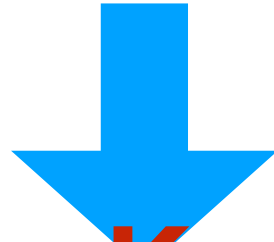


~~Specializing Word Embeddings (for Parsing) by Information Bottleneck~~ Threshing Word Embeddings

Xiang Lisa Li and Jason Eisner



Pre-trained Word Embeddings



Syntactic Knowledge



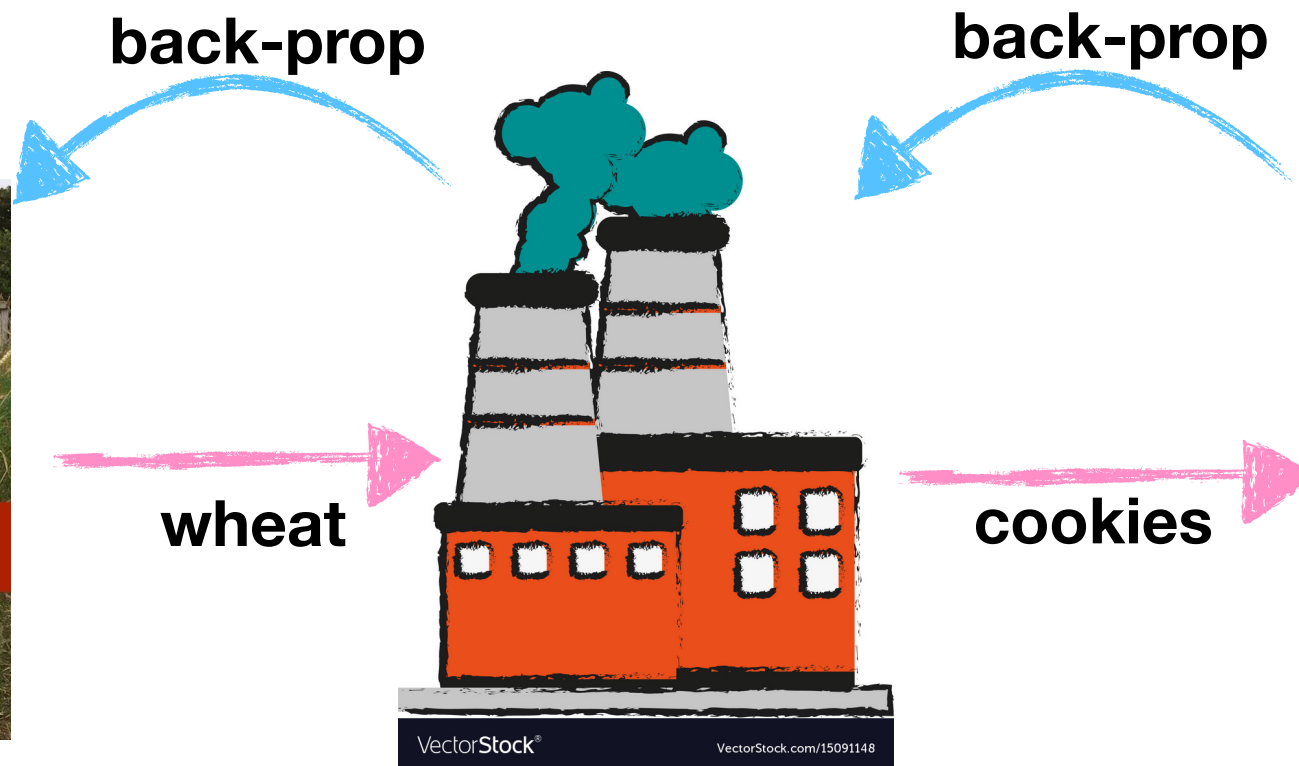
**How to extract syntax
automatically?**

Specialize the Embeddings!

for the task (parsing)



wheat field



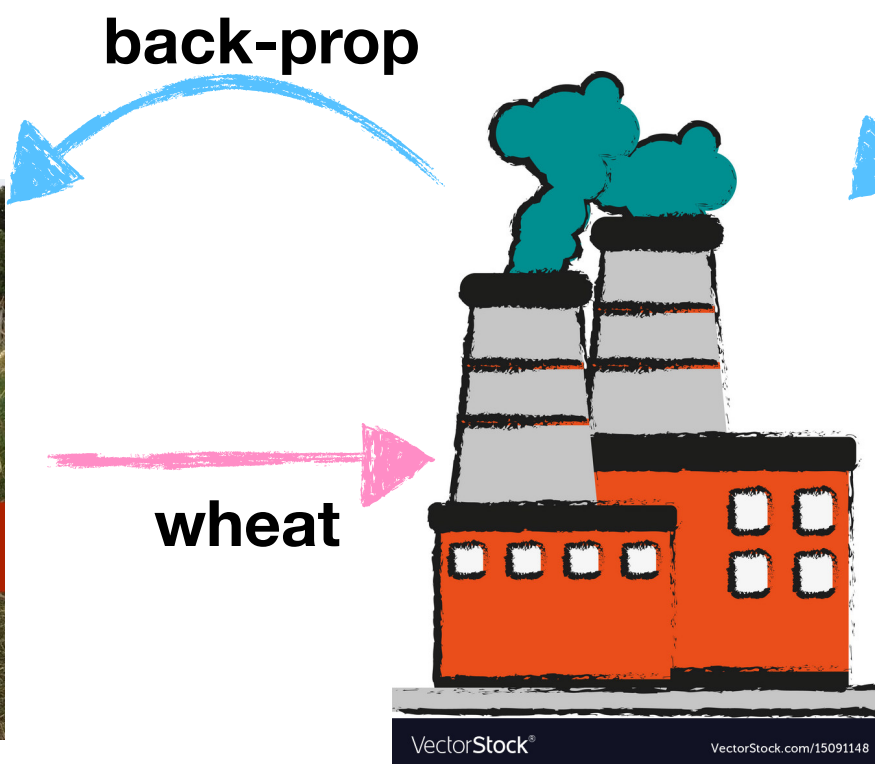
Factory



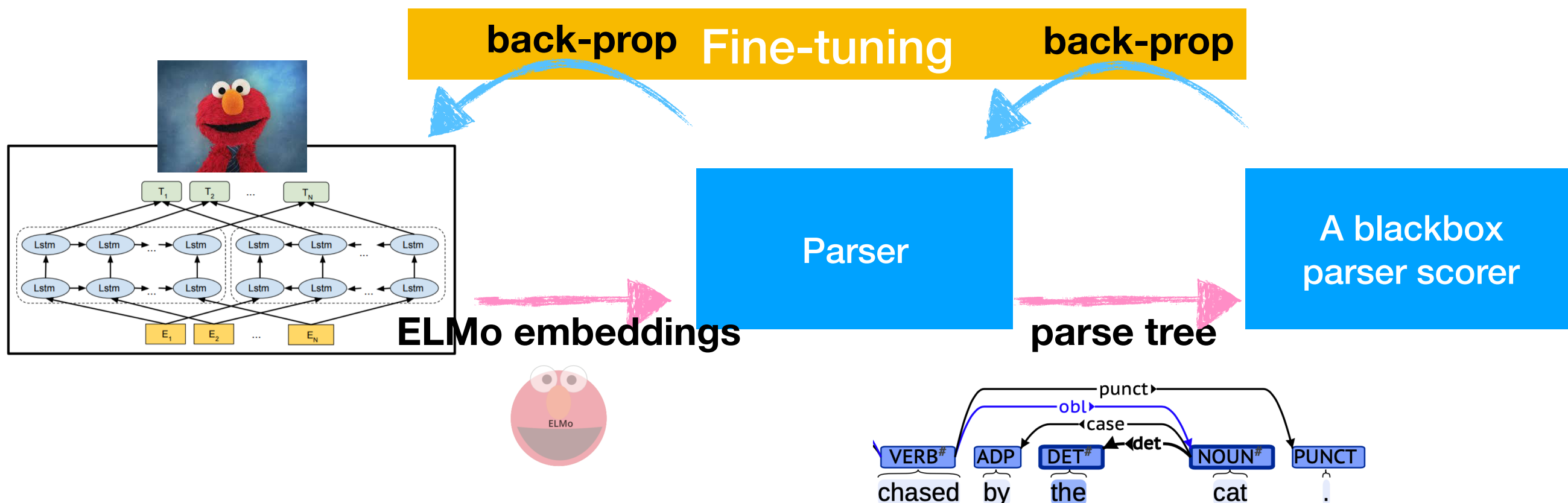
Fine-tuning



wheat field



Factory



What Is ~~Threshing~~?

**Specializing
Word Embeddings**

~~Chaff~~

**Sentiment,
Topic, Semantics ...**

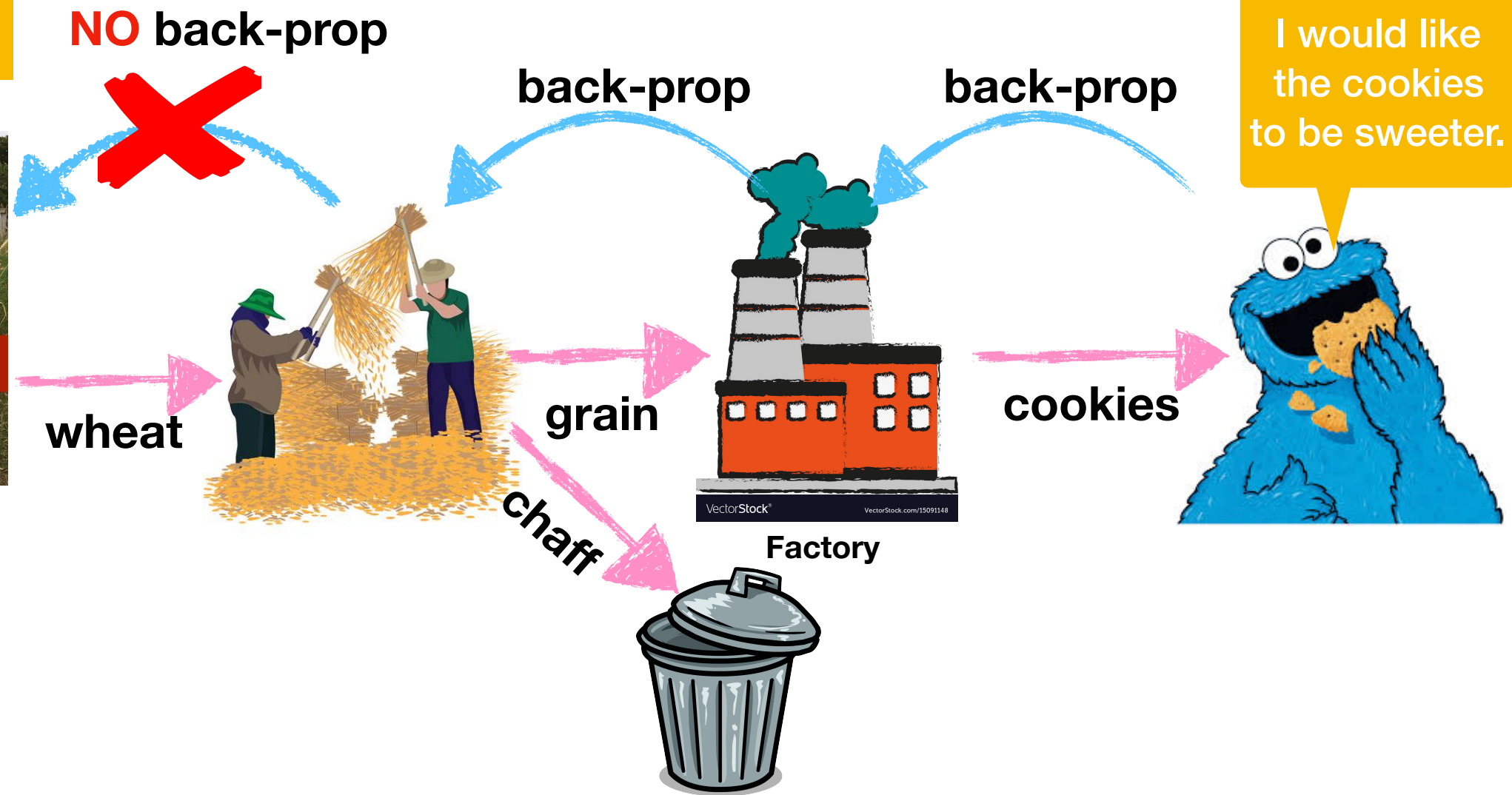
~~Grain~~

Syntax

Specialization:



wheat



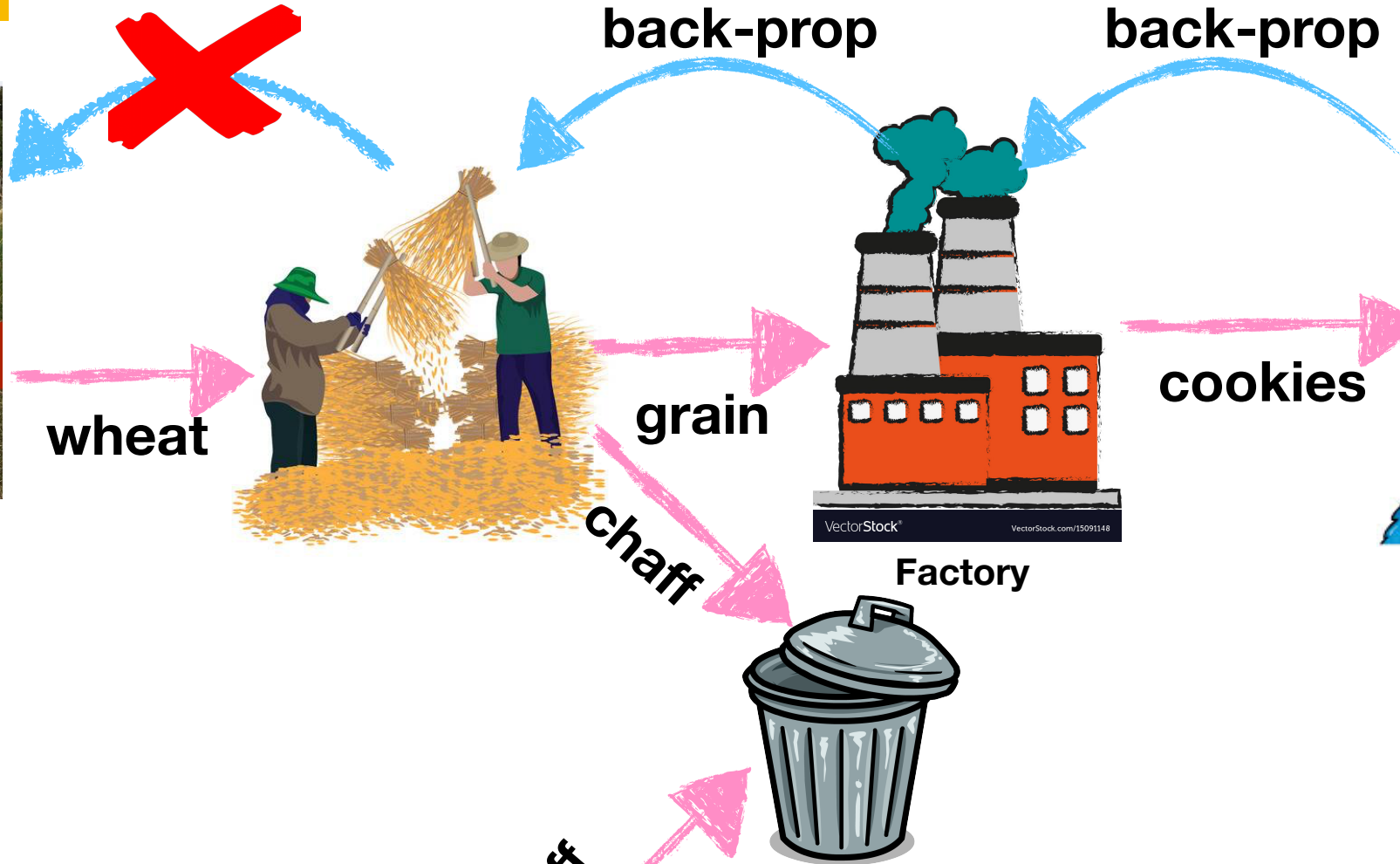
Specialization:



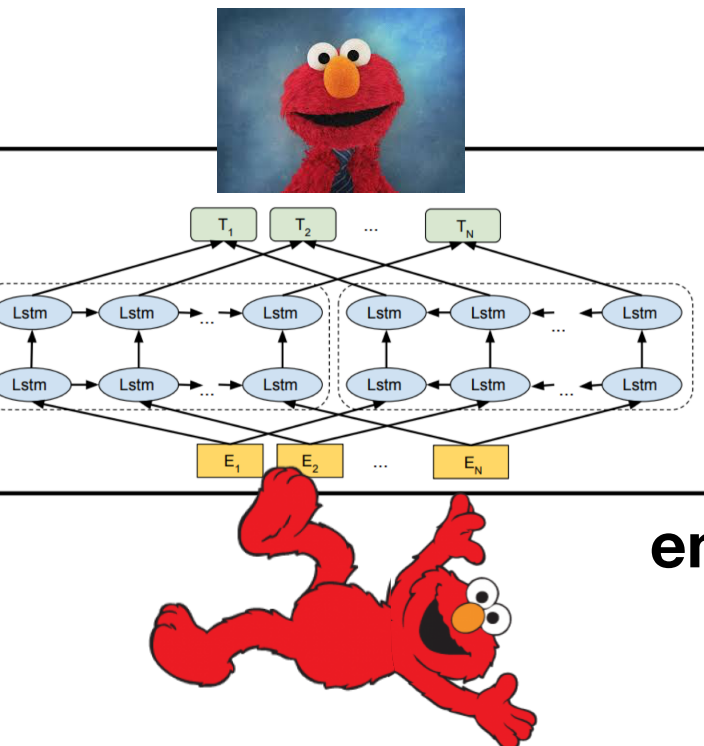
Growing Wheat

wheat

NO back-prop



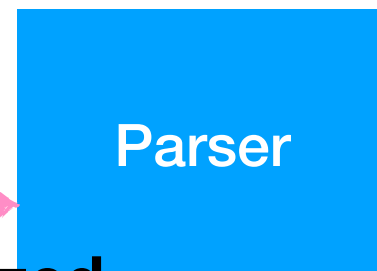
I would like the cookies to be sweeter.



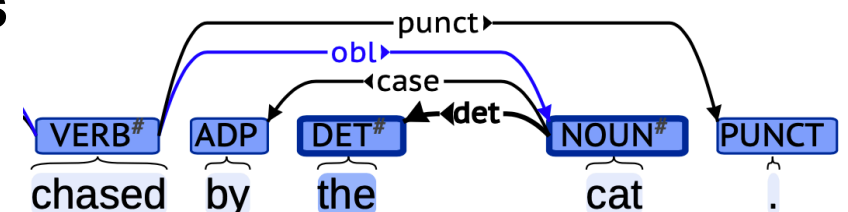
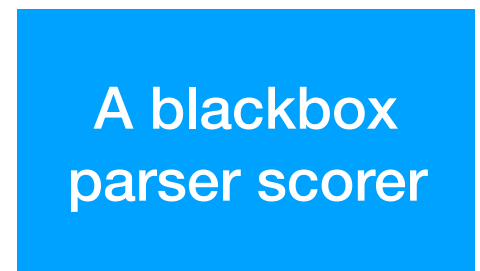
ELMo embeddings



specialized embeddings or tags



parse tree



Specialization:



Growing Wheat

wheat

NO back-prop

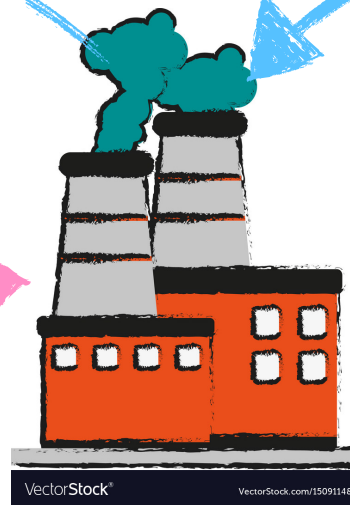


wheat



grain

chaff



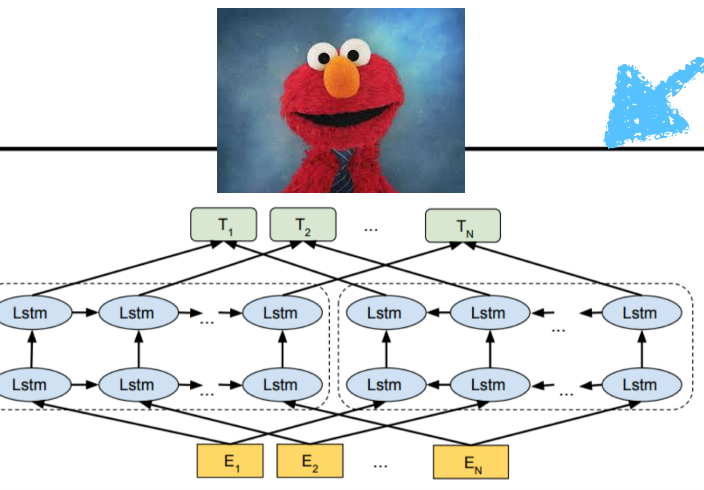
Factory

cookies



I would like the cookies to be sweeter.

NO back-prop



ELMo embeddings



specialized embeddings or tags

chaff

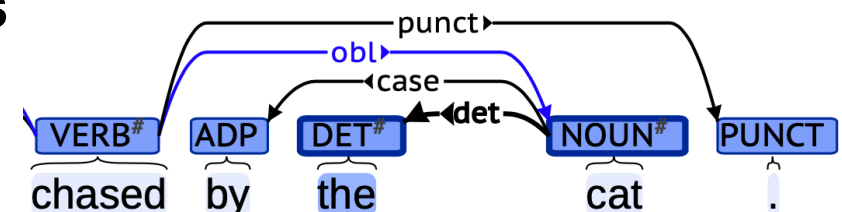
back-prop

back-prop

Parser

A blackbox parser scorer

parse tree



Specialization:



Growing Wheat

wheat

NO back-prop

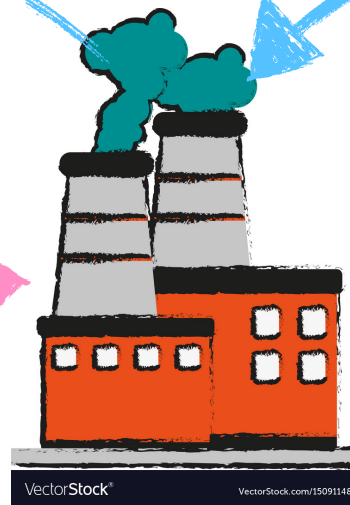


wheat



grain

chaff



Factory

cookies

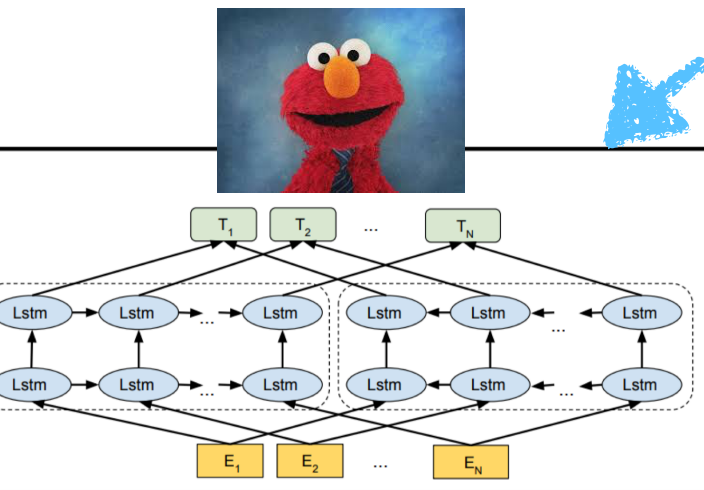


I would like the cookies to be sweeter.

back-prop

back-prop

NO back-prop



ELMo

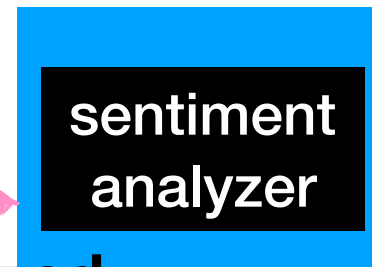
embeddings



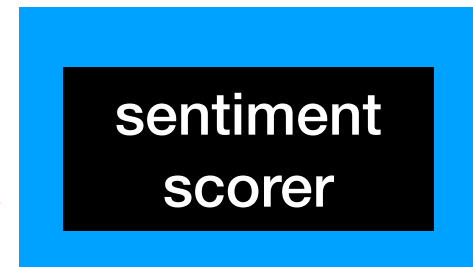
chaff

back-prop

back-prop



sentiment analyzer



sentiment scorer

sentiment

specialized

embeddings

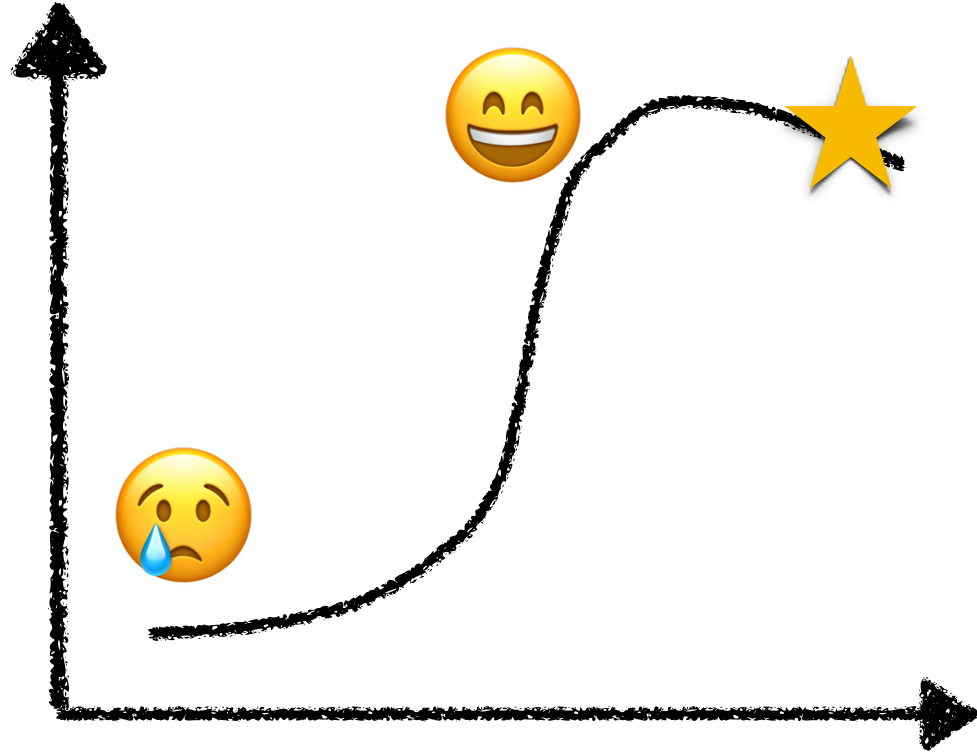
Our specialization technique is easily adaptable by swapping blackboxes !

Pros of Specialization

vs. Fine Tuning

- Faster — trains on 100 sentences per second
- Better generalization — few parameters, so harder to overfit.

