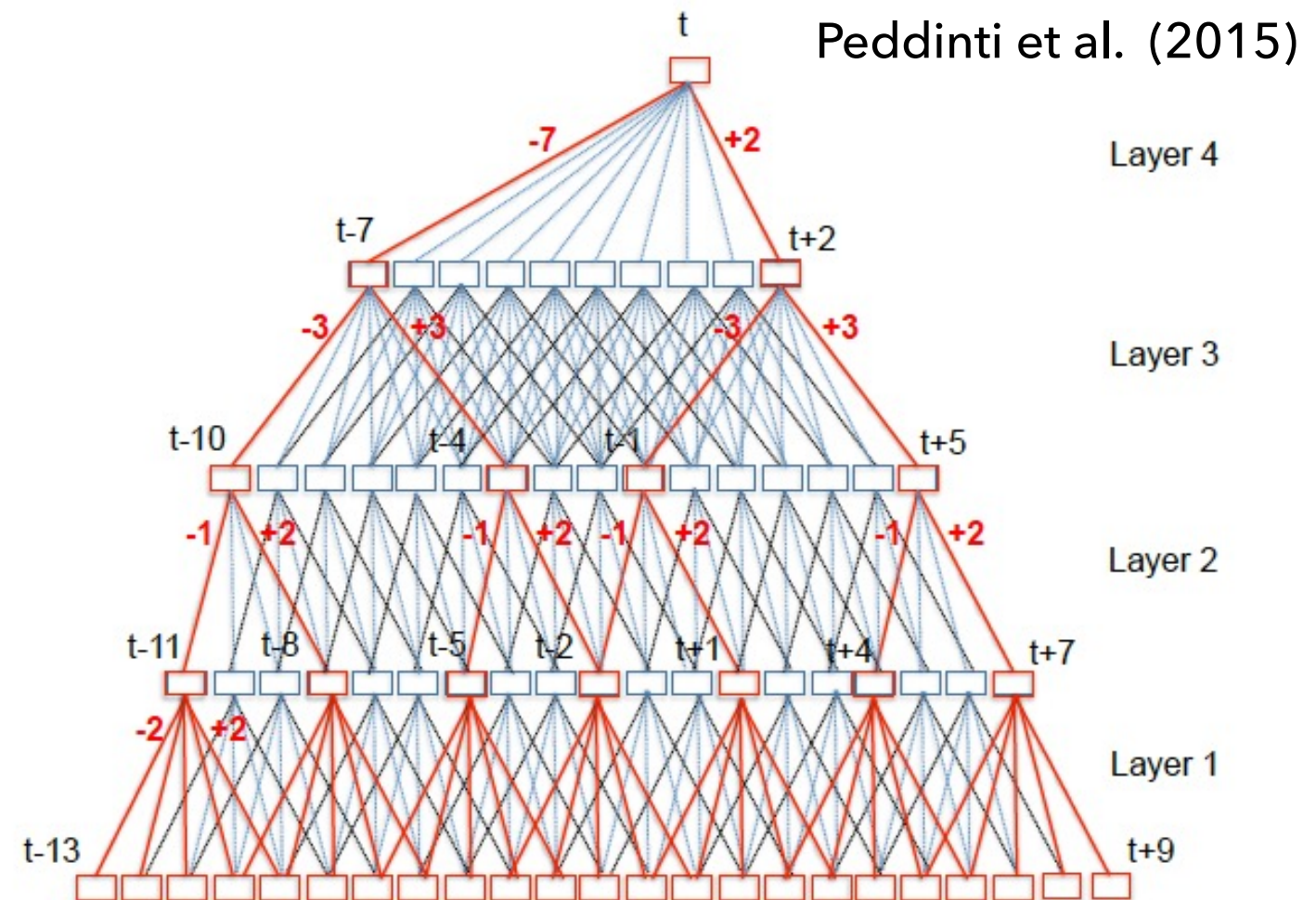
