

Lecture 6/4

✓ • Recap NN / BackProp - Brilliant

✓ • Activation(s)

• Auto Encoder

• word vectors

~~✓~~ Demo/playground

✓ • Dropout

~~✓~~ HW 2B - implement 3 ex NNets

- compare yours - vs - pytorch

demo next wed

Course Project

- Proposal by July 4th
- Similar to HW load/hours
- 3 options (Standard)
 - replicate research paper results
 - use class method on new dataset
 - use new method on class-dataset
- non standard: something you are already working on.

Dropout: Randomly (small probs) choose some neurons
to not update / consider for few iterations (iter = 1, 2)
⇒ ignore that neuron. (usually for update)
⇒ regularization effect on not letting NN
concentrate weights on few neurons.
• use heuristic for sampling "idle" neurons.

Activation		
Relu	$\max(0, x) = \begin{cases} x & \text{if } x > 0 \\ 0 & \text{if } x \leq 0 \end{cases}$	not diff in 0 ?
Leaky Relu	$\begin{cases} x & \text{if } x > 0 \\ \alpha x & \text{if } x \leq 0 \end{cases}$	reduce RELU risk for ignored neurons
Sigmoid	$\frac{1}{1 + e^{-x}}$	Bin classif
Tanh	$\frac{e^x - e^{-x}}{e^x + e^{-x}}$	RNN
ELU	$\begin{cases} x & \text{if } x > 0 \\ \alpha(e^x - 1) & \text{if } x \leq 0 \end{cases}$	
swish	$x \cdot \sigma(x)$	slow, but diff
Soft plus	$\ln(1 + e^x)$	