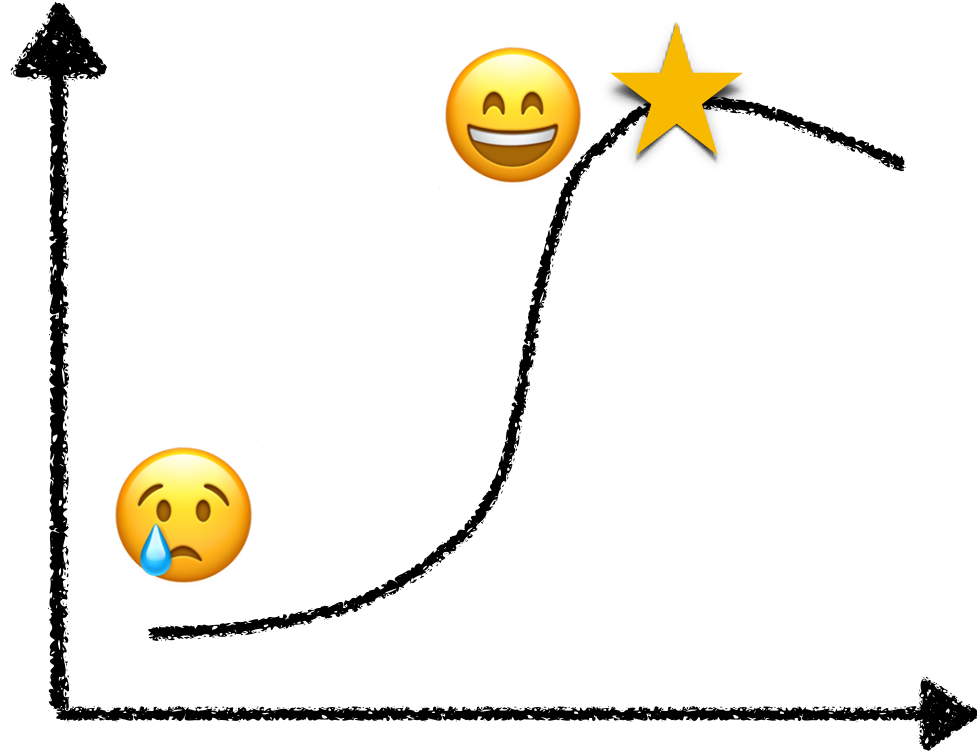
Accuracy



Information Kept



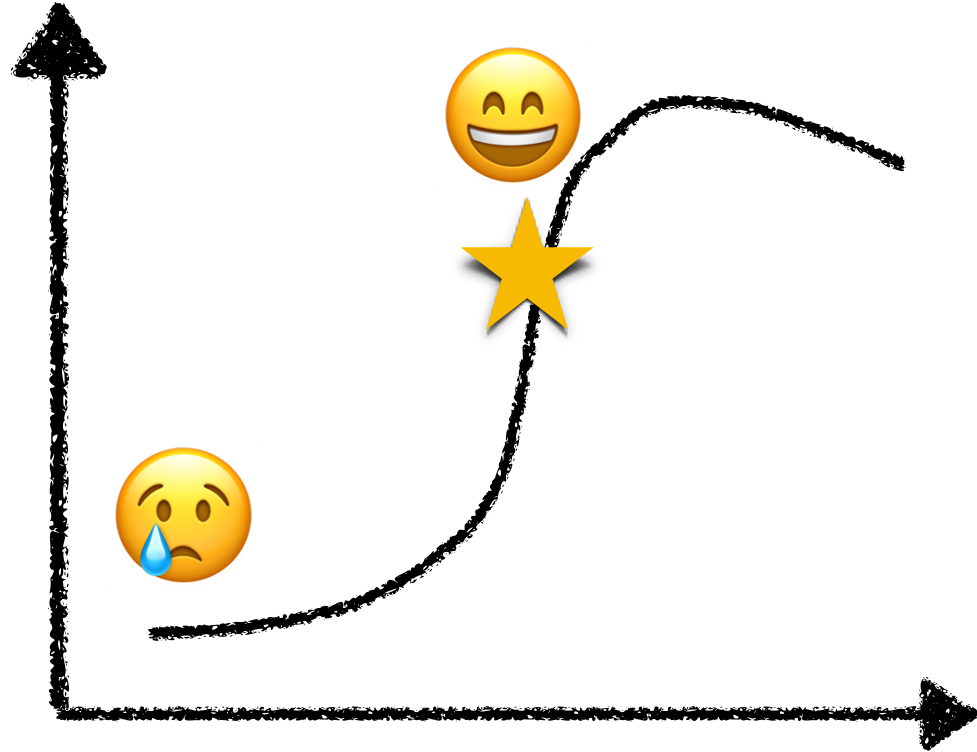
Accuracy



Information Kept



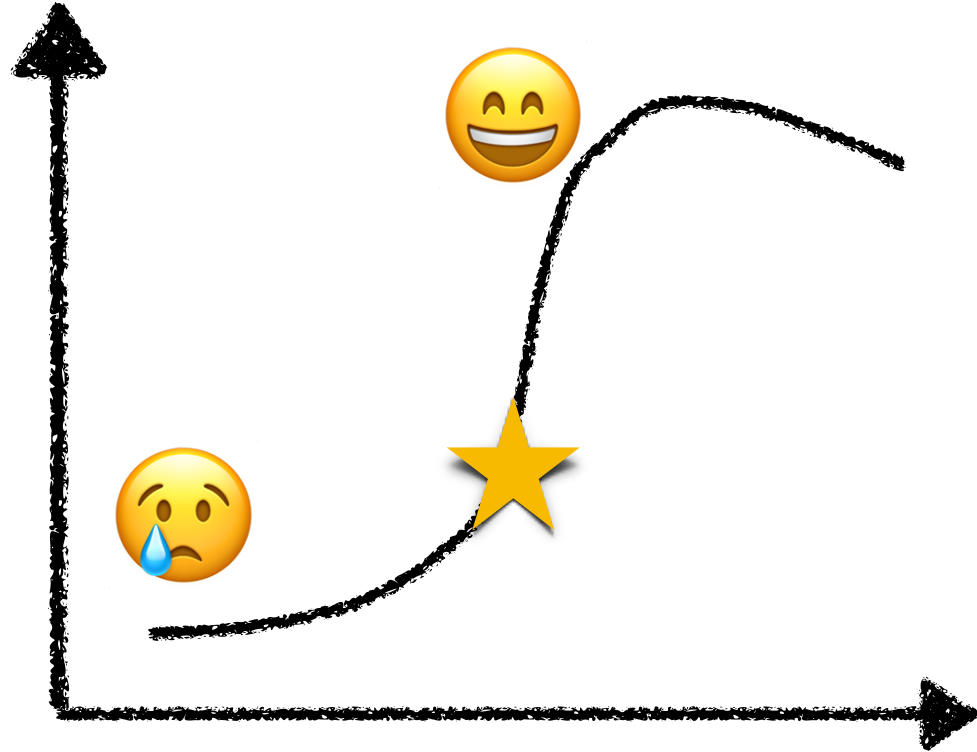
Accuracy



Information Kept



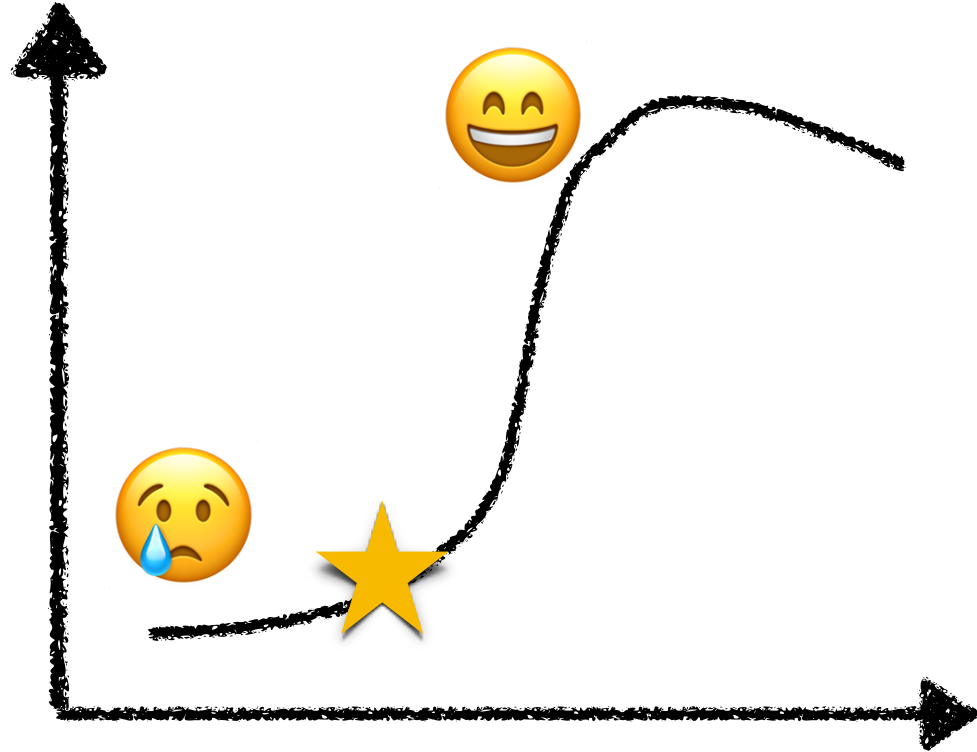
Accuracy



Information Kept



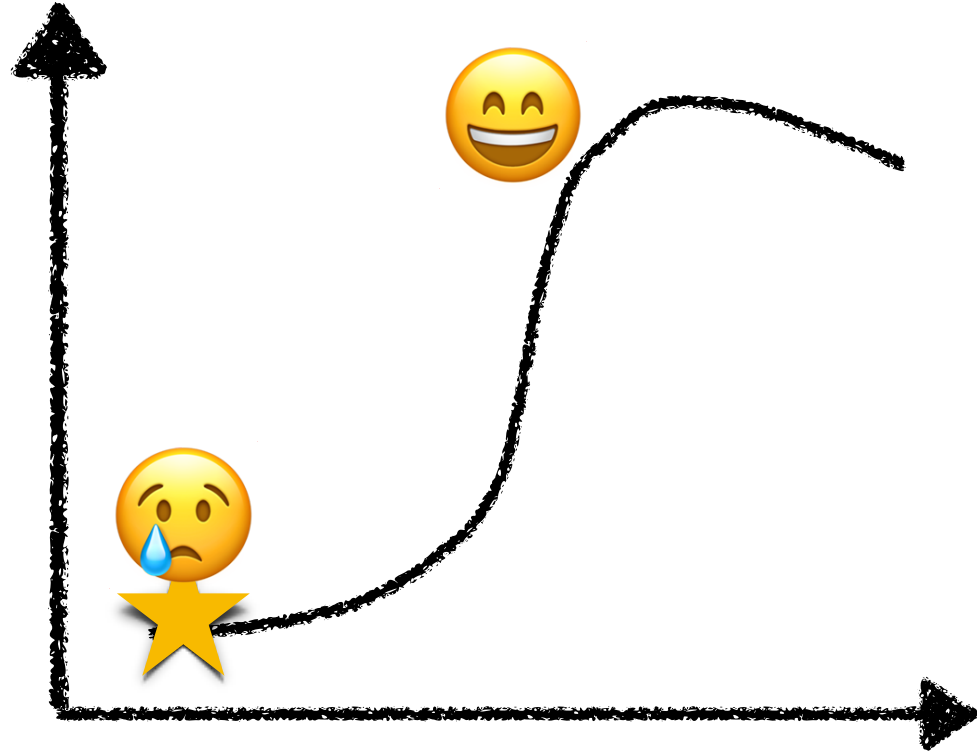
Accuracy



Information Kept



Accuracy



Information Kept



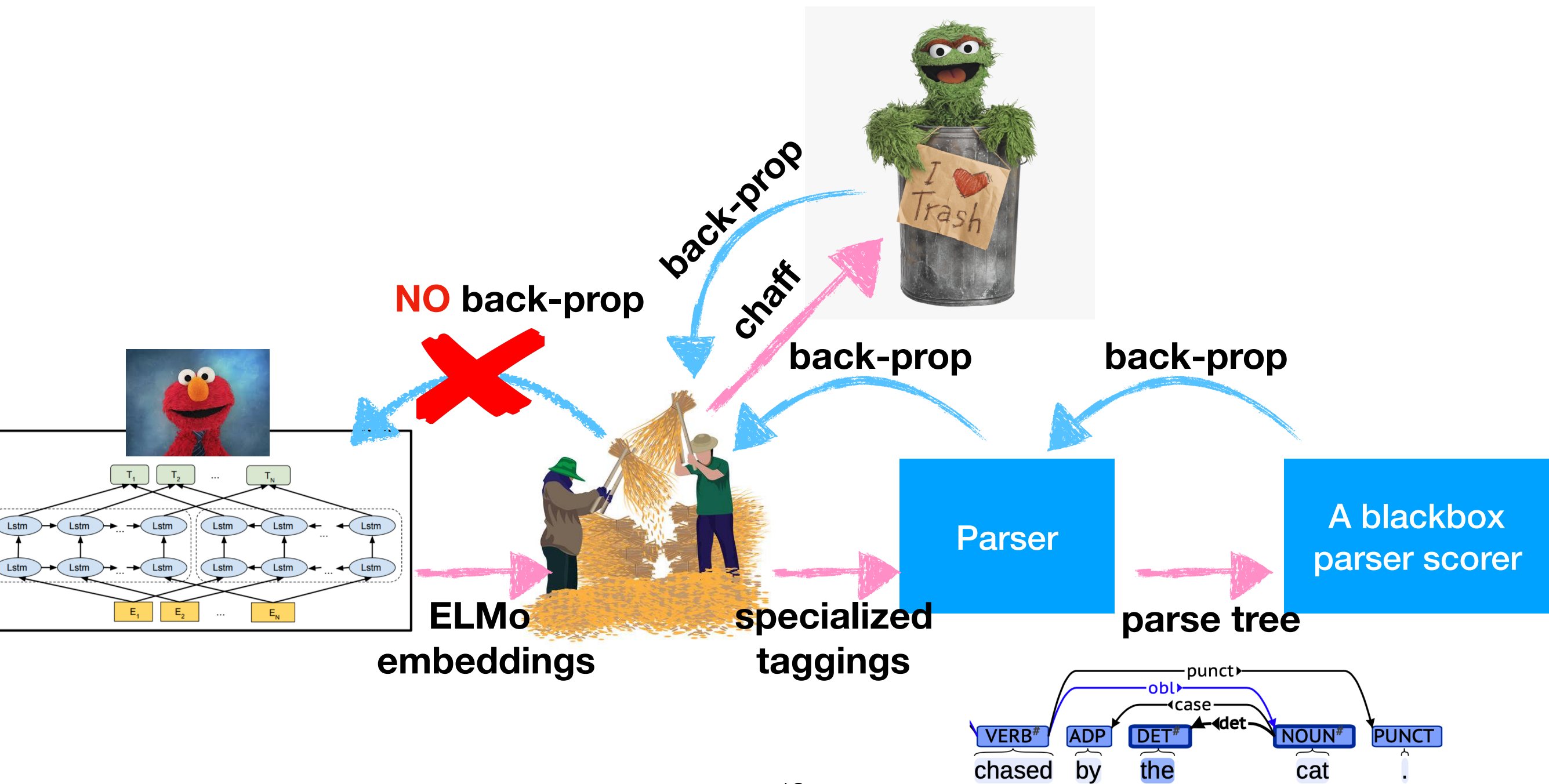
Accuracy

How exactly does the
specialization work?



Information Bottleneck





Information Bottleneck Objective

$$\mathcal{L} = - \underbrace{I(\text{dependency parse}; \text{specialized representation})}_{\text{mutual information between dependency parses and specialized representation.}} + \beta \cdot \underbrace{I(\text{Elmo}; \text{specialized representation})}_{\text{mutual information between pre-trained ELMo embeddings and specialized representation.}}$$

keep more syntactic information

discard more information

bigger β —  compression
Oscar happy



$$\mathcal{L} = -I(\text{chased by the cat.}; \text{spoon}) + \beta I(\text{Elmo}; \text{spoon})$$



$$\cancel{H(\text{chased by the cat.})} - H(\text{chased by the cat.} \mid \text{spoon})$$

constant

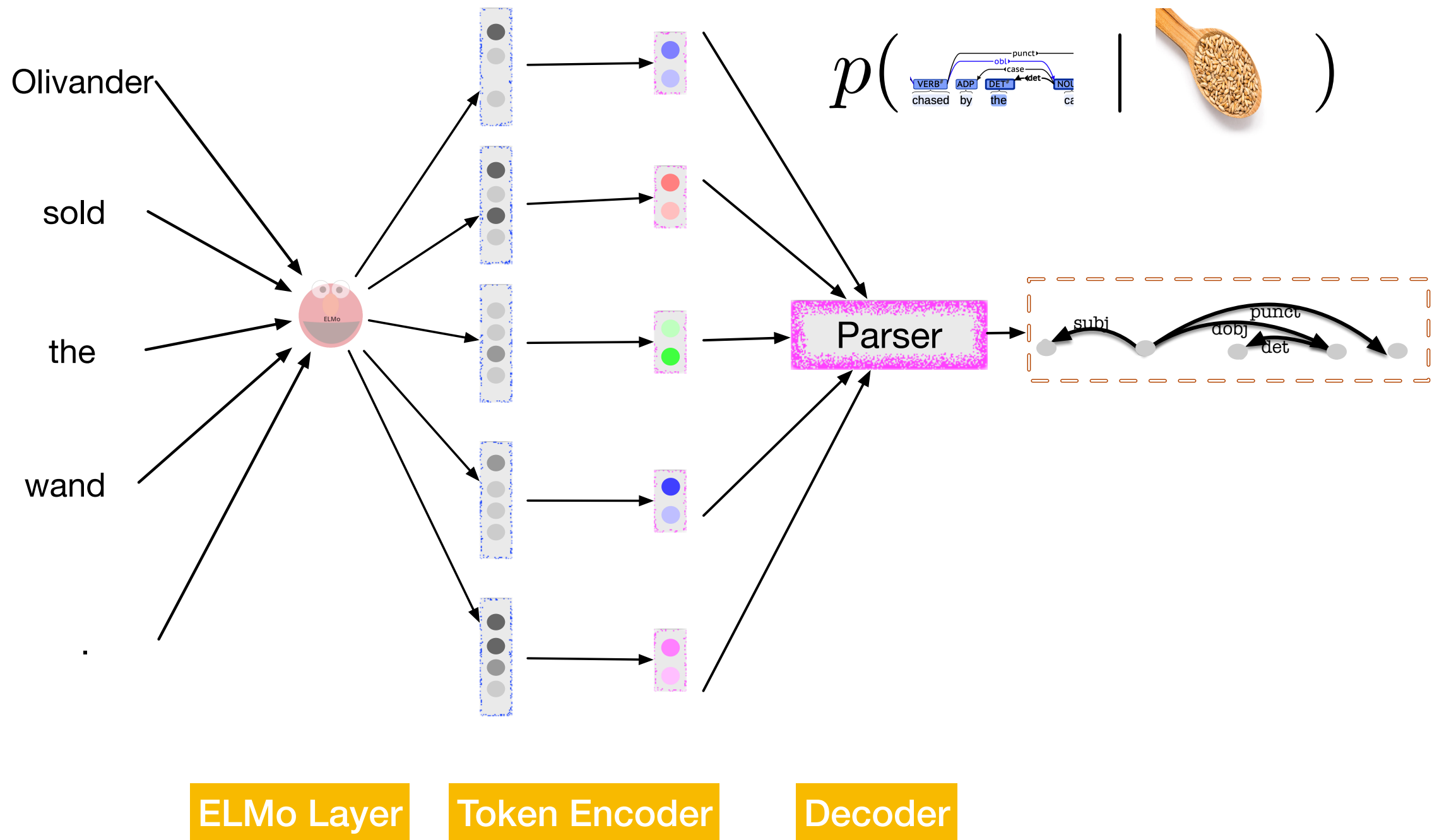
need to know

$$p(\text{chased by the cat.} \mid \text{spoon})$$

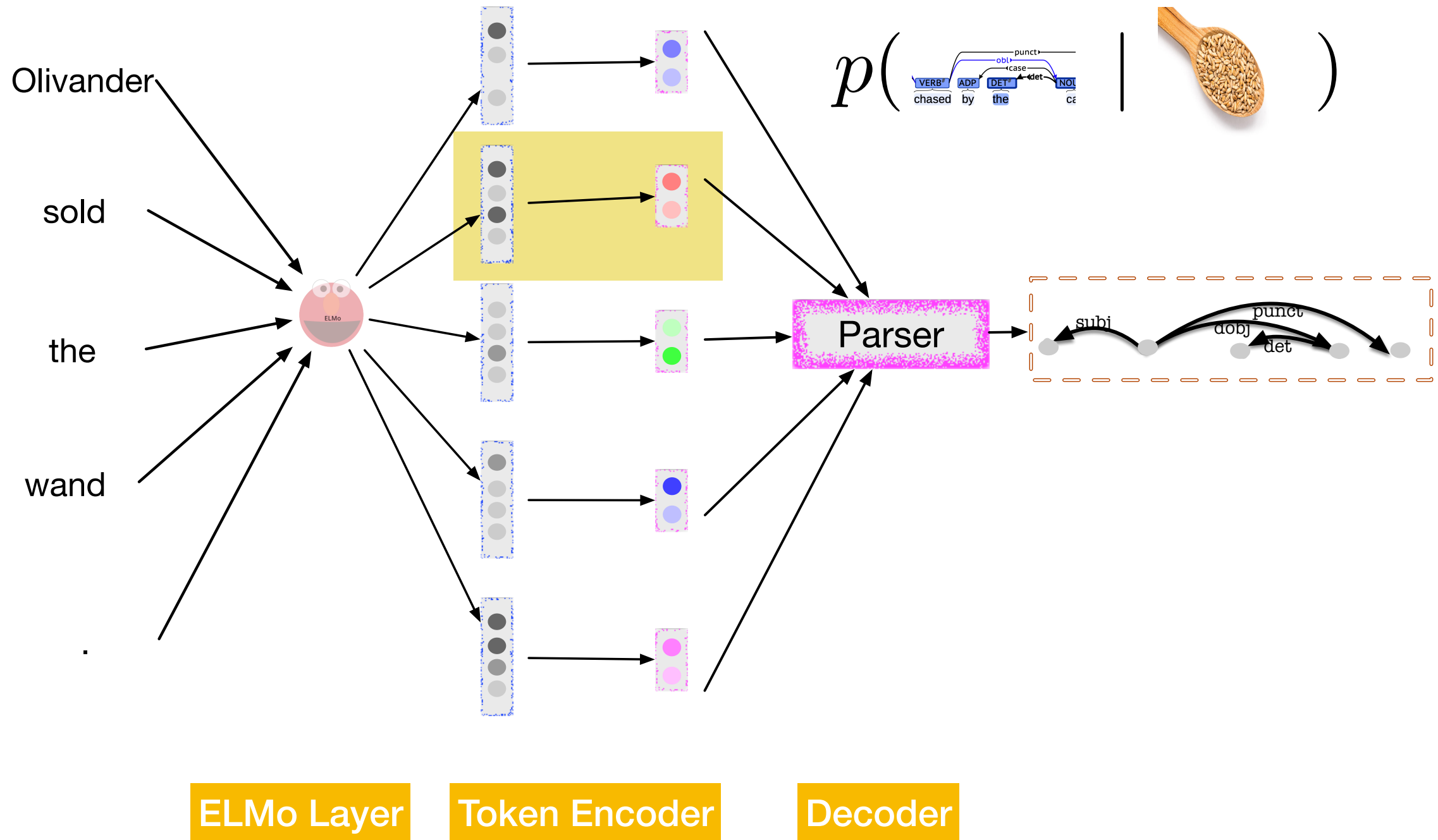
$$KL(p(\text{spoon} \mid \text{Elmo}) \parallel p(\text{spoon}))$$

need to know $p(\text{spoon} \mid \text{Elmo})$ and $p(\text{spoon})$

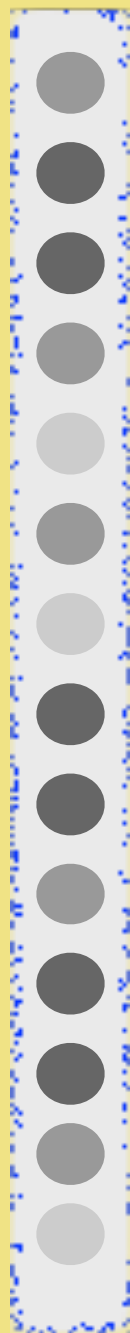
$$p(\text{spoon} \mid \text{Elmo})$$



$$p(\text{spoon} \mid \text{Elmo})$$



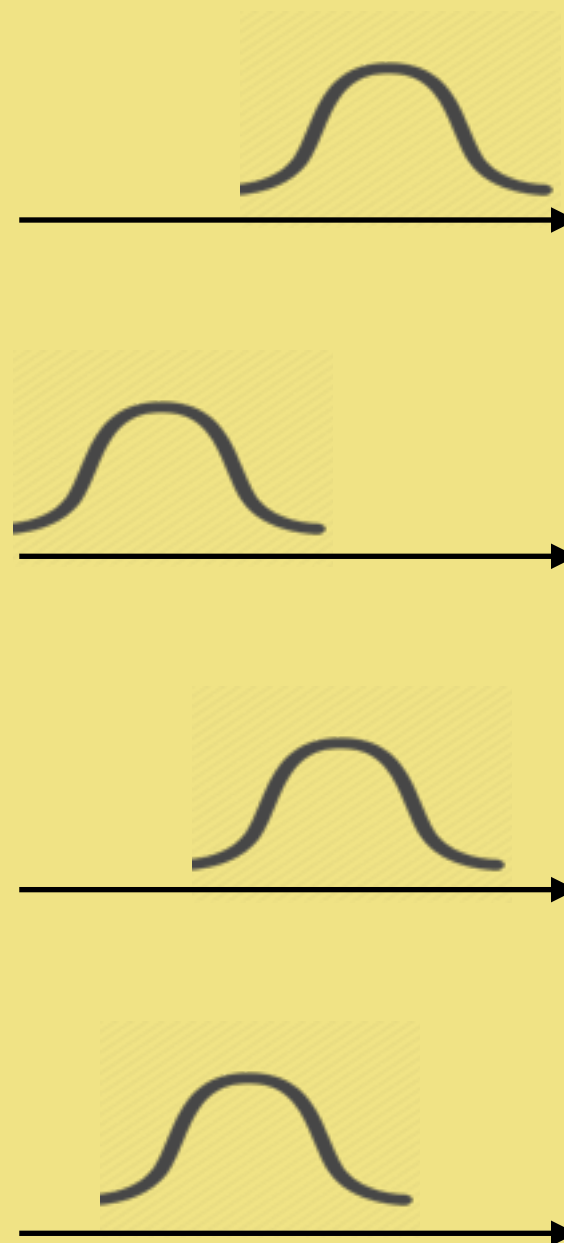
ELMo
Layer-1



CCA
~~**PCA**~~

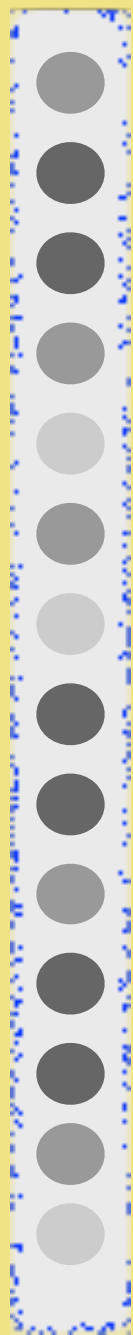


IB
(stochastic)



one word token

**ELMo
Layer-1**

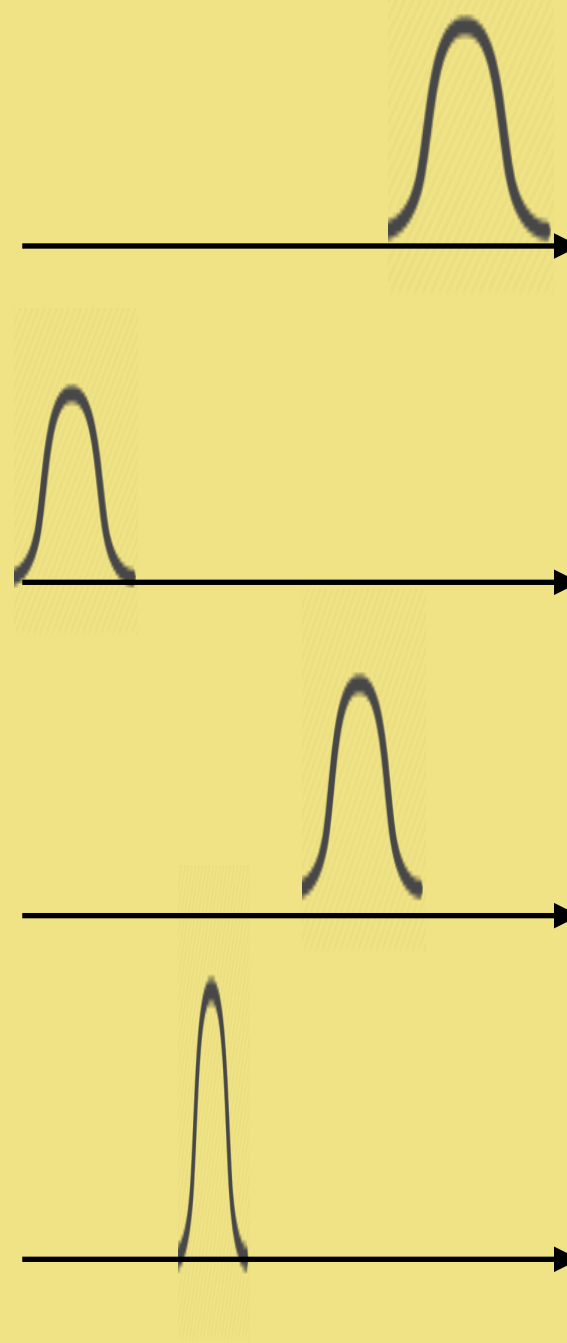


one word token

CCA

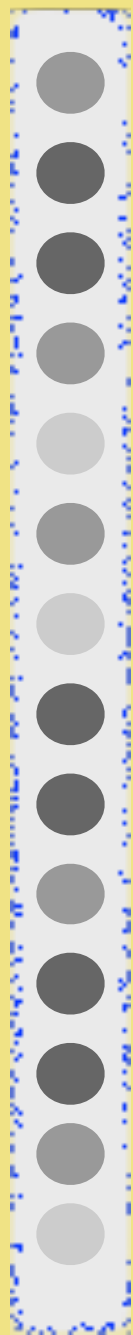


**IB
(stochastic)**



**For some tokens, we
need all 4 dimensions**

**ELMo
Layer-1**

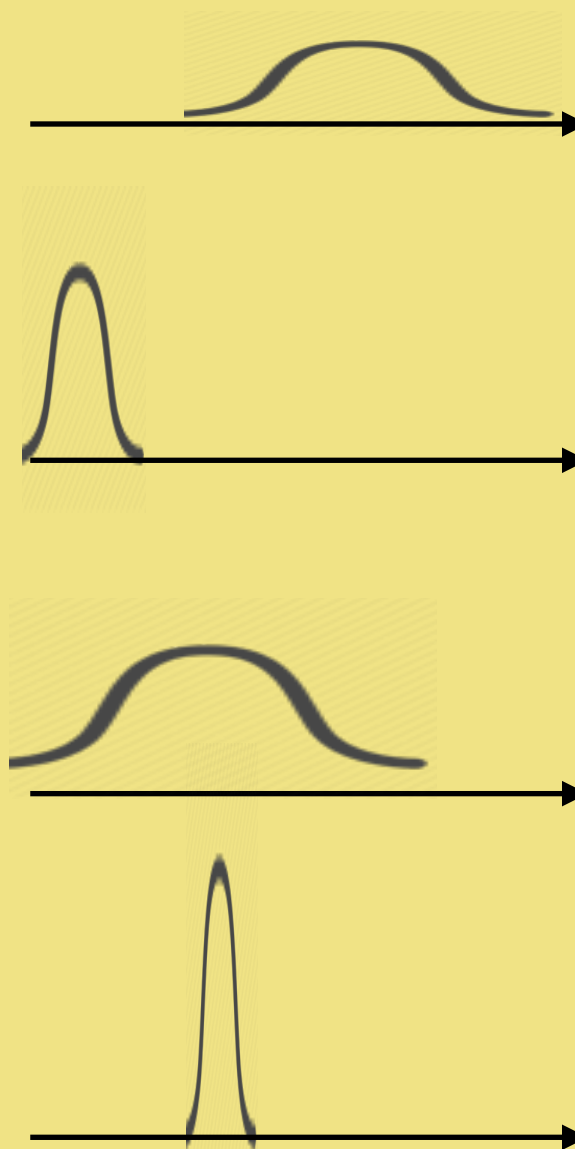


one word token

CCA

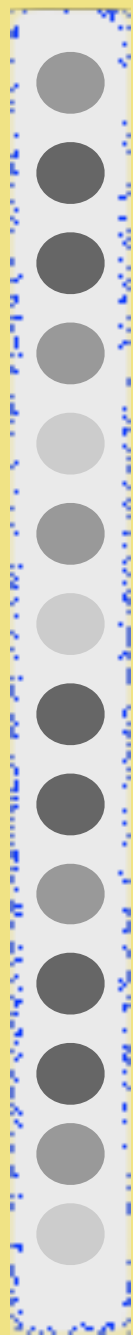


**IB
(stochastic)**



**For some tokens,
retaining less information
does not hurt parsing.**

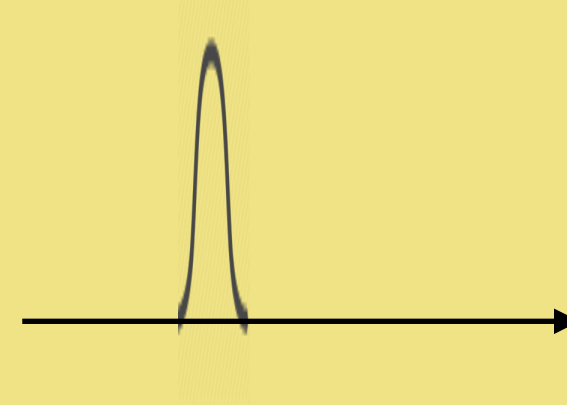
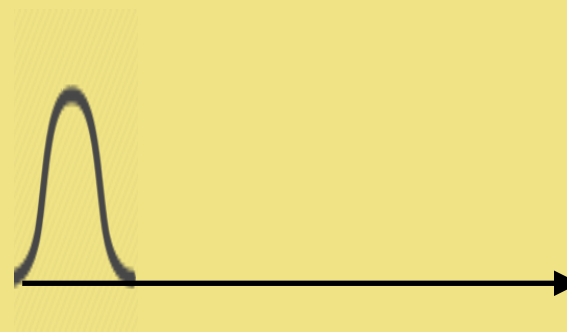
**ELMo
Layer-1**



CCA



**IB
(stochastic)**



number:

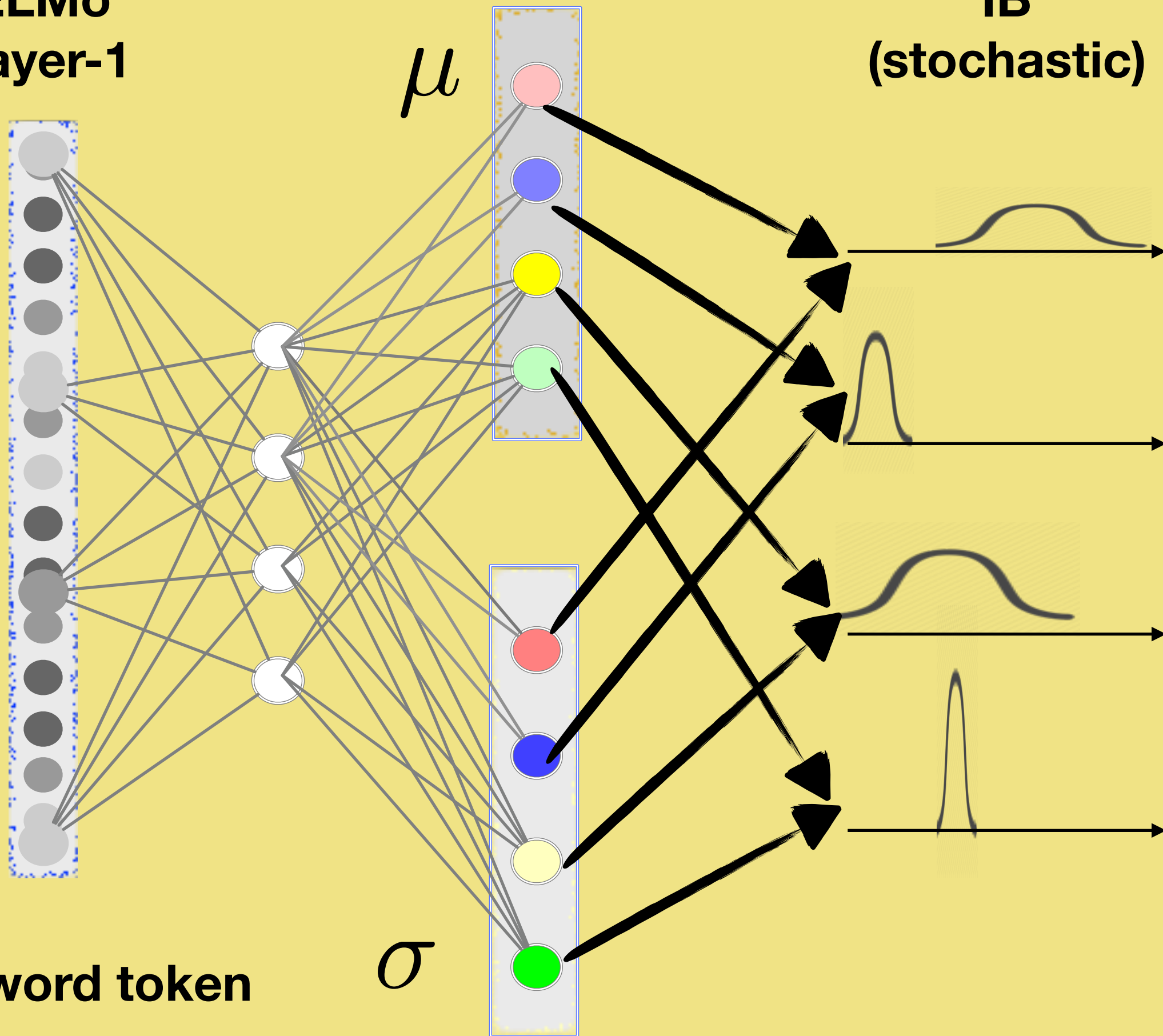
**accusative/
nominative:**

**Example: In English, we don't care whether an
object pronoun is singular or plural.**

ELMo
Layer-1

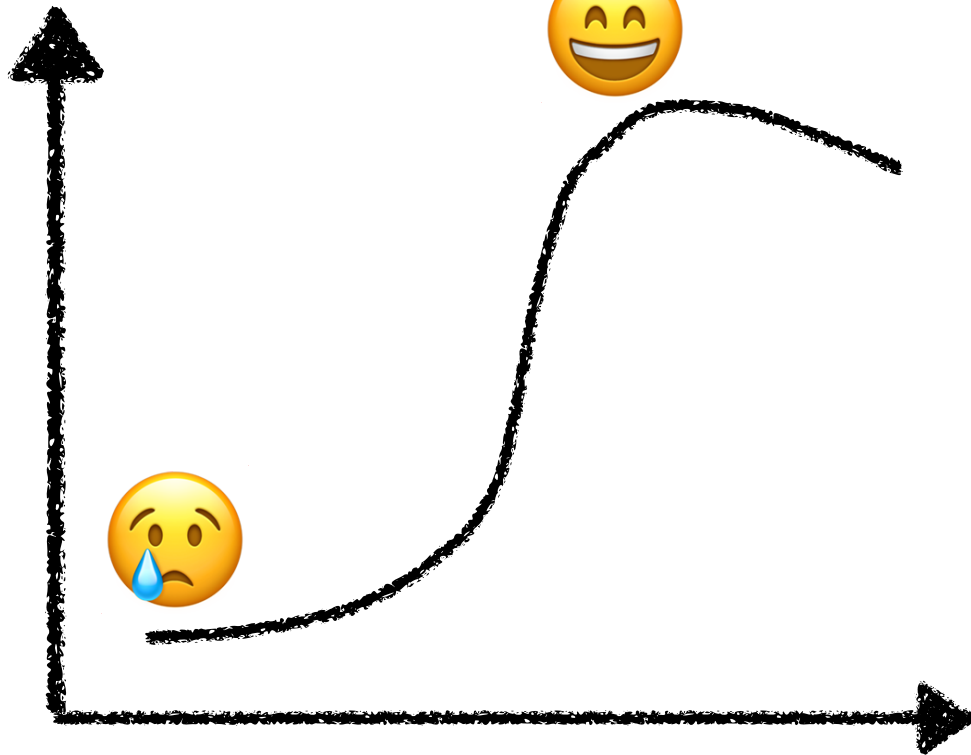
IB
(stochastic)

one word token



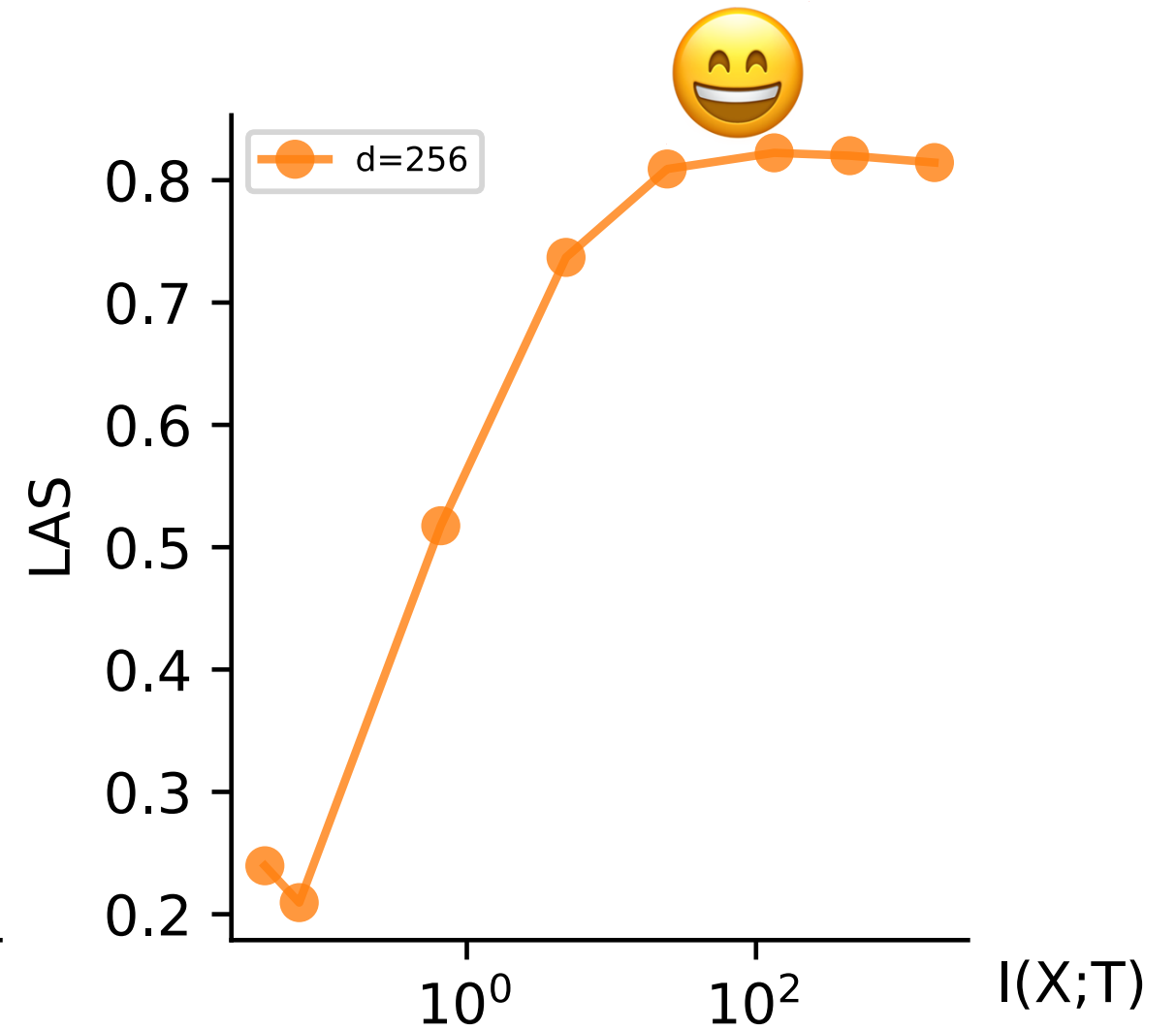
Cartoon

LAS



Information Kept

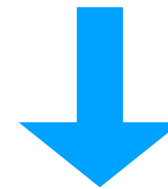
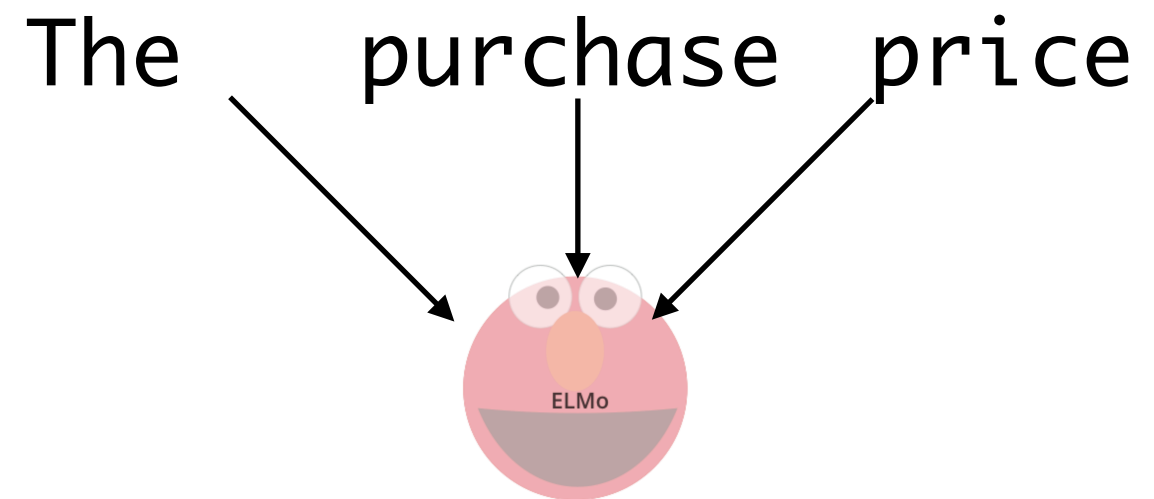
English



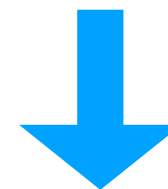
Q: Is our specialization contextual?

A: YES, because we compress ELMo layer 1, which depends on the context.

ELMo Embeddings

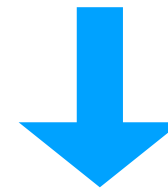
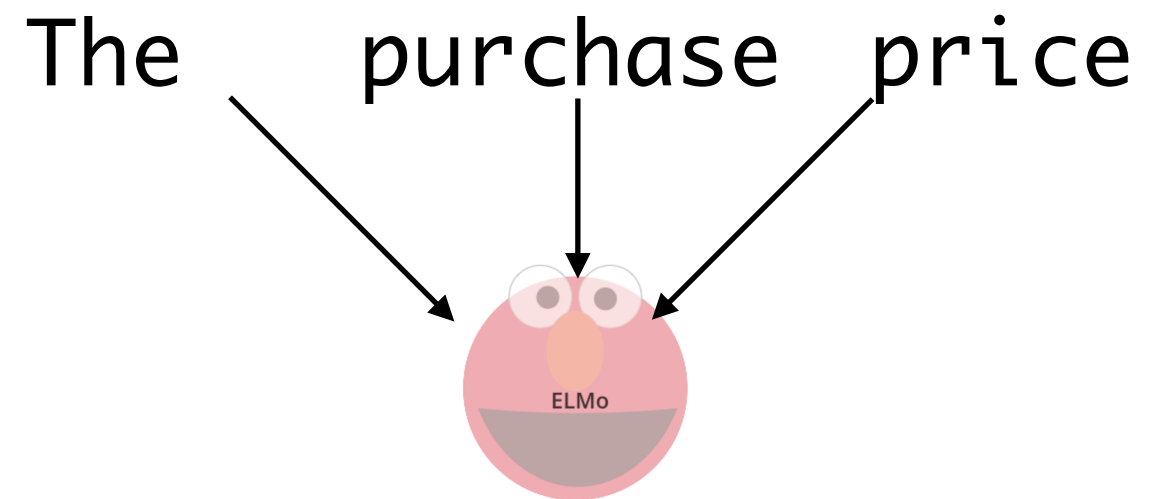


Structured
Prediction

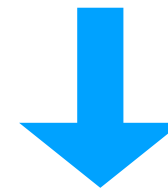


part-of-speech taggings Det Noun Noun

ELMo Embeddings



**Structured
Prediction**



part-of-speech taggings

CCG taggings

NP/NN

NN/N

N

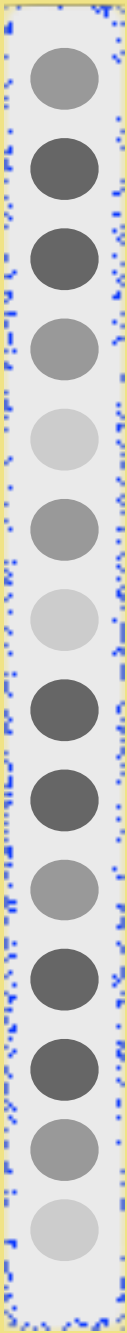
Discrete

Continuous

ELMo
Layer-1

HMM
IB
(stochastic)

PCA
IB
(stochastic)



Tag 1

~~DET~~

0

0.05

Tag 2

~~ADJ~~

0

0.1

Tag 3

~~VERB~~

1

0.5

Tag 4

~~NOUN~~

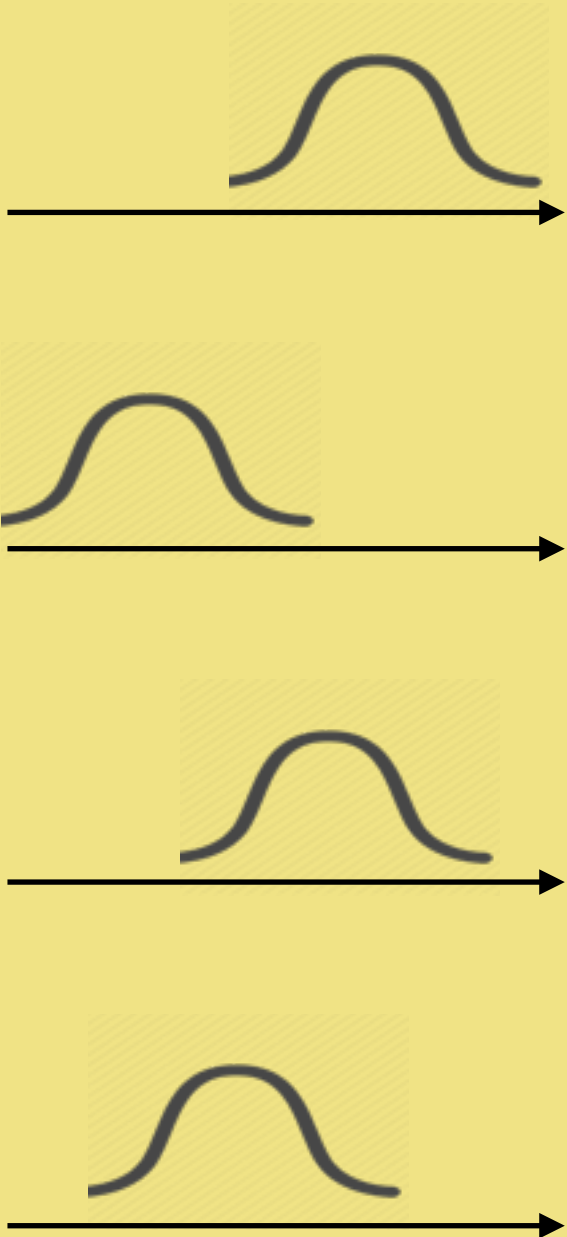
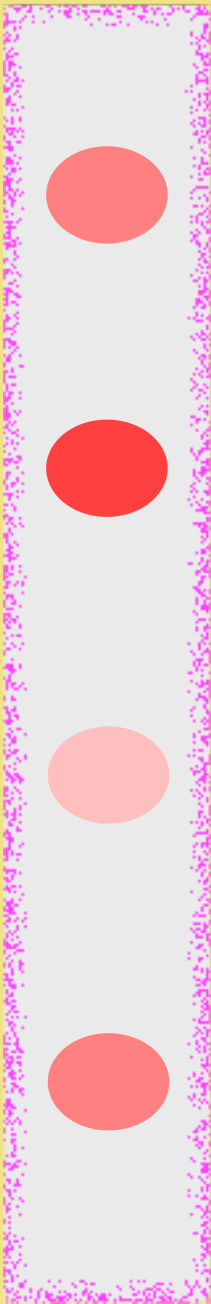
0

0.35

...

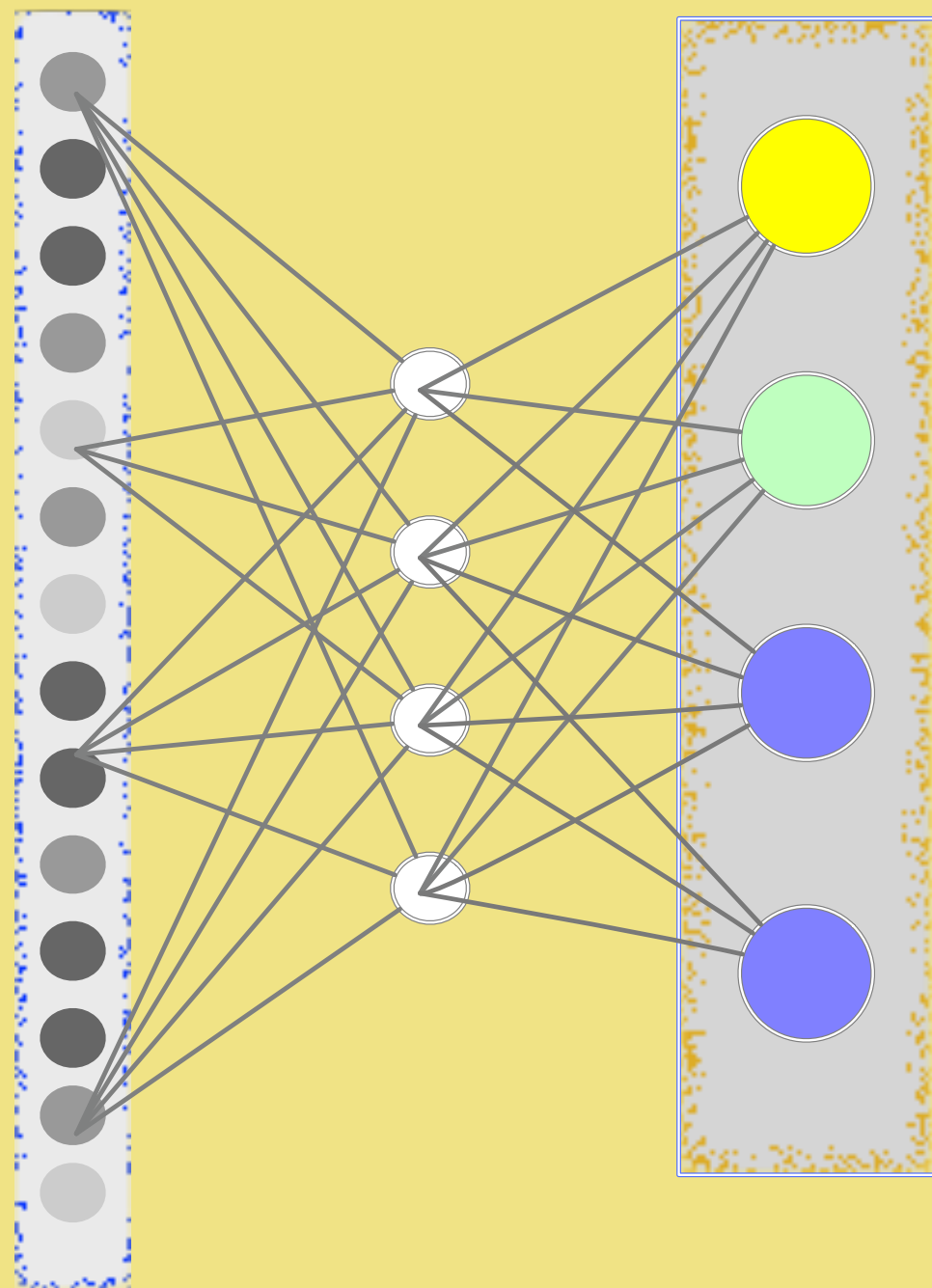
0

0.0



one word token

ELMo Layer-1



one word token

IB (stochastic)