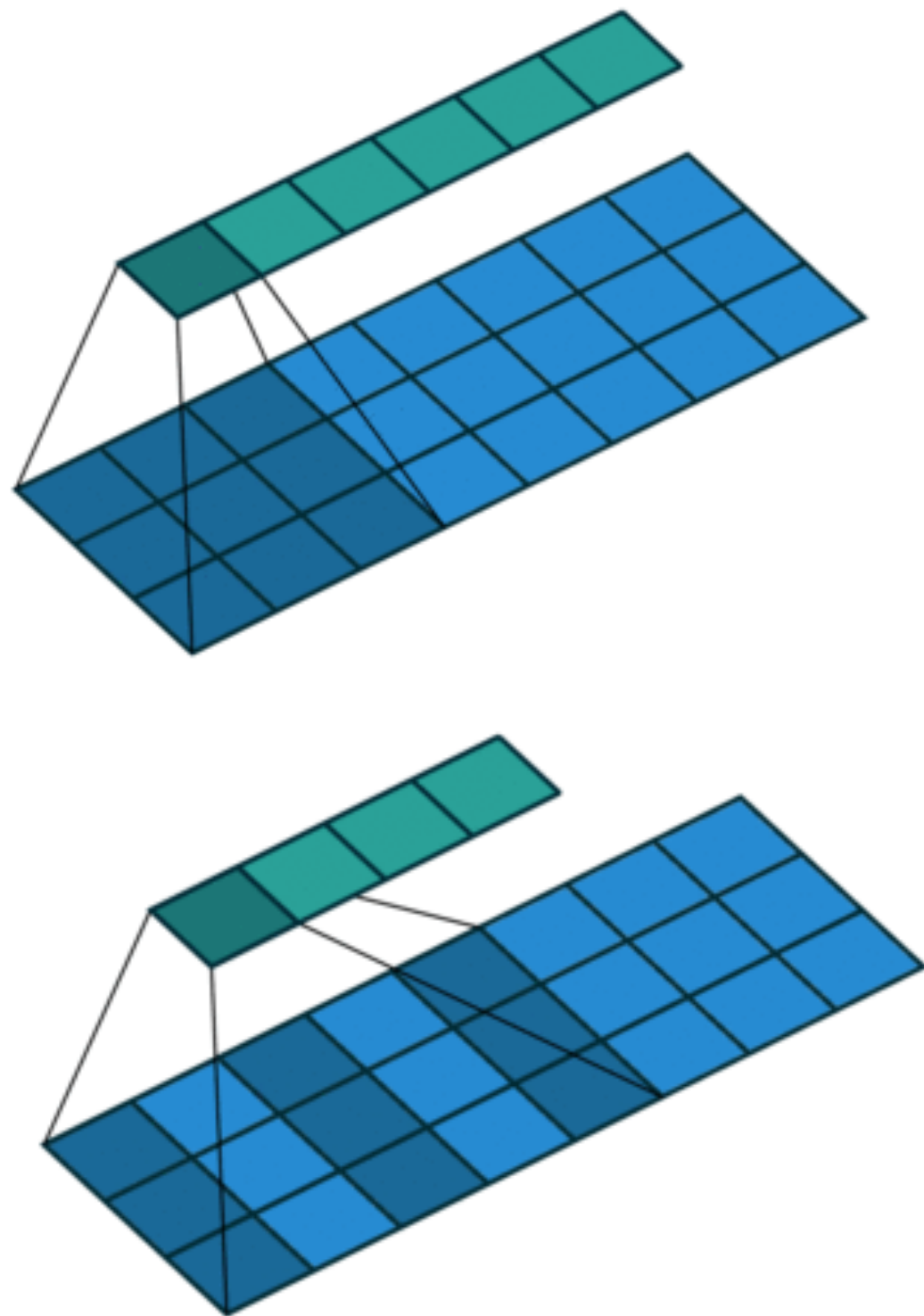