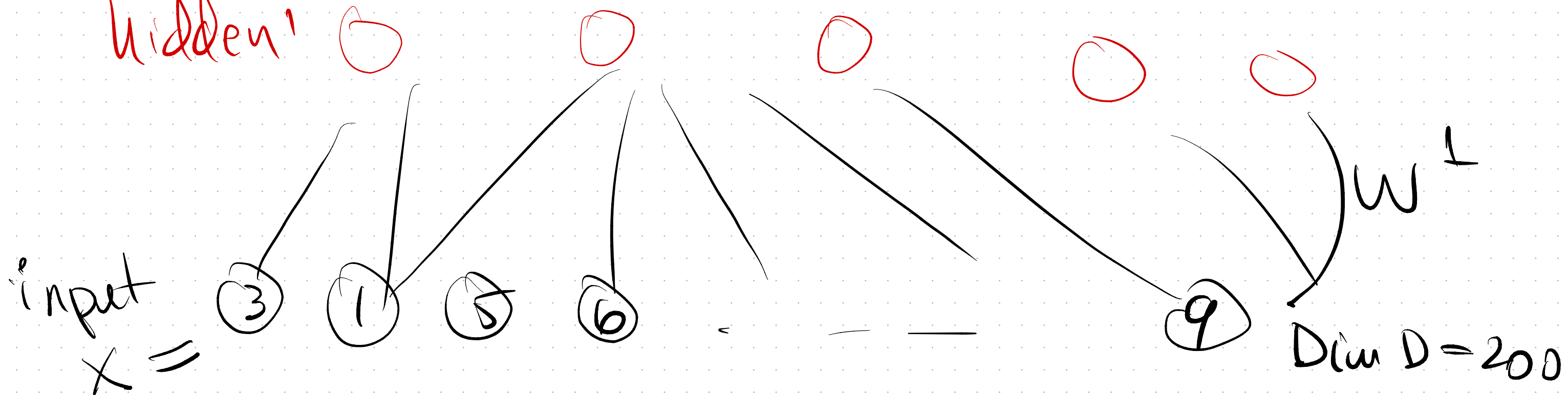
Auto Encoder NN

$X =$
output
= input

hidden

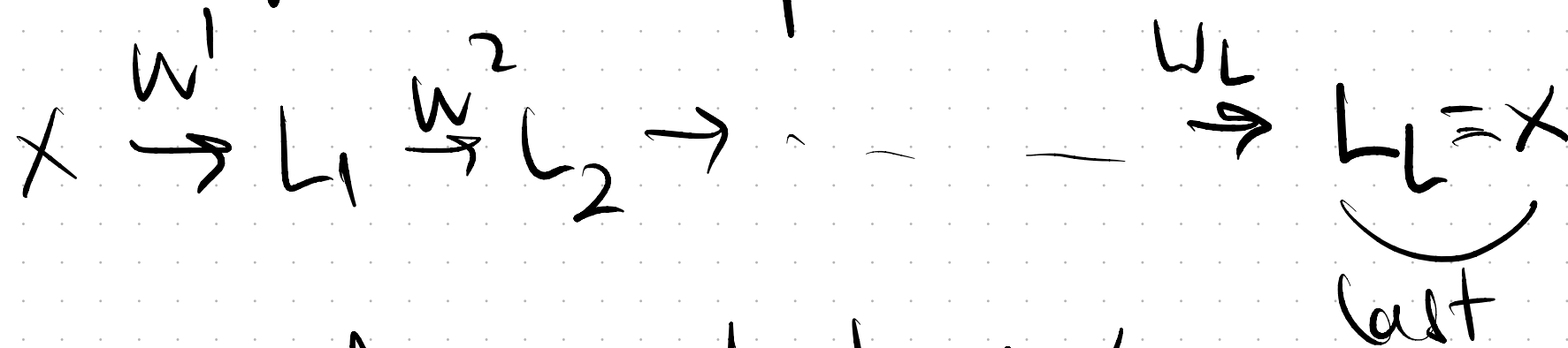
hidden2

hidden1



Task find w^k ($k=1:t$) matrices s.t.

Forward Comp



Success!

Output $\approx$ input

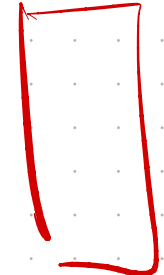
All datapoints
measure $\rightarrow$
Loss (OBS)

produces output $\approx$ input

restriction: hidden layers (at least one) is much smaller
small hid. layer

datap x
0
0
0
0

input
 $D=200$



$D=3$

datap x

0

0

0

0

0

$D=200$

compress = represent data x
on 3-dim

decompress = recover orig
data x
from 3-dim represent

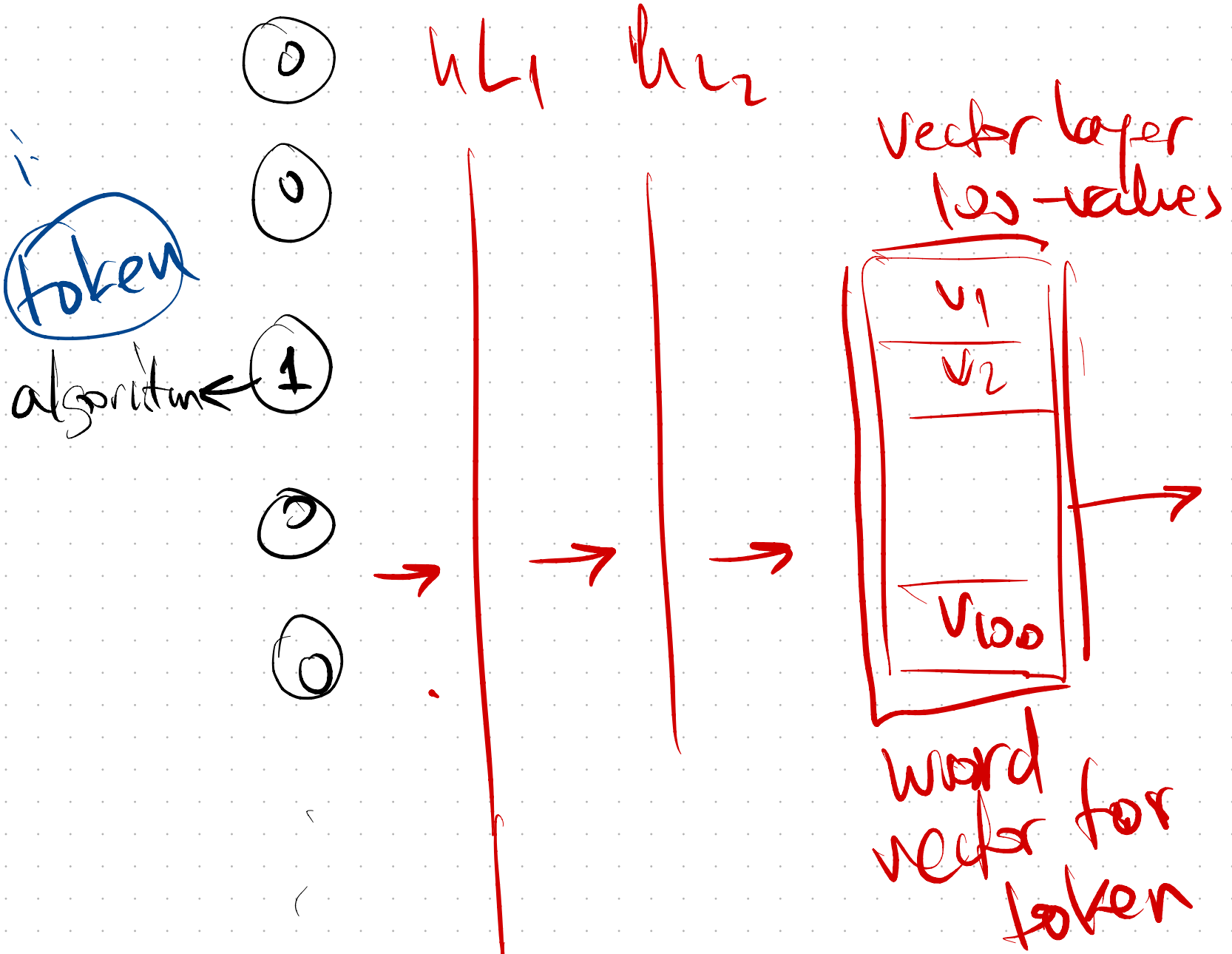
word vectors

word "Algorithm" $\rightarrow$ vector

$[1.2 \mid 13 \mid -5.71 \mid \dots]$

input 1-hot: $D \times D = \# \text{ words } 50,000?$

output $D = 50,000$ (every word)
300-values
100-values



$w_1 \rightarrow \text{prob}(\text{token} \approx w_1)$

$w_2 \rightarrow \text{prob}$

$\rightarrow \text{prob}$

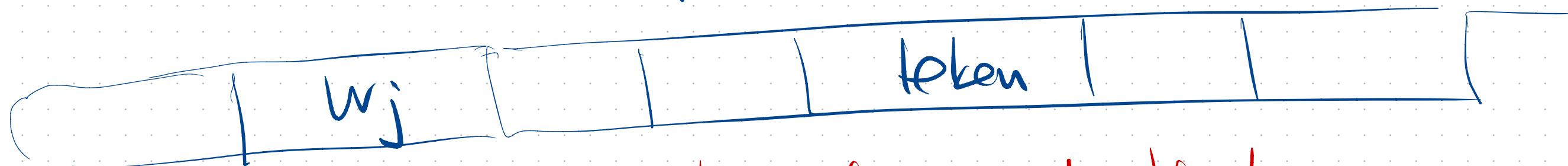
$w_N \rightarrow \text{prob}$
 $\downarrow$
all vocabulary

Labels
given probab

Obj: $\text{prob}(\text{token} \approx w_j) \approx$

- $\text{label}(\text{token}, w_j) = \text{prob/likelihood of synonym. (same meaning)}$
 $\Rightarrow$ word vector similar for words similar meaning.

- $\text{label}_{\text{prob}}(\text{token}, w_j) = \text{co-occurrence in text is span-window of 5 words}$



- $\Rightarrow$ word vectors similar for words that co-occur.