

[1]

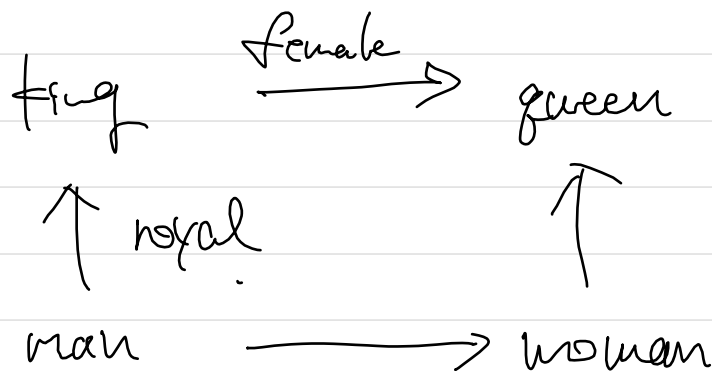
distributed representation

(a) informally state

advantage of dist rep
in terms of generalization

By using distributed representation

we can calc similarity of a certain data
and make it represent some
feature. e.g



(b) n-gram only consider probability

conditioned by n-1 element as follows

$$P(w_1 \dots w_m) = \prod_{i=1}^m P(w_i | w_1 \dots w_{i-1})$$

$$\approx \prod_{i=1}^m P(w_i | w_{i-(n-1)} \dots w_{i-1})$$

20A_Q4

2

e.g. (she was happy prob 0%
he was joyful prob 90%.

consider such a part of sentence

if we can use distributed
representation w/
NN

above example has similarity and
dealt with properly.

but for n -gram, it sometimes
happened that
sentence 1 is appeared
frequently but
sentence 2 is not at all.

2

if we set all w_i to 0,
we cannot get gradient.

and we should avoid

to set large value
it leads grad vanishing

and to set same value it leads to decrease

representation
capacity

20A_Q4

3

勾配消失と表現中の制限に気づいた

for initialization

if we set w too large

activation goes to 0 or 1

(histogram of activation )

it leads grad vanishing.

if we set variance < 0

then all activation value

goes to 0.5

There is no variety of activation

it leads to lost representation capacity

20A_Q4

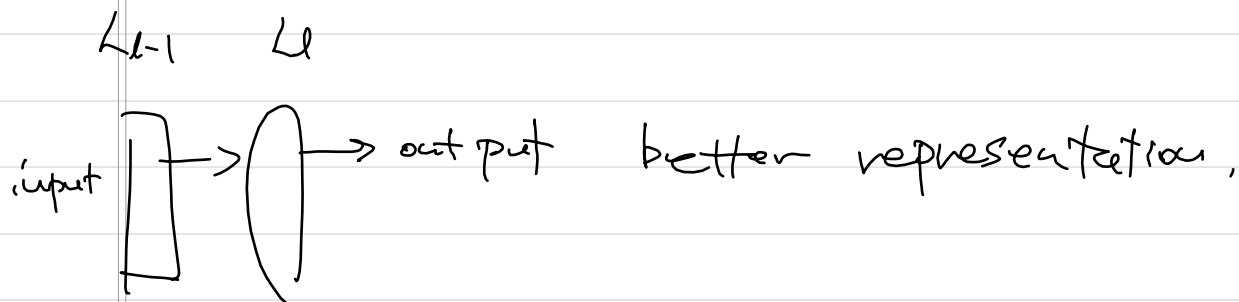
7

[3] (1) pre-training
use trained model by supervised training
to predict other region, and tune
to specific scenario by training w/
limited label.

(2) unsupervised greedy layer-wise
pre-training

$\Rightarrow$ Auto Encoder
which initialize w properly and regularize
model.

(3) supervised greedy layer-wise
pre-training



this enable to avoid vanishing grad
by setting appropriate initialization
for RBM, DMN.

20A_Q4

5

(4)

FFN is affected by initialization.
in terms of grad vanishing
and exploding.

RNN is affected by
eigenvalue of Affine
transformations
which is repeated.

it can be trained

iff $|A| < 1$