Figure 1: Computation in TDNN with sub-sampling (red) and without sub-sampling (blue+red)

Layer	Input context	Input context with sub-sampling
1	$[-2, +2]$	$[-2, 2]$
2	$[-1, 2]$	$\{-1, 2\}$
3	$[-3, 3]$	$\{-3, 3\}$
4	$[-7, 2]$	$\{-7, 2\}$
5	$\{0\}$	$\{0\}$

SUB-SAMPLED TIME-DELAY NEURAL NETS

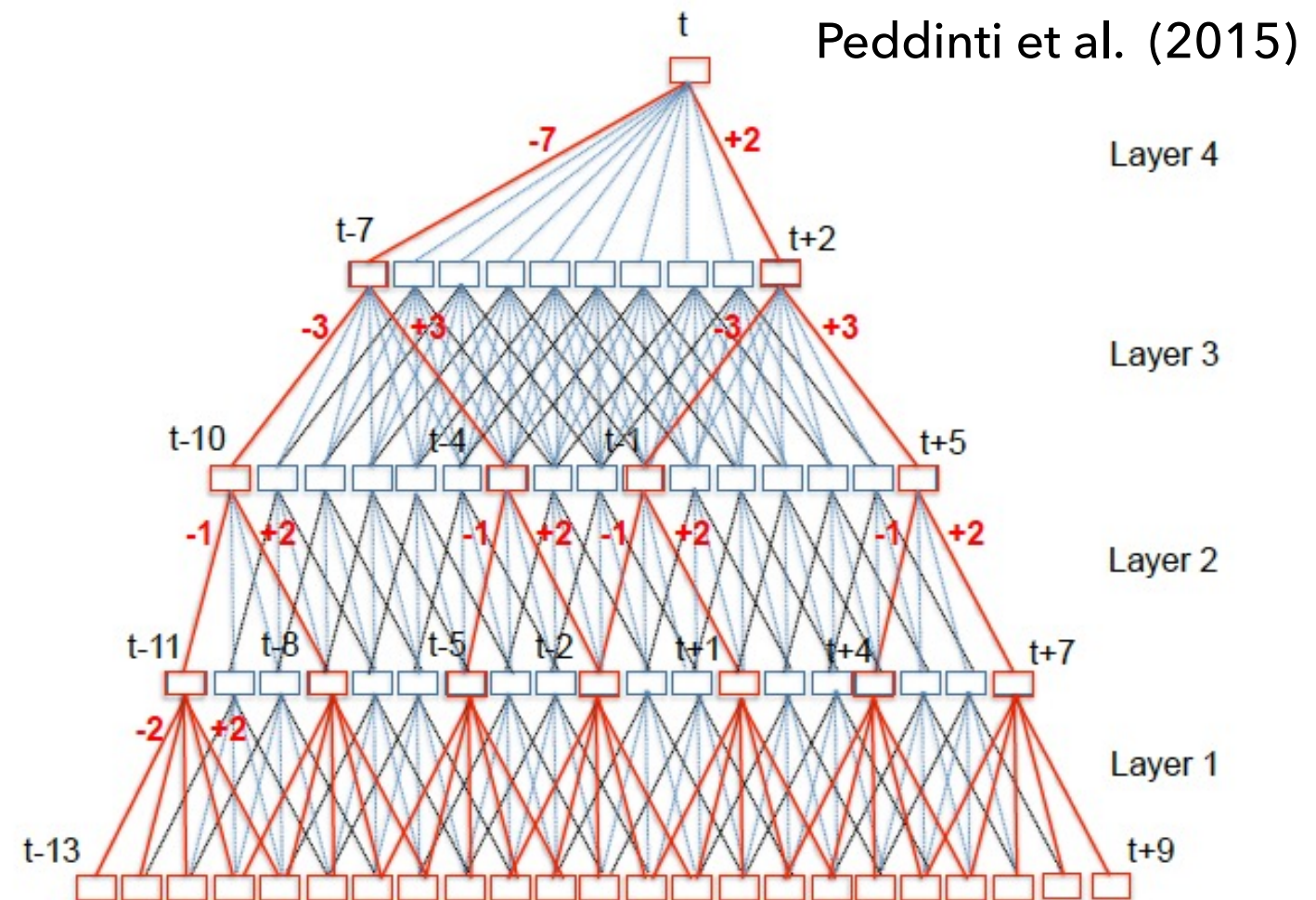
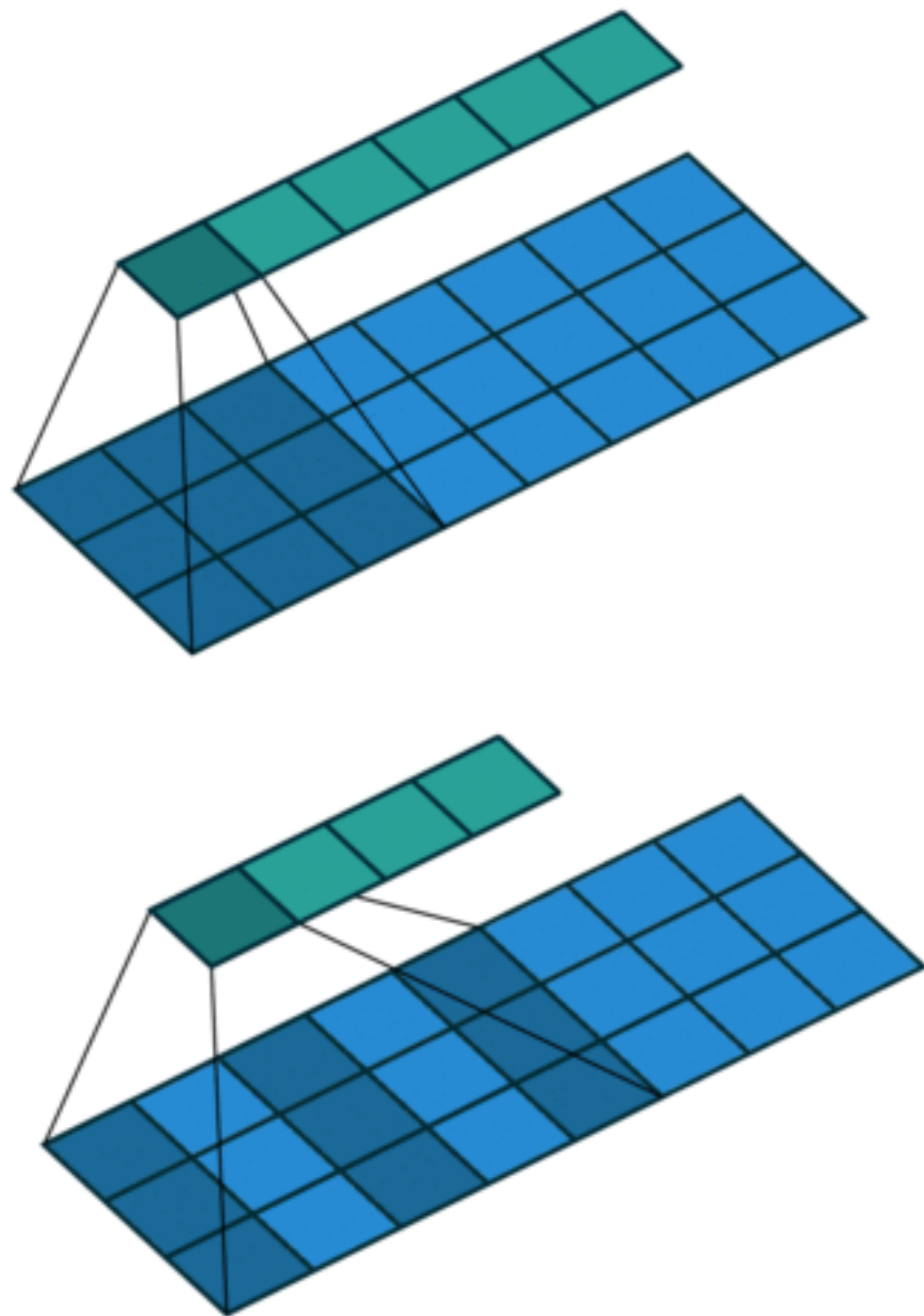


Figure 1: Computation in TDNN with sub-sampling (red) and without sub-sampling (blue+red)

Layer	Input context	Input context with sub-sampling
1	$[-2, +2]$	$[-2, 2]$
2	$[-1, 2]$	$\{-1, 2\}$
3	$[-3, 3]$	$\{-3, 3\}$
4	$[-7, 2]$	$\{-7, 2\}$
5	$\{0\}$	$\{0\}$

SUB-SAMPLED TIME-DELAY NEURAL NETS

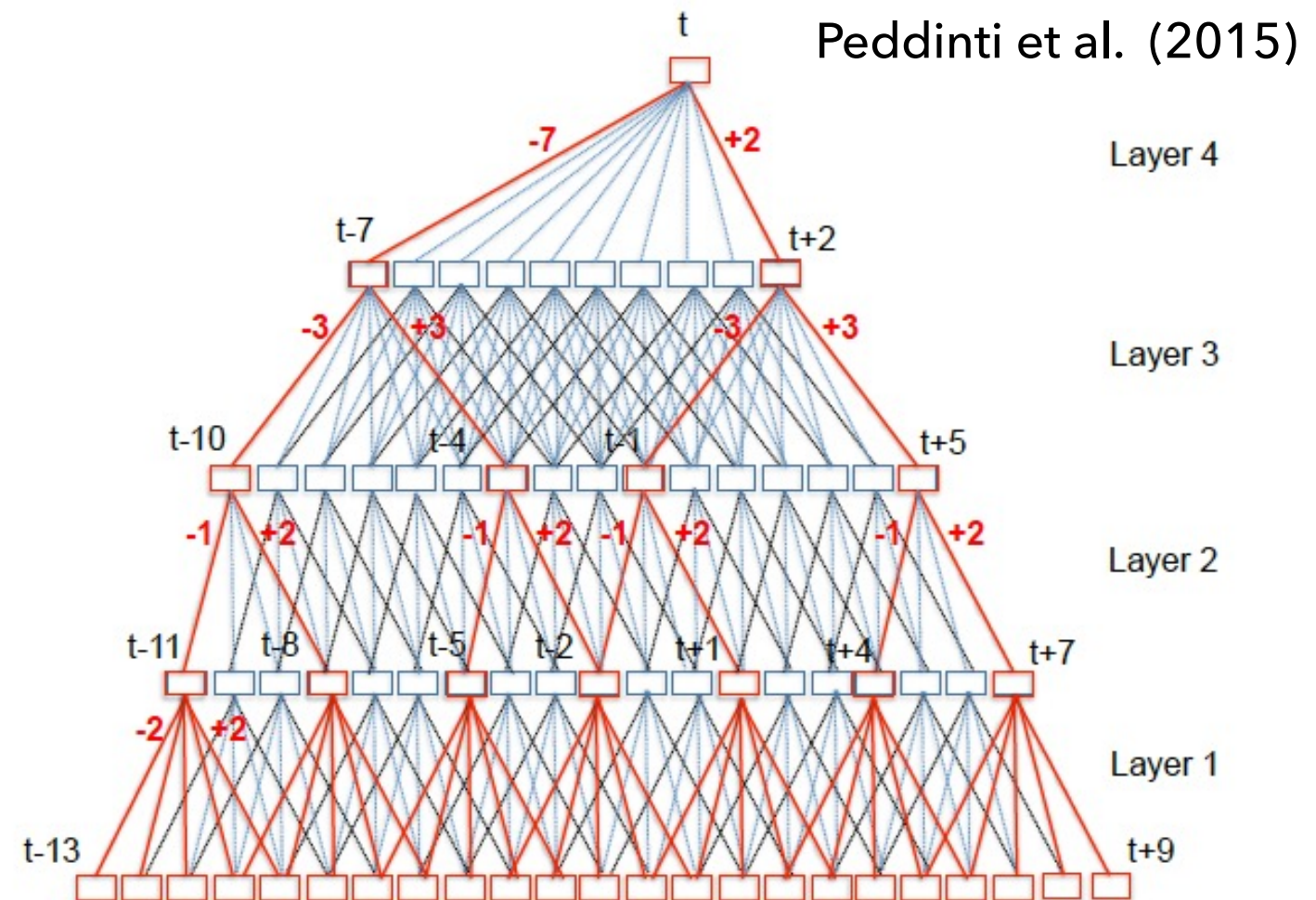
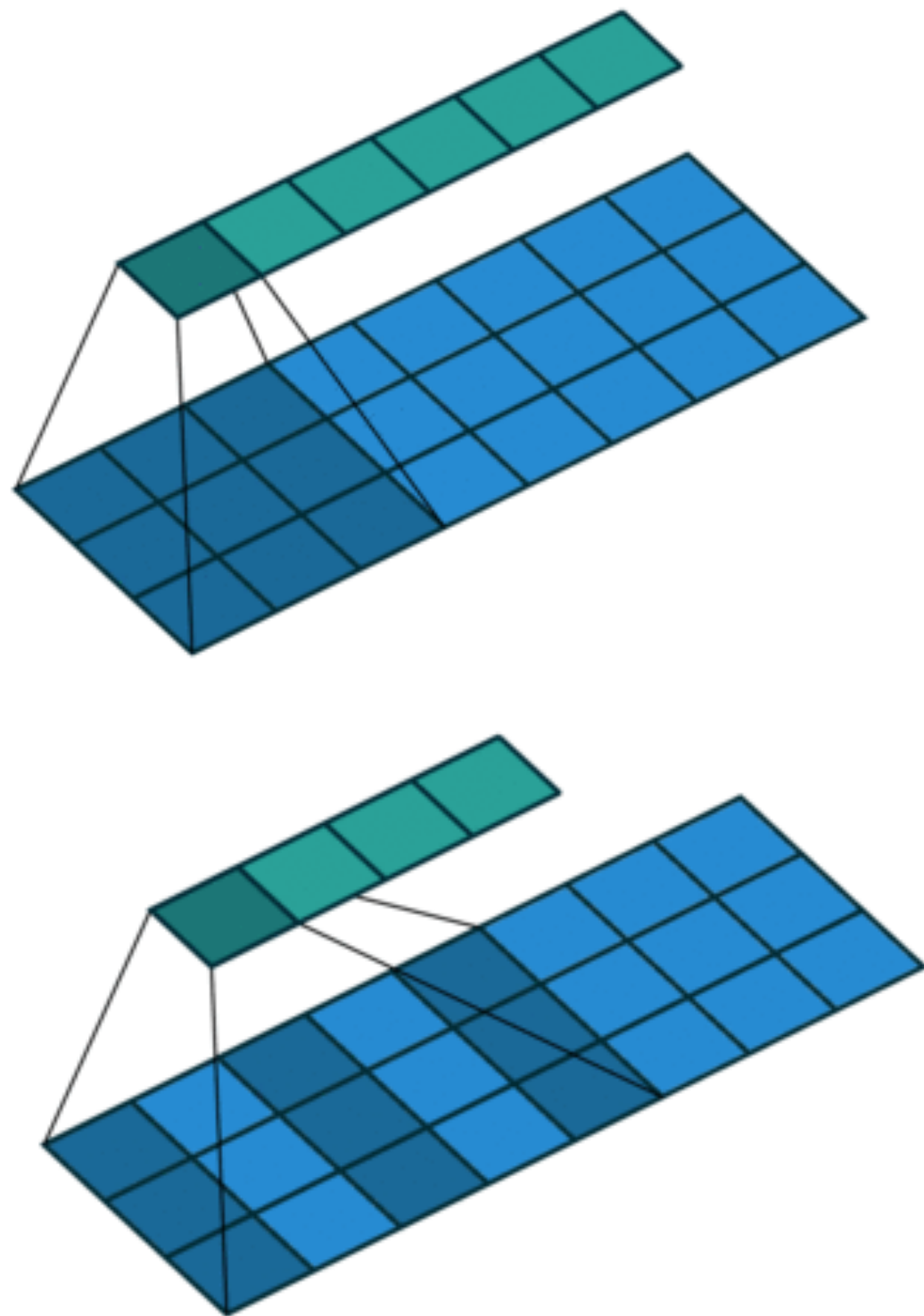


Figure 1: Computation in TDNN with sub-sampling (red) and without sub-sampling (blue+red)

Layer	Input context	Input context with sub-sampling
1	$[-2, +2]$	$[-2, 2]$
2	$[-1, 2]$	$\{-1, 2\}$
3	$[-3, 3]$	$\{-3, 3\}$
4	$[-7, 2]$	$\{-7, 2\}$
5	$\{0\}$	$\{0\}$

“2D” CONVOLUTION BY HAND

https://leonardoaraujosantos.gitbook.io/artificial-intelligence/machine_learning/deep_learning/convolution#doing-by-hand-1