Tag 1 0.05

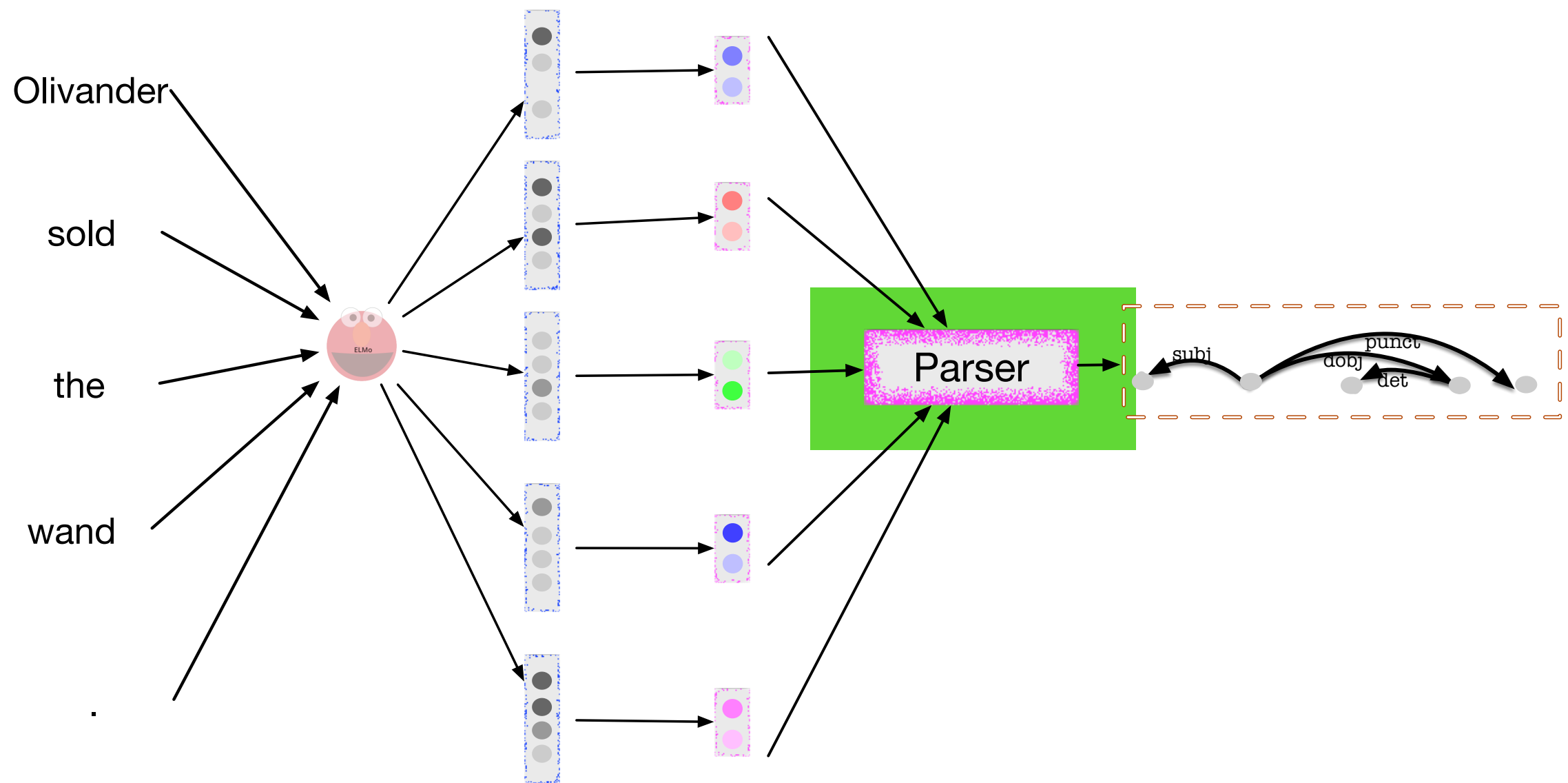
Tag 2 0.1

Tag 3 0.7

Tag 4 0.15

Softmax





ELMo Layer

Token Encoder

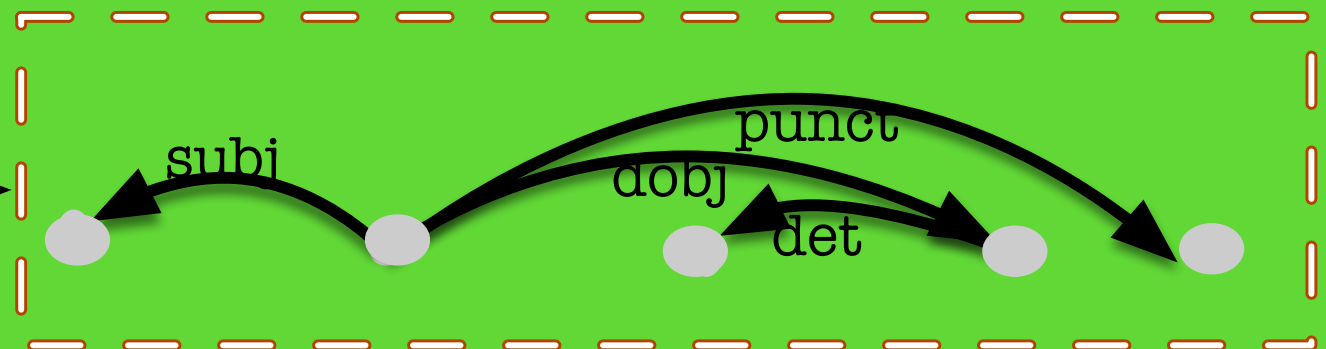
Decoder

$$p(\text{spoon of seeds} \mid \text{Elmo})$$

$$p(\text{chased by the cat} \mid \text{spoon of seeds})$$

Decoder

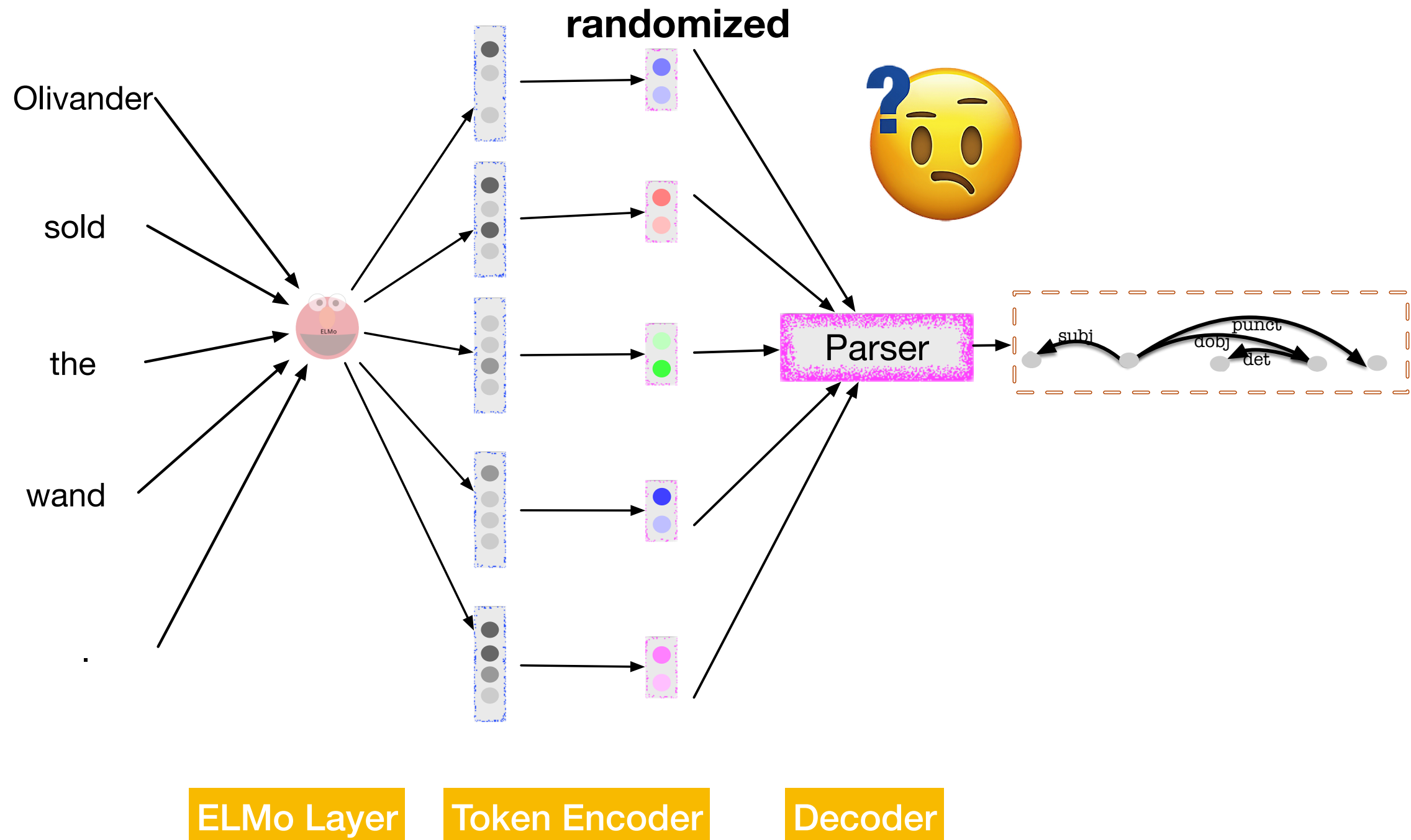
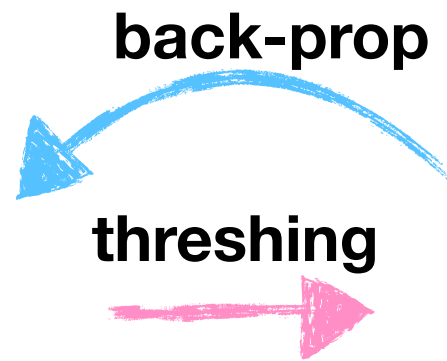
Dependency Parsing



the deep biaffine dependency parser (Dozat and Manning, 2016)

reparametrization trick
or Gumbel-softmax trick

Sampling



Variational Upper Bound

- To minimize the IB objective directly, some terms are intractable.

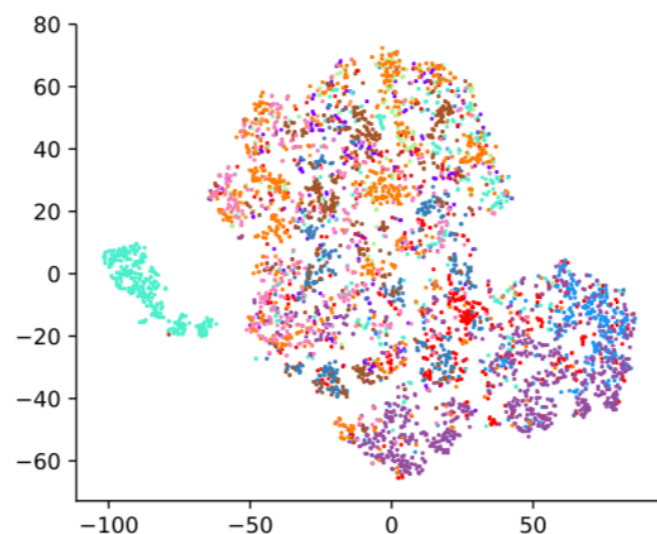
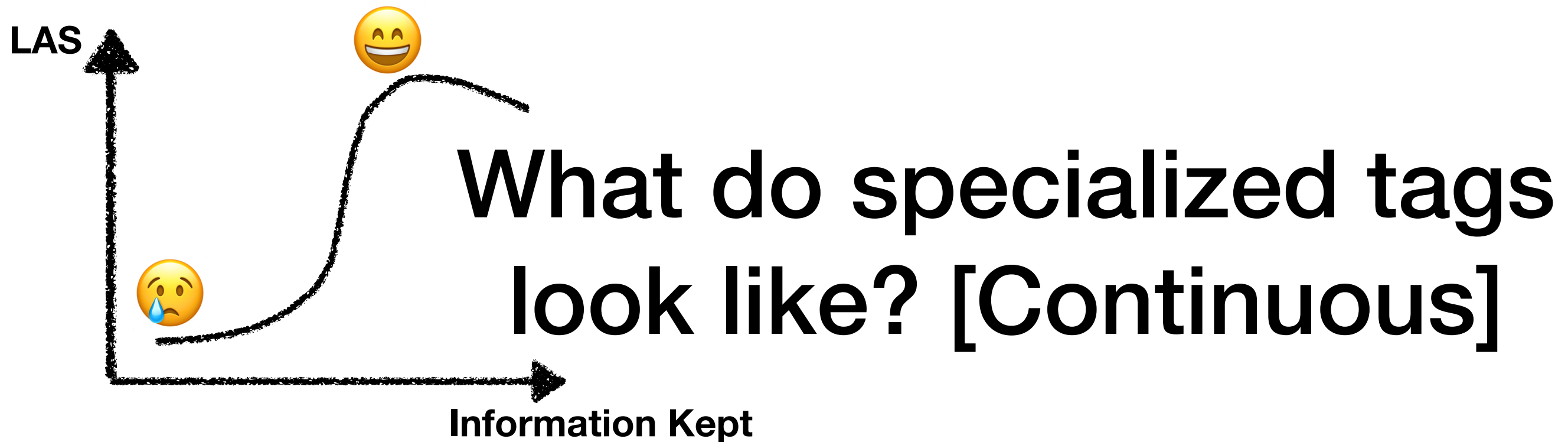
- $p(\text{chased by the cat} \mid \text{img})$ and $p(\text{img})$

- So we instead minimize a variational upper bound.
 - Which is what the previous slide estimated by sampling
- For mathematical details, please refer to our paper ;)

Do our specialized tags correlate with traditional POS tags?

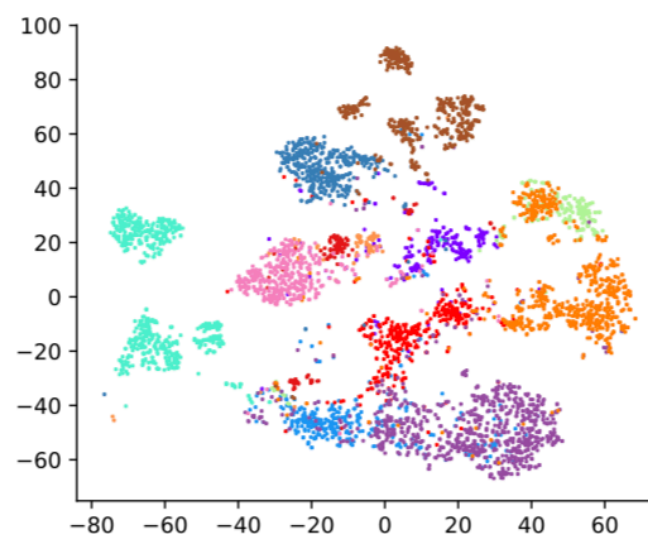
Yes!

Results



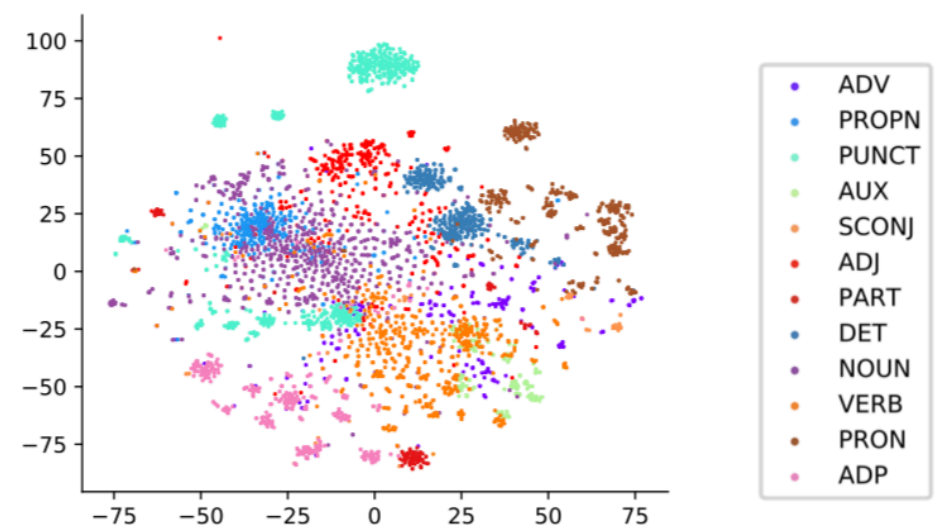
(c) $I(X;T) \approx 0.069$

**Too much
Compression**



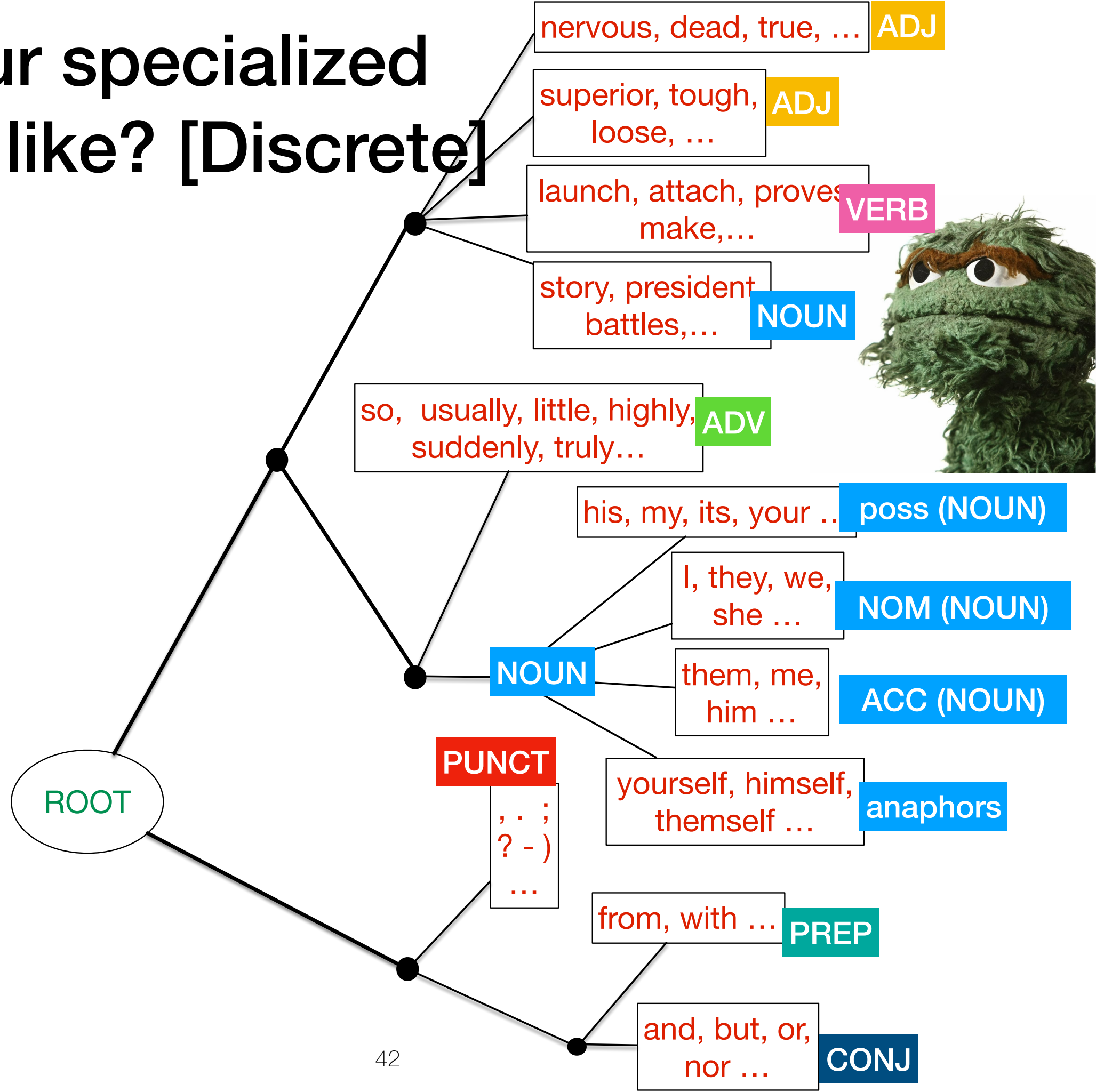
(b) $I(X;T) \approx 24.3$

**Moderate
Compression**



(a) ELMo, $I(X;T) = H(X) \approx 400.6$

**All Information
Kept**



Do our specialized tags correlate with POS tags?

Yes!

- Our results agree with the intuition that POS keeps information about parsing.
- POS is contextual. Our taggings are contextual as well !!!

WAIT!? The word type (without context) should also be a strong predictor of POS.

How do we specialize so that we mostly depend on type information?

1. We could use ELMo Layer-0.

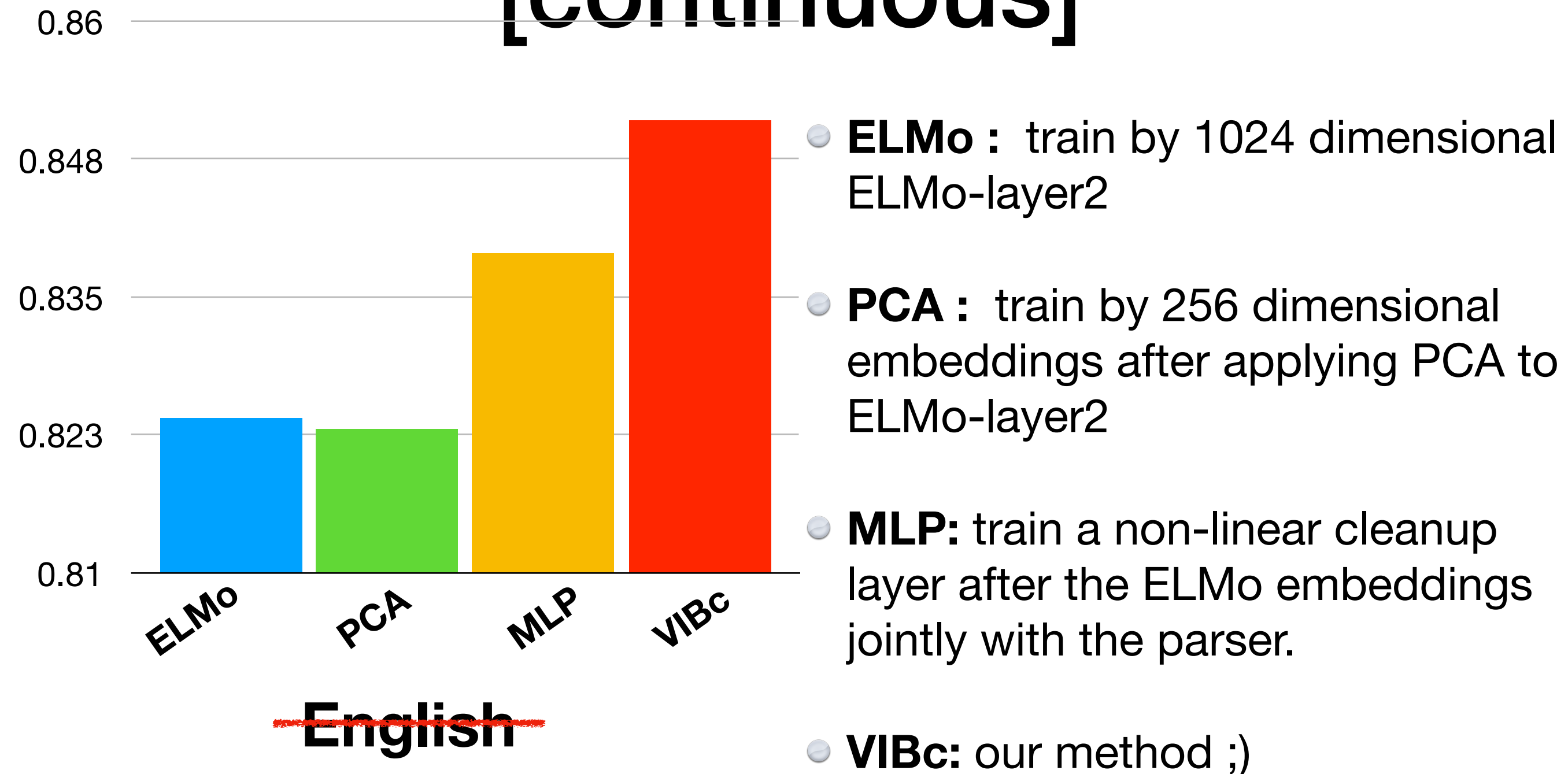
ELMo layer-0 is based on a character-level convolutional network. Thus, it inherently does not contain contextual information. Specializing is to extract from existing embeddings, which is also non-contextual.

2. We could do a softer version — make the specialized tagging depend “mostly” on its word type.

$$\mathcal{L} = -I(\text{chased by the cat .}; \text{spoon}) + \beta \cdot I(\text{Elmo}; \text{spoon}) + \gamma \cdot \sum_{\text{word } i \text{ in sentence}} I(\text{Elmo}_i; \text{spoon}_i \mid \text{layer-0, } i)$$

Parsing Performance [continuous]

LAS



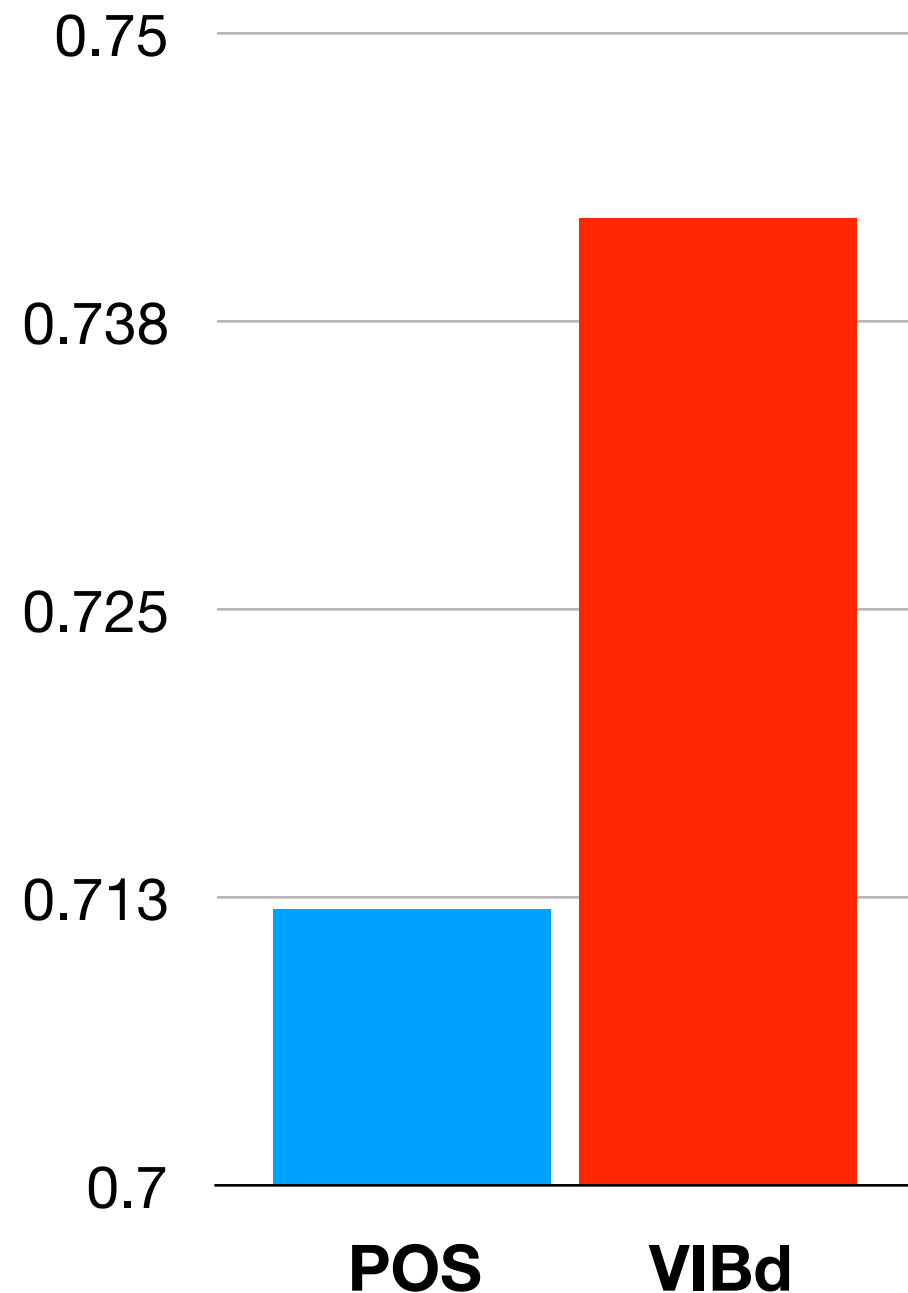
~~English~~

Wins or ties across 9 languages.

≈ 1.0 point improvement on average

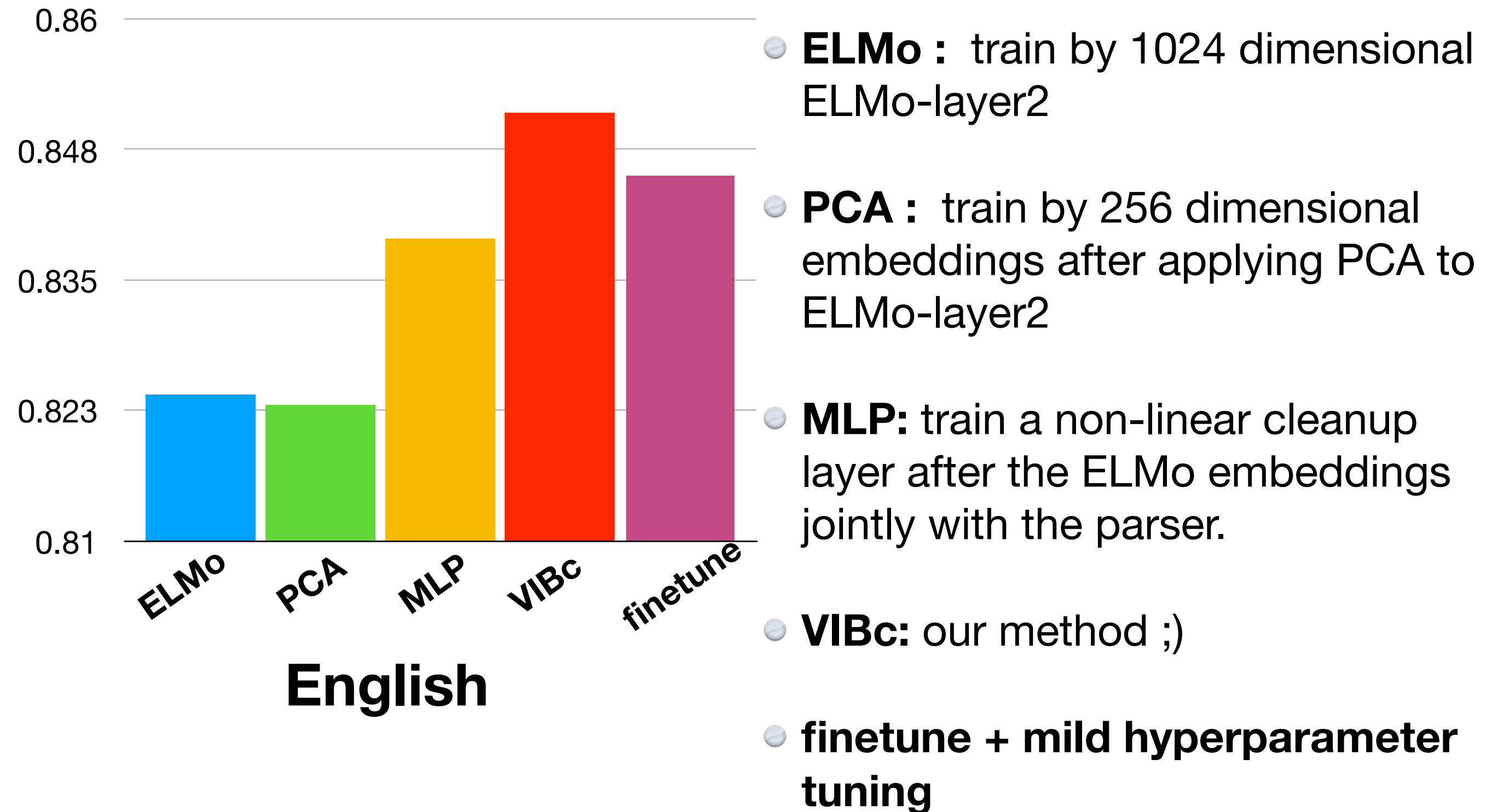
Parsing Performance [discrete]

LAS



- POS: train the parser on gold POS tags
- VIBd: Our discrete version.

Parsing Performance [continuous]





Thanks