

Лекция 4.

Seq2seq (Sequence-to-Sequence) и Attention

Seq2seq (Sequence-to-Sequence) и Attention

- Seq2seq
- Attention
- Transformer
- Сегментация по подсловам: BPE
- Анализ и интерпретируемость

Задача перевода текста

- Человеческий перевод:

вероятность $y^* = \arg \max_y p(y|x)$ интуитивна и определяется опытом человека.

- Машинный перевод:

вероятность $y' = \arg \max_y p(y|x, \theta)$, где θ — параметр модели.

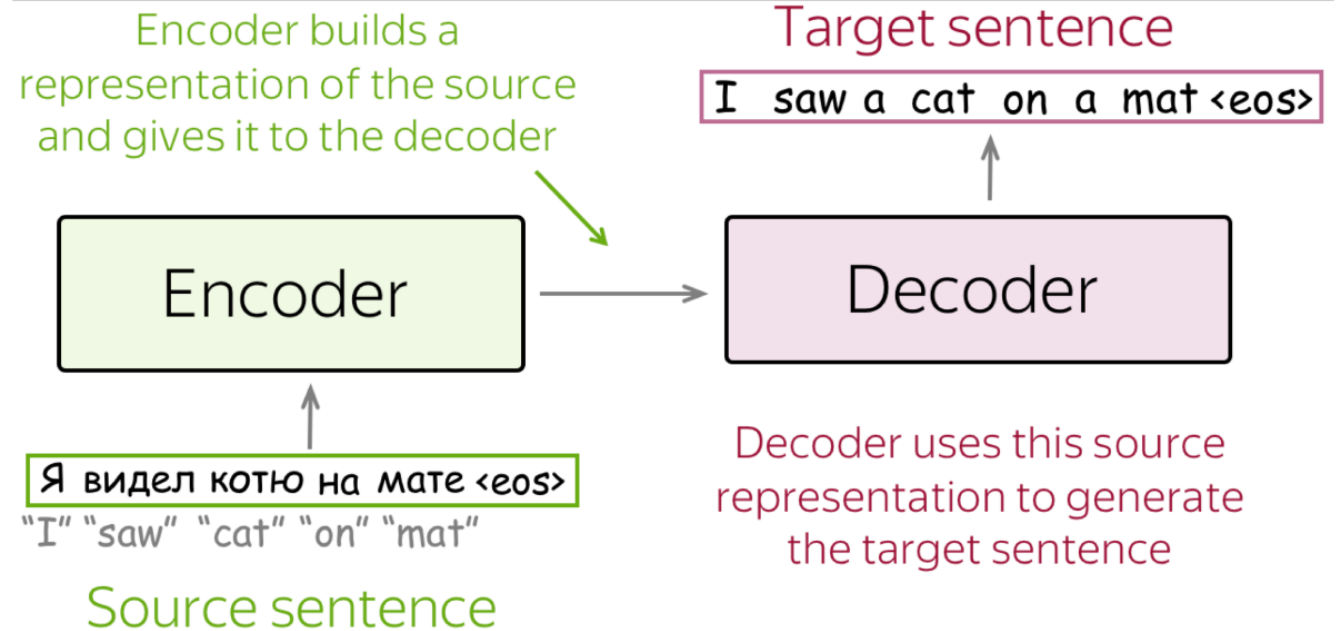
Возникают вопросы **моделирования** (как выглядит модель), **обучения** (как найти θ) и **поиска** (как найти $\arg \max$).

- modeling

How does the model for $p(y|x, \theta)$ look like?

Архитектура «кодировщик-декодировщик» (encoder-decoder)

- Кодировщик – считывает исходное предложение и создаёт его представление
- Декодировщик – использует исходное представление от кодировщика для генерации целевой последовательности.



Модели условного языка

Языковая модель: $P(y_1, y_2, \dots, y_n) = \prod_{t=1}^n p(y_t | y_{<t})$

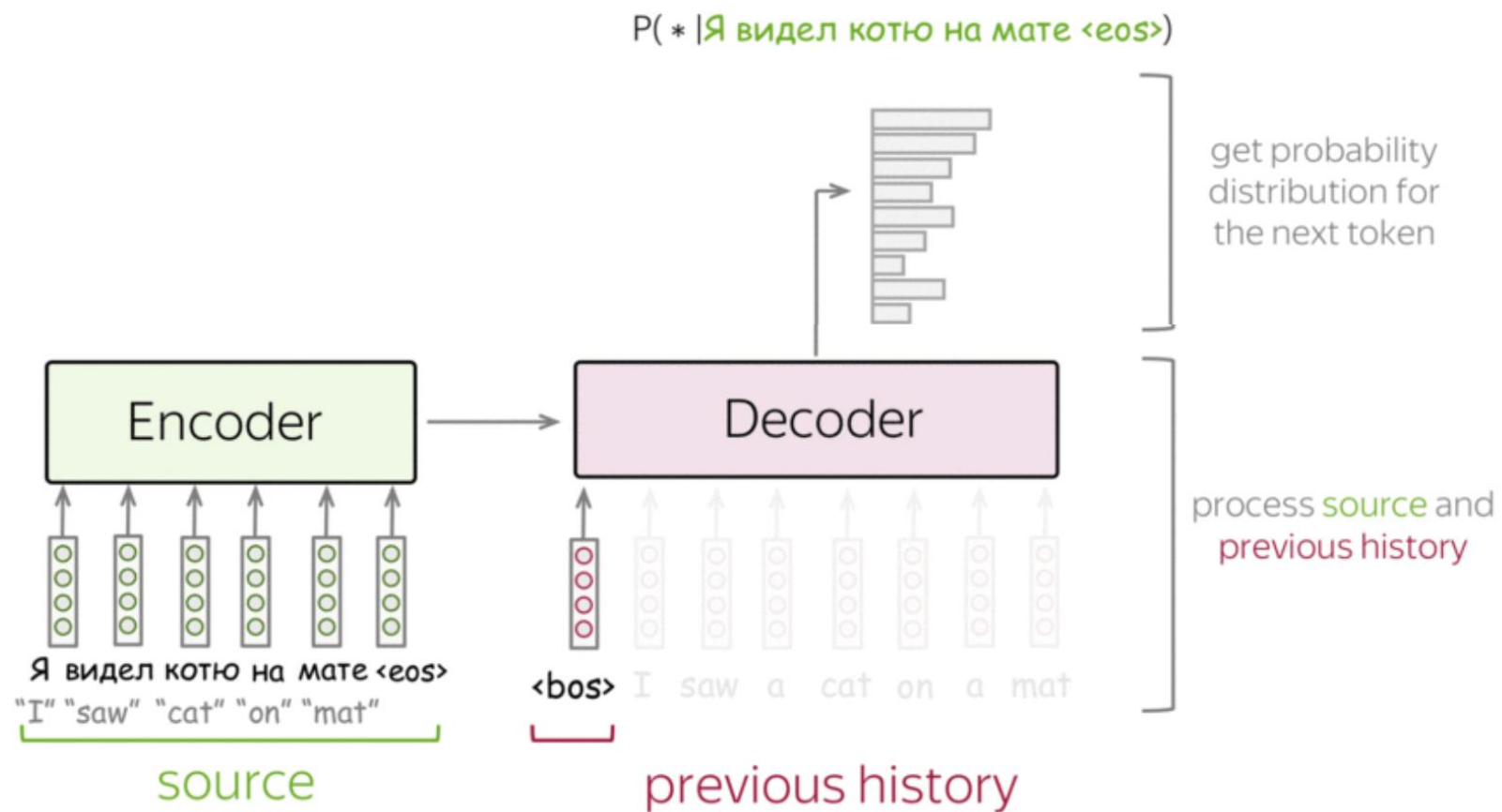
Условная языковая модель: $P(y_1, y_2, \dots, y_n, |x) = \prod_{t=1}^n p(y_t | y_{<t}, x)$
condition on source x

Общий вид

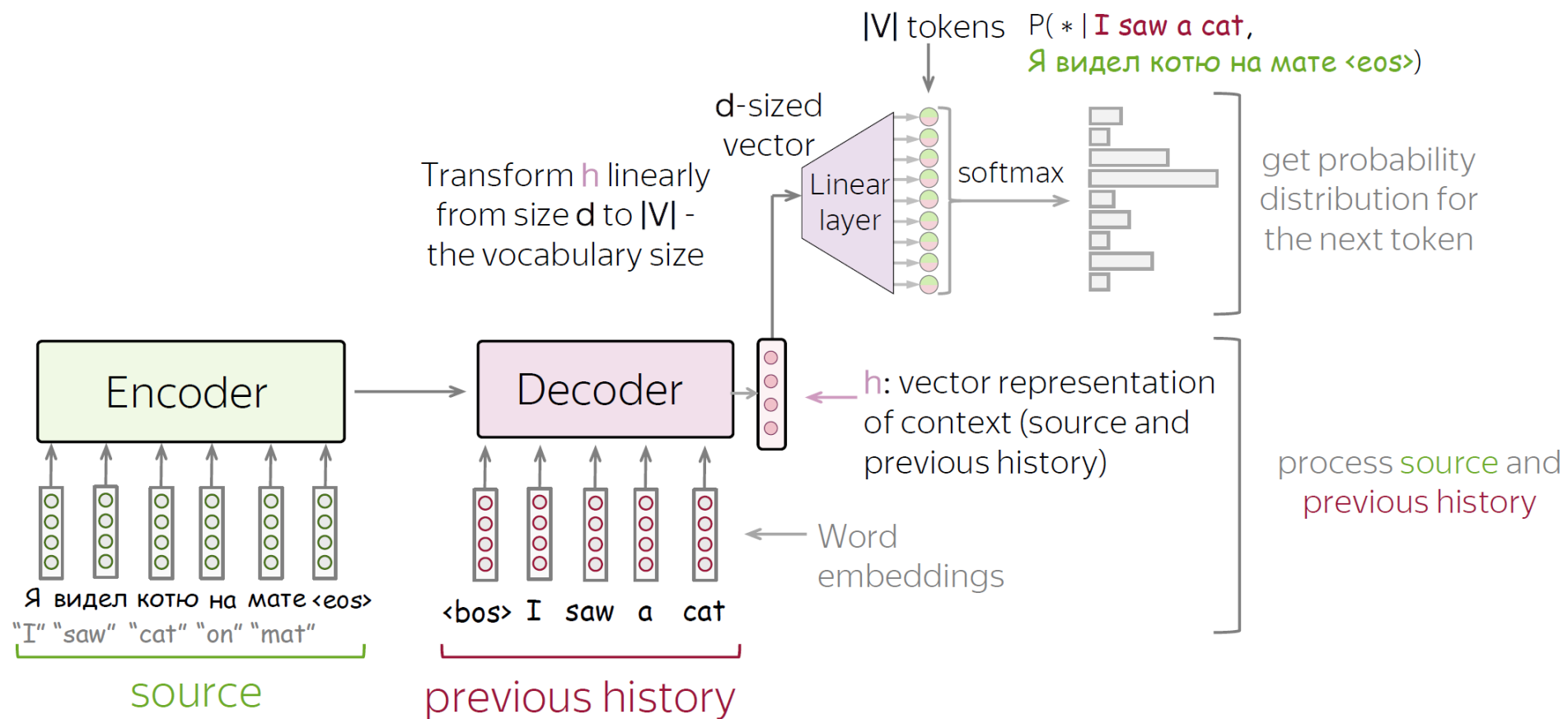
Общий процесс имеет вид:

- источник данных и ранее сгенерированные слова передаются в сеть;
- получаем векторное представление контекста (как исходного, так и предыдущего);
- на основе этого векторного представления прогнозируем распределение вероятностей для следующего токена.

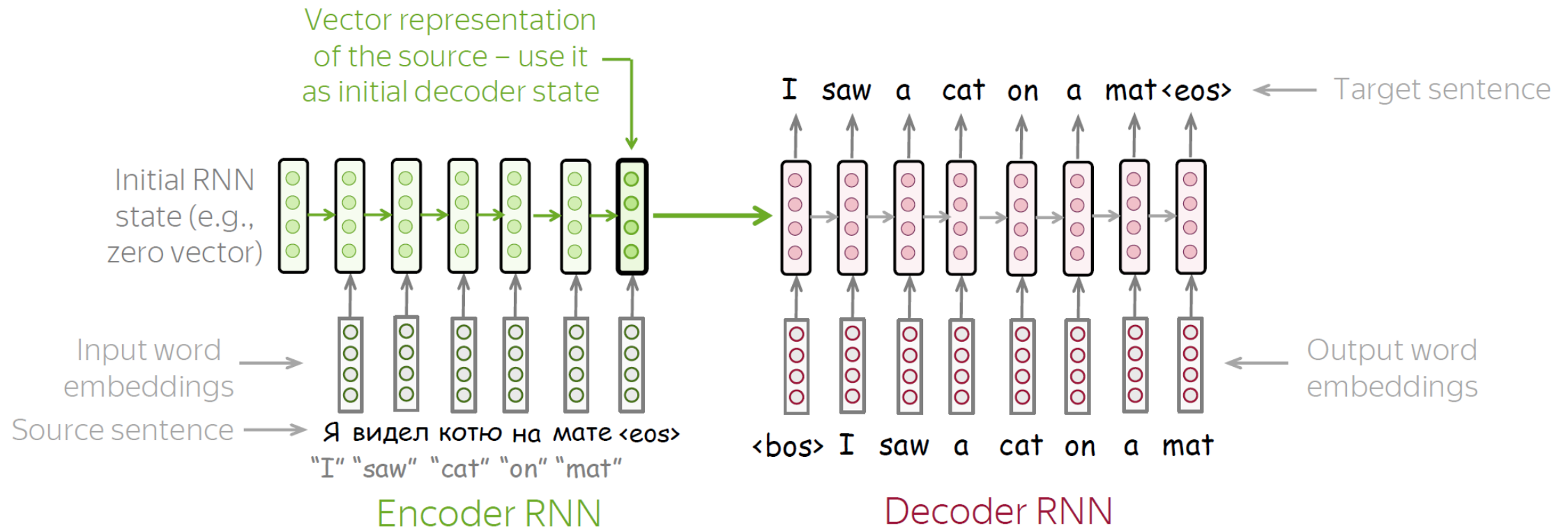
Общий вид



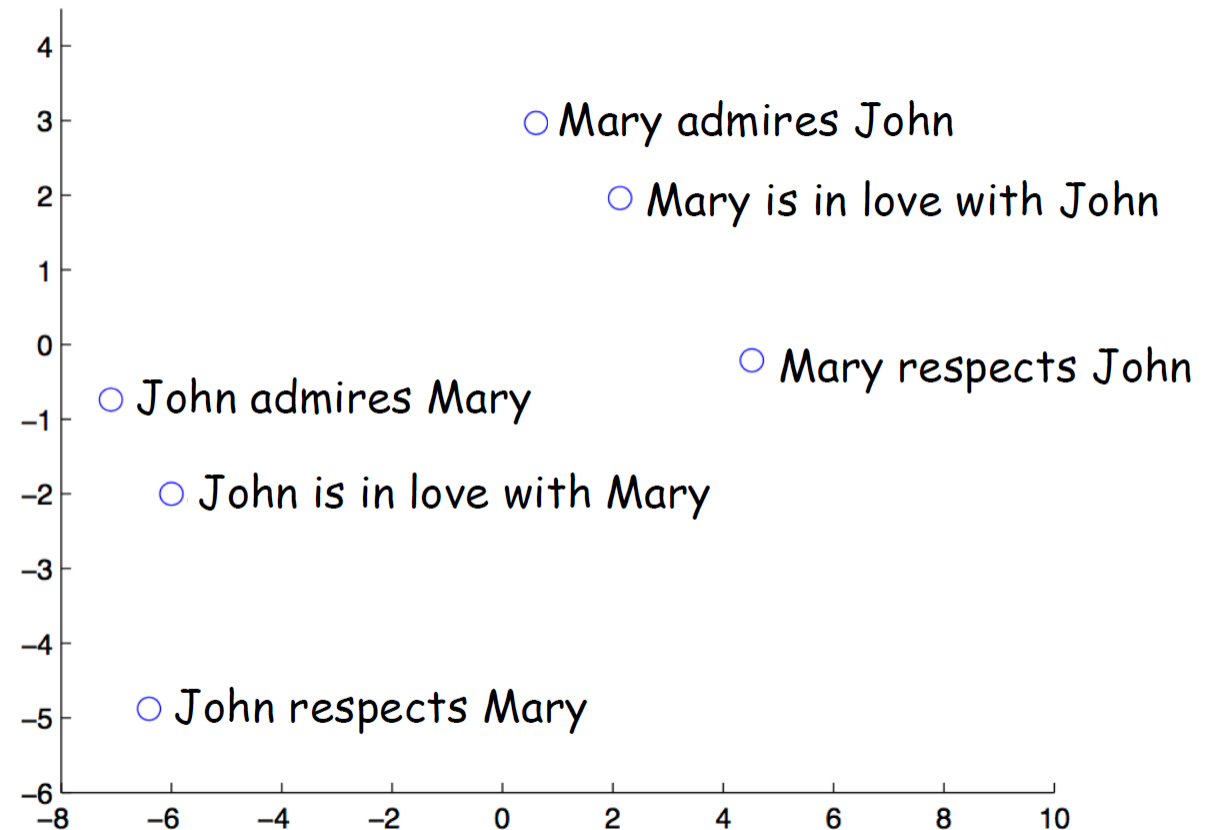
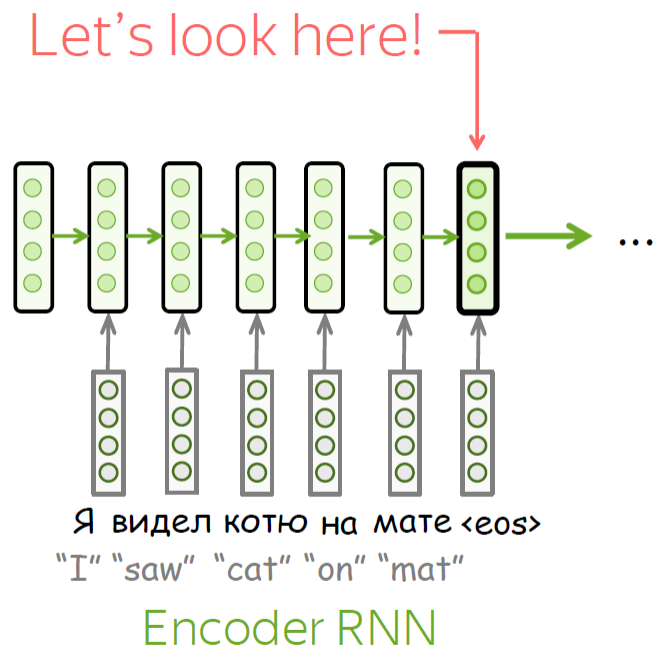
Высокоуровневая схема



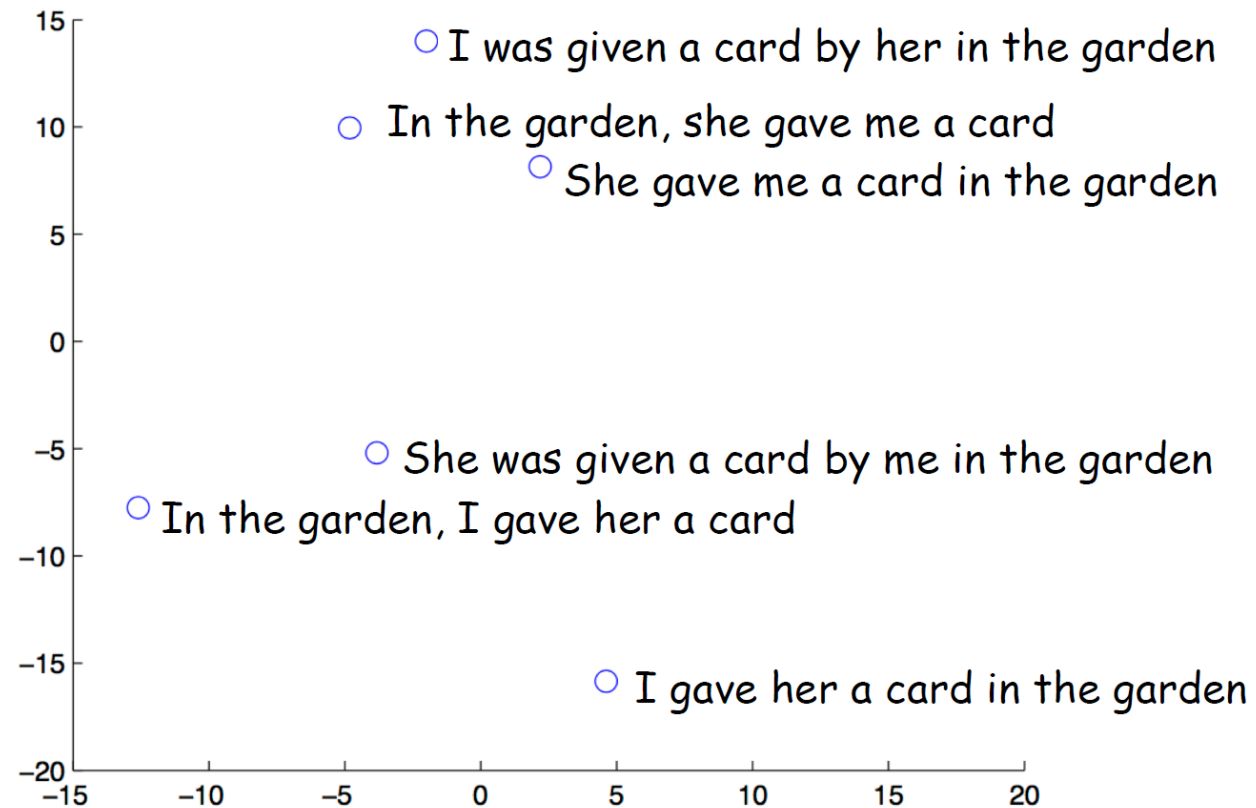
Простейшая модель: RNN-кодер и декодер



Конечное состояние кодировщика



Конечное состояние кодировщика



- learning

How to find θ ?

Кросс-энтропия

Source sequence:

Я видел котю на мате <eos>

"I" "saw" "cat" "on" "mat"

Target sequence:

I saw a cat on a mat <eos>

previous tokens

we want the model
to predict this

← one training example

← one step for this example

Model prediction: $p(* | \text{I saw a, Я ... <eos>})$



cat

Target



Loss = $-\log(p(\text{cat})) \rightarrow \min$



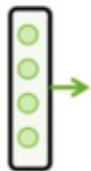
decrease

increase

decrease

Кросс-энтропия

Encoder: read source

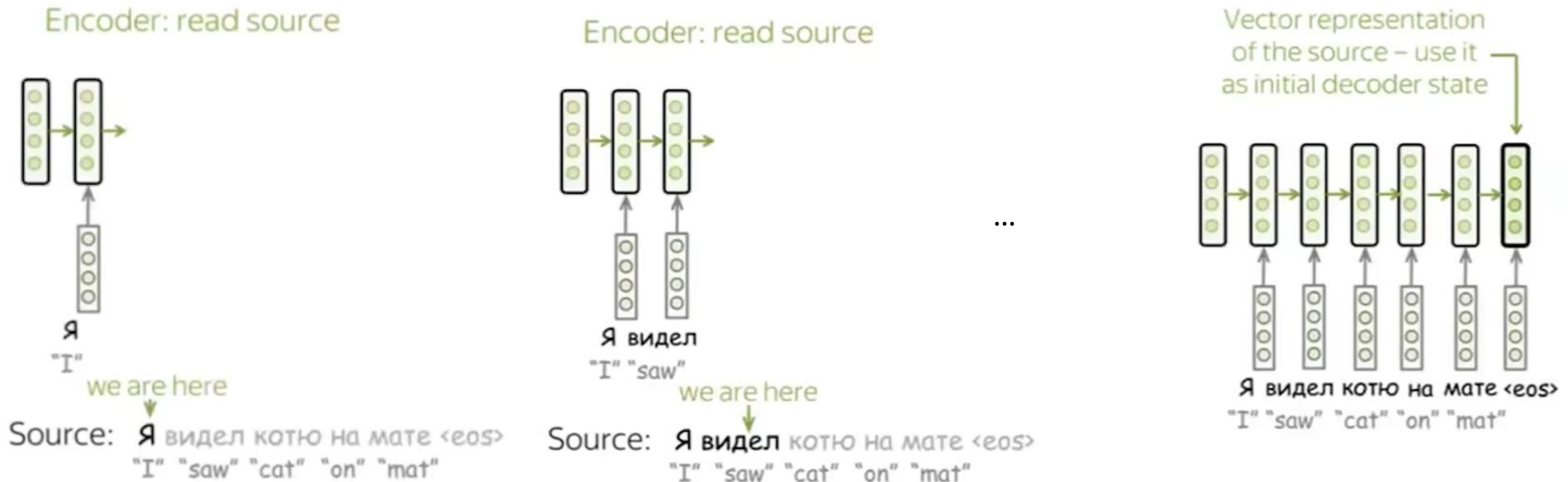


we are here

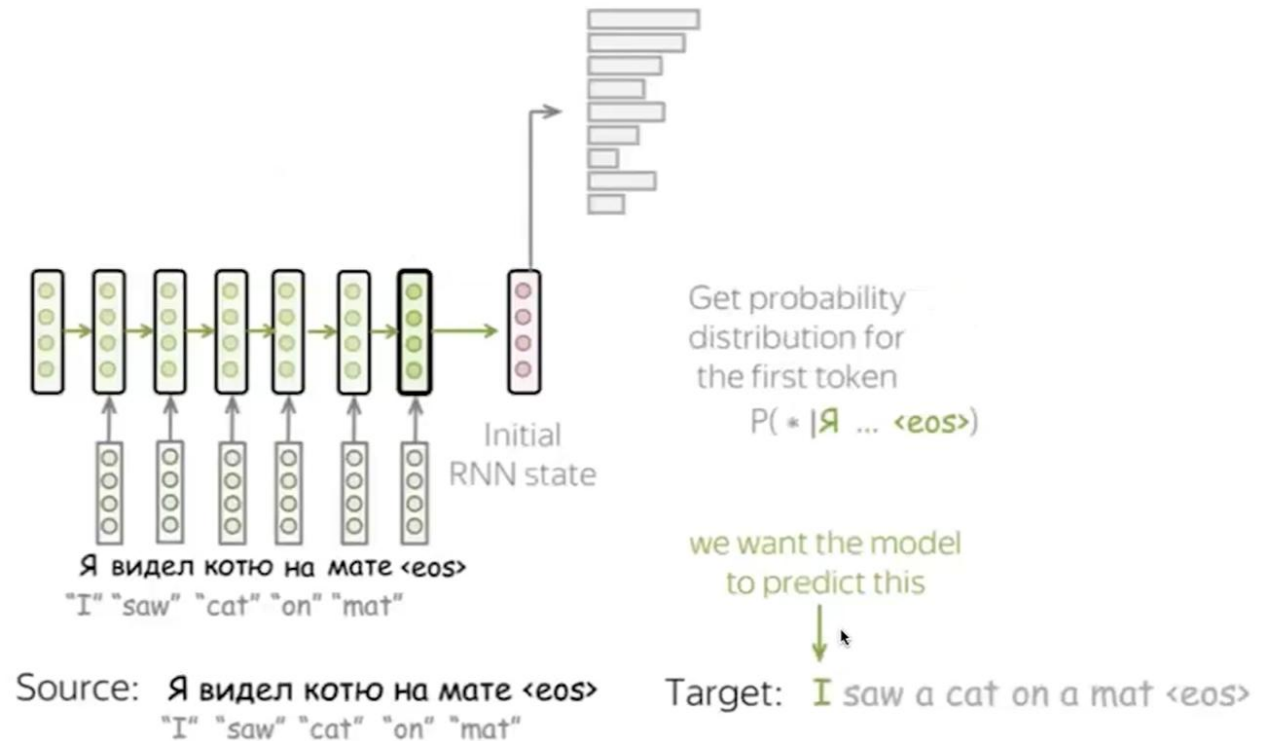
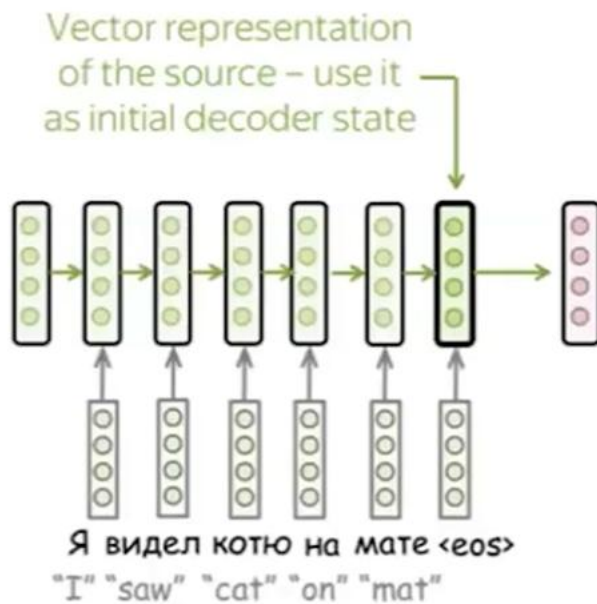
Source: Я видел котю на мате <eos>
"I" "saw" "cat" "on" "mat"

Target: I saw a cat on a mat <eos>

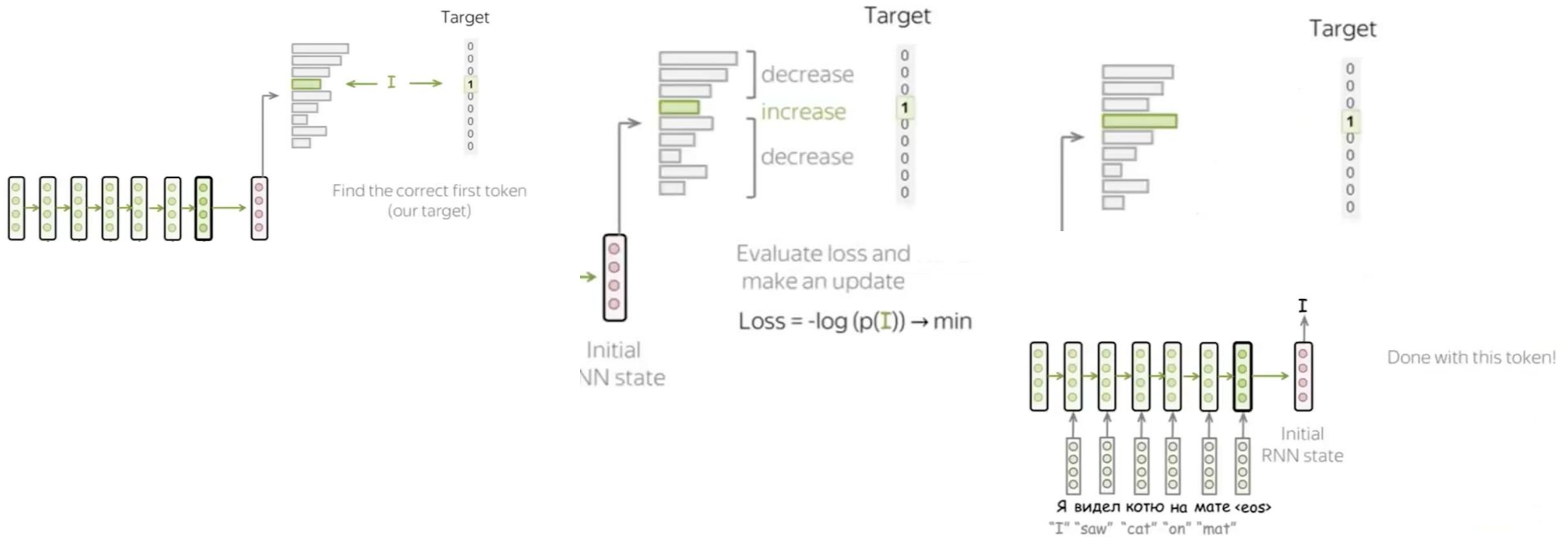
Кросс-энтропия



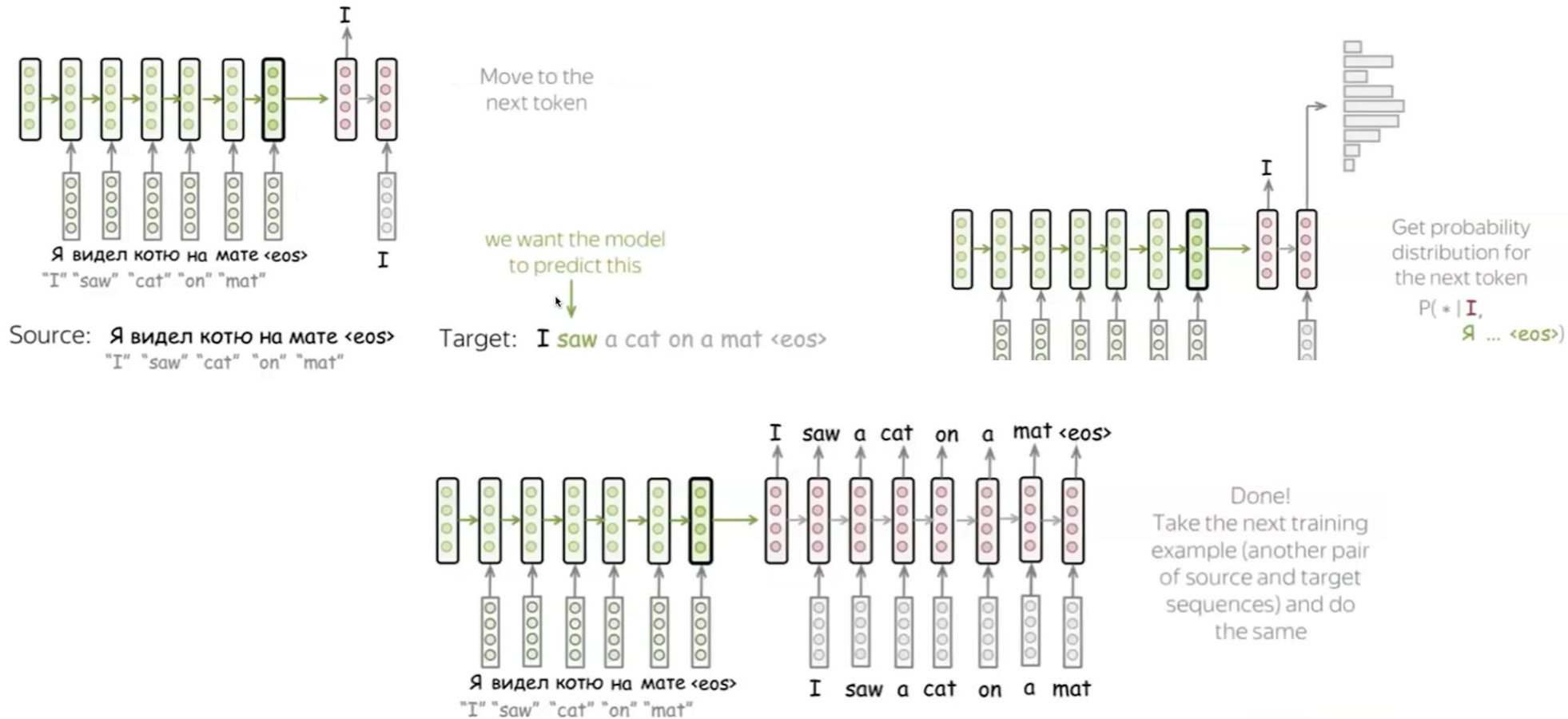
Кросс-энтропия



Кросс-энтропия



Кросс-энтропия



- search

How to find
the argmax?

Как создать перевод?

$$y' = \arg \max_y p(y|x) = \arg \max_y \prod_{t=1}^n p(y_t|y_{<t}, x)$$

Простой вариант:

- жадное декодирование — на каждом шаге выбираем токен с наибольшей вероятностью

НО

$$\arg \max_y \prod_{t=1}^n p(y_t|y_{<t}, x) \neq \prod_{t=1}^n \arg \max_{y_t} p(y_t|y_{<t}, x)$$

Beam Search (лучевой поиск)

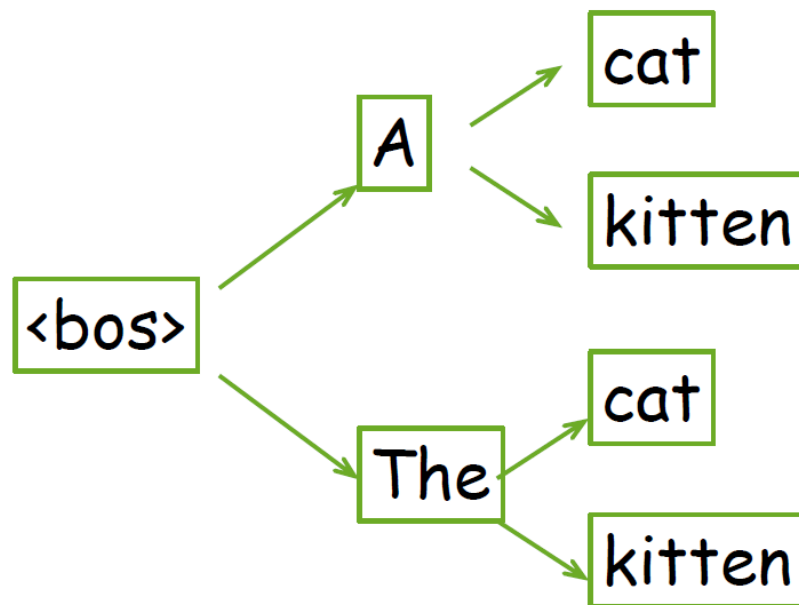
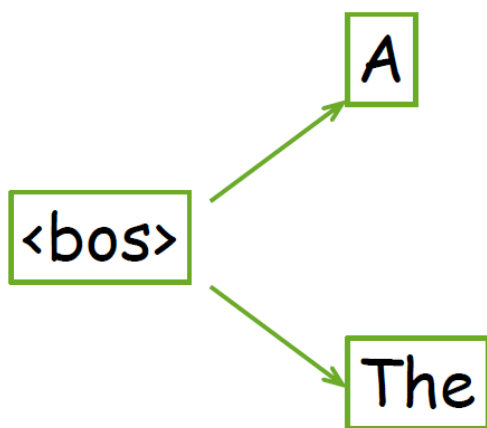
На каждом этапе сохраняем несколько лучших гипотез.

Начинаем с маркера начала предложения или с пустой последовательности.

`<bos>`

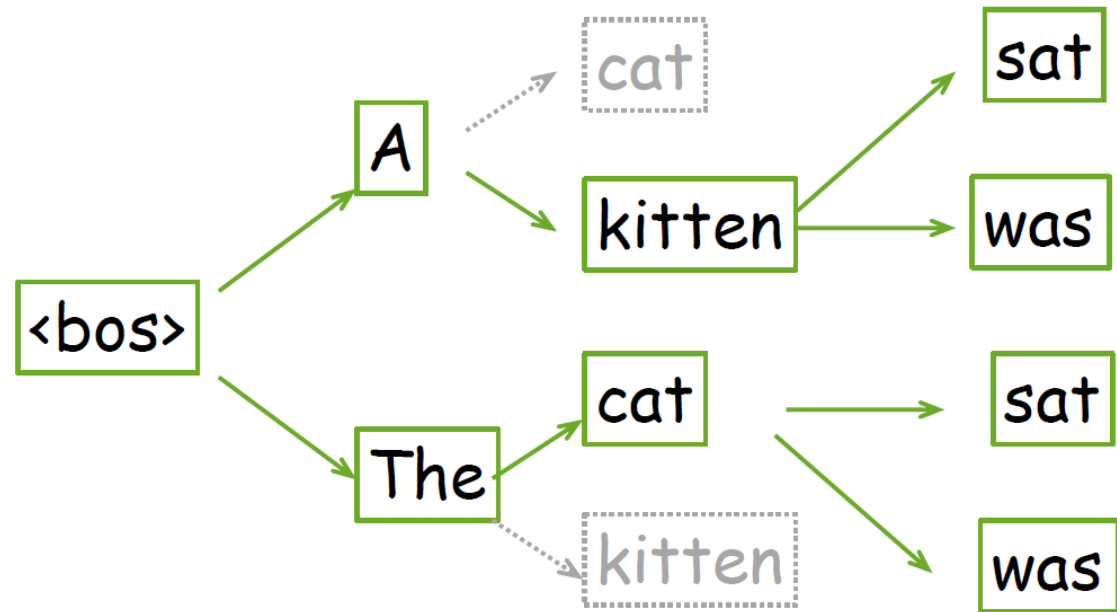
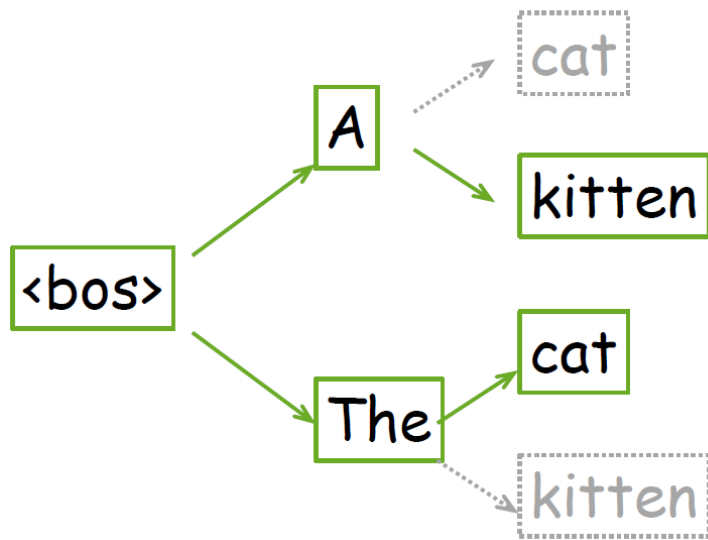
Beam Search (лучевой поиск)

Сгенерировать: beam_size наиболее вероятных токенов

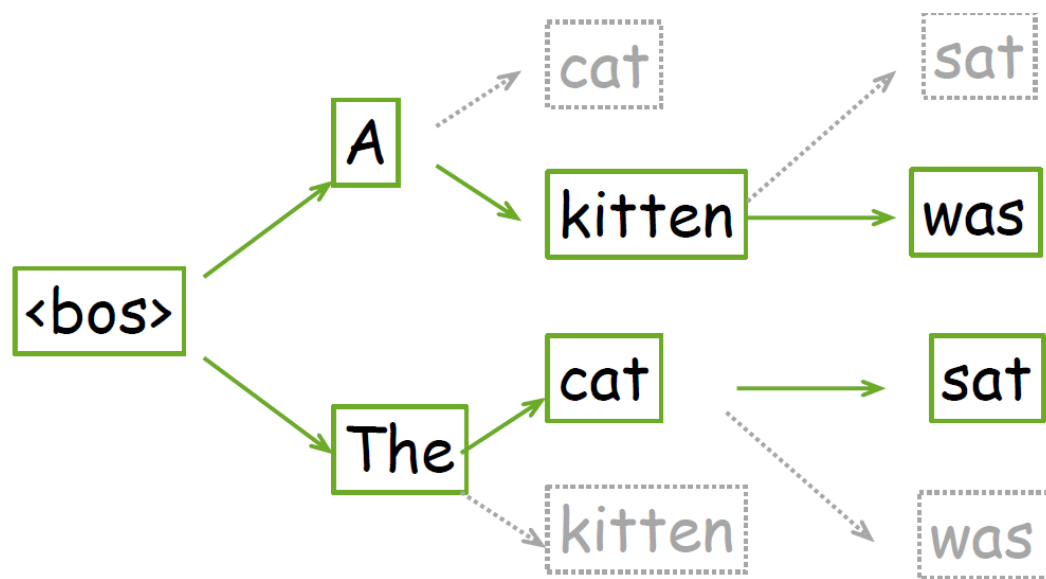


Beam Search (лучевой поиск)

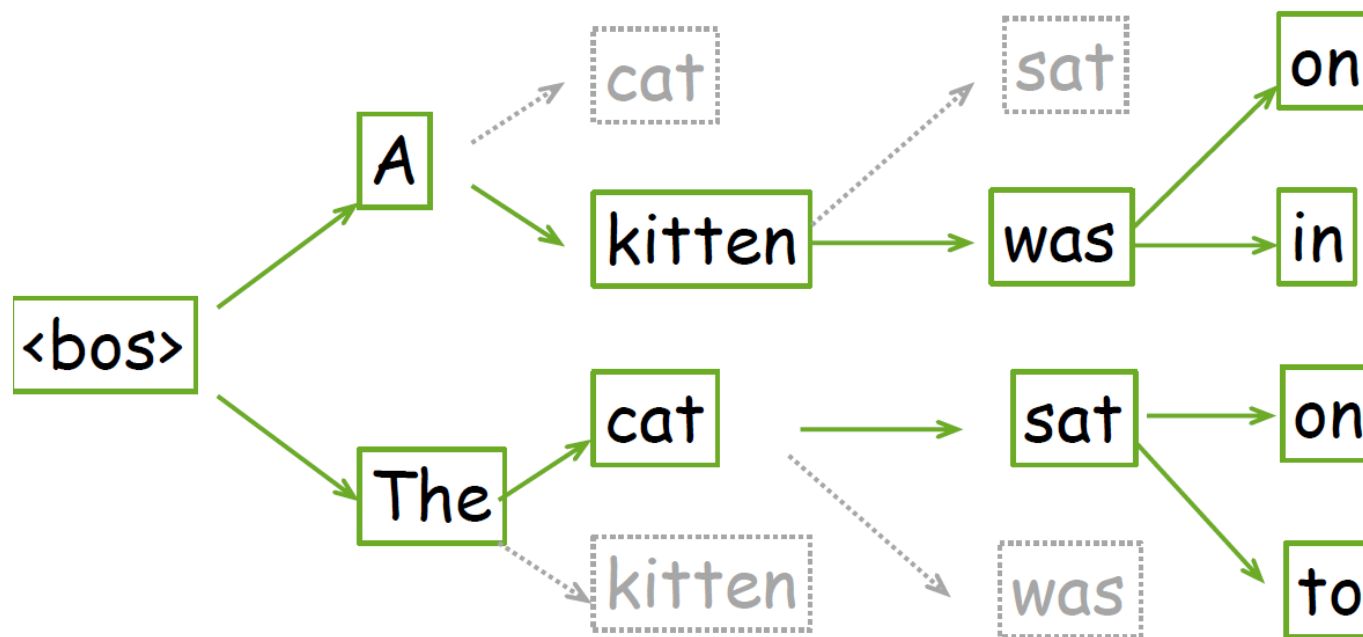
Посмотрим на вероятности: $P(\text{The cat}) > P(\text{A kitten}) > P(\text{A cat}) > P(\text{The kitten})$



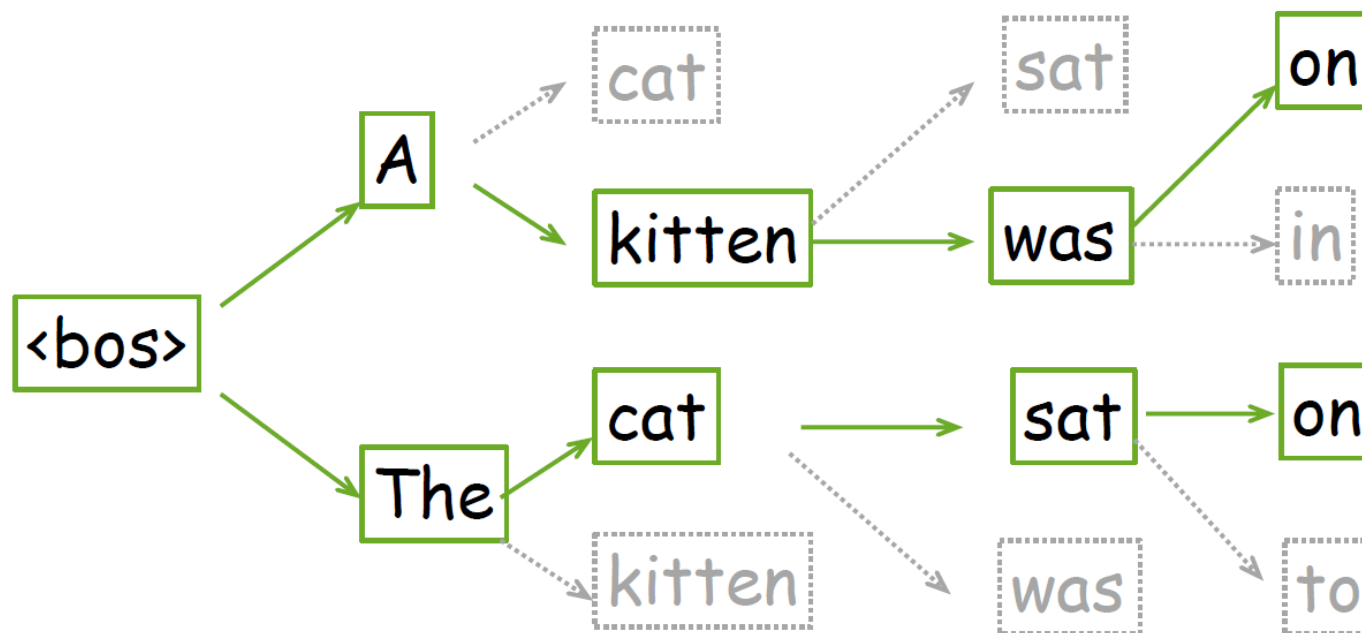
Beam Search (лучевой поиск)



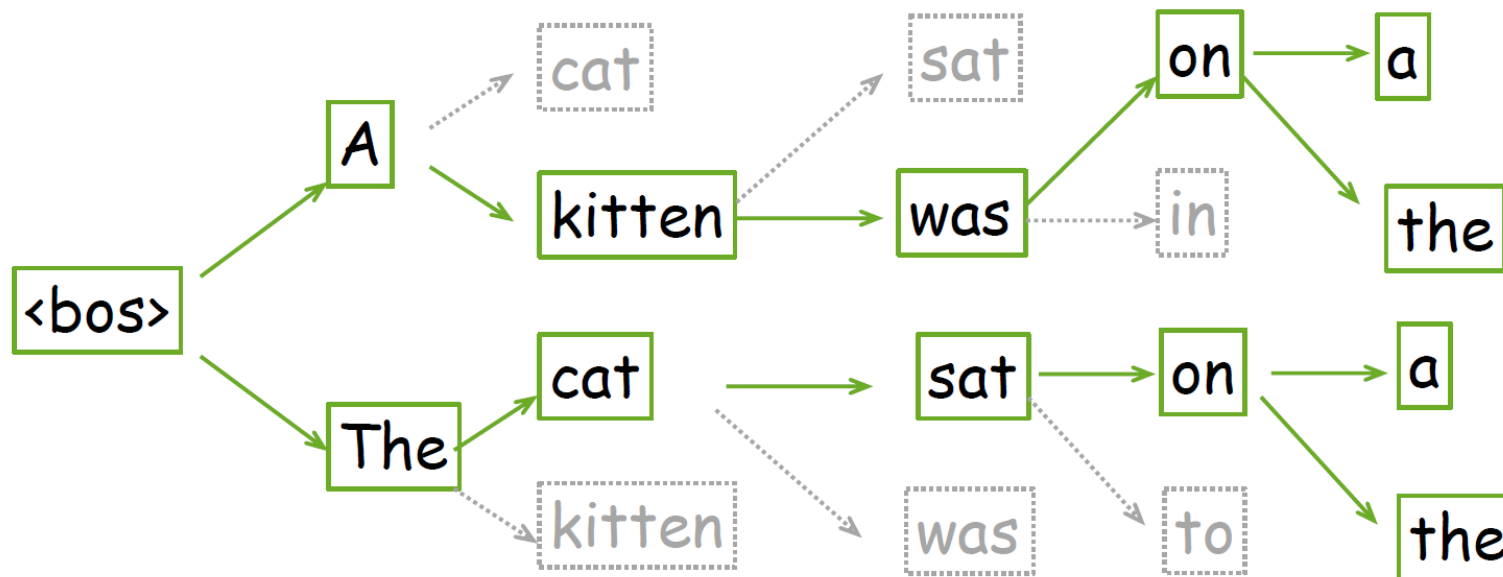
Beam Search (лучевой поиск)



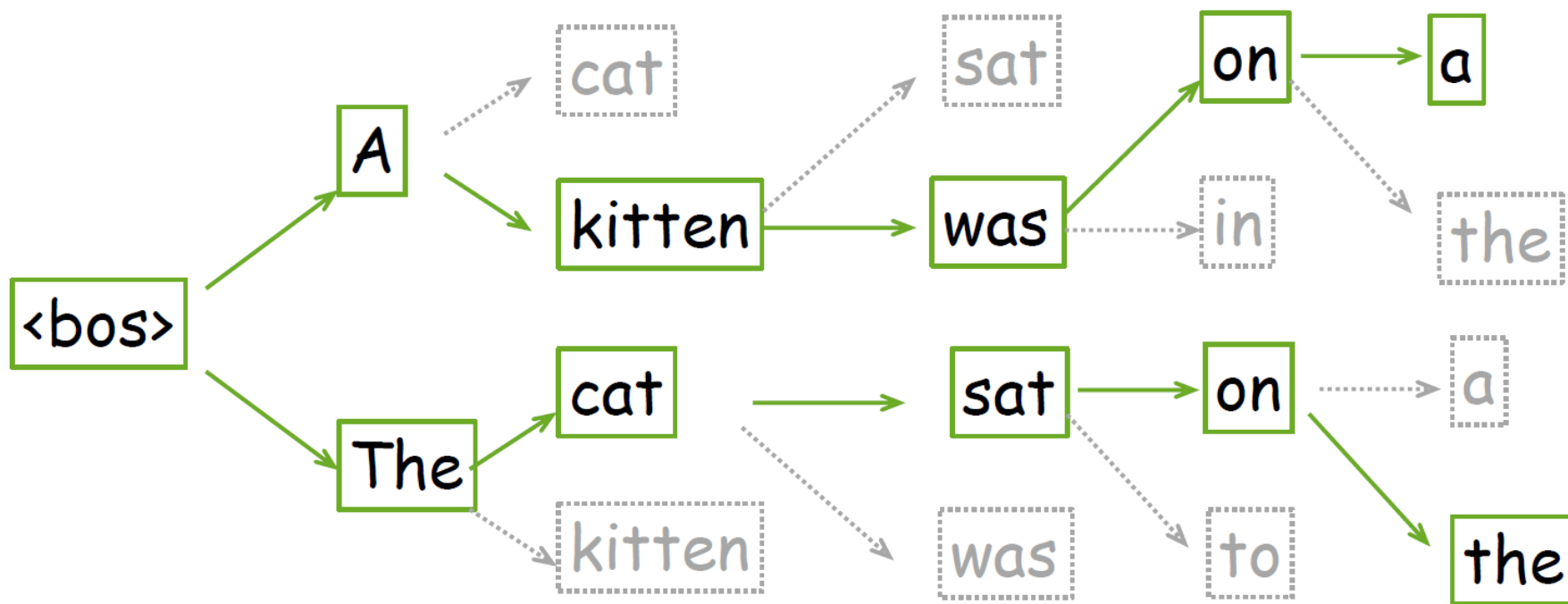
Beam Search (лучевой поиск)



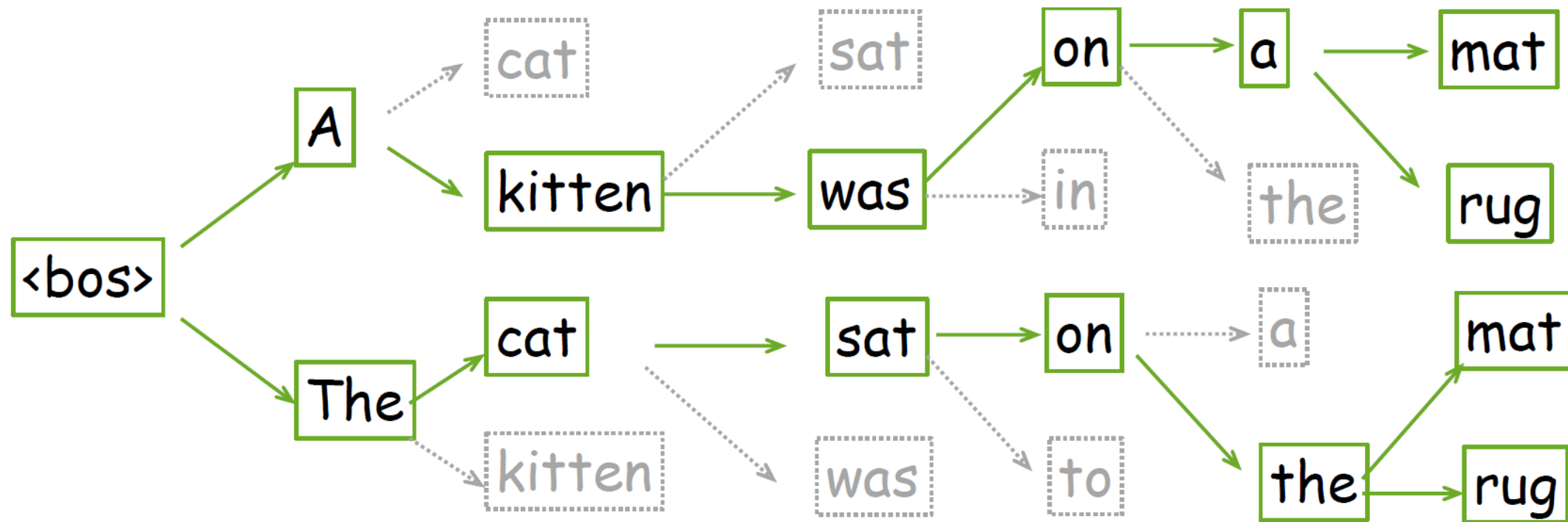
Beam Search (лучевой поиск)



Beam Search (лучевой поиск)

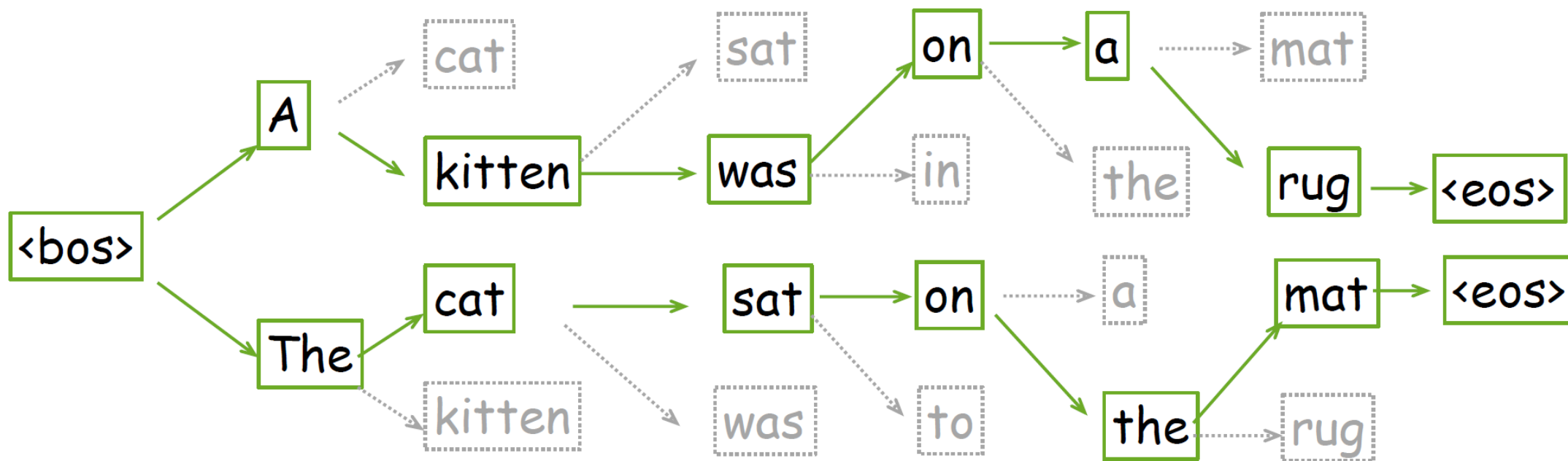


Beam Search (лучевой поиск)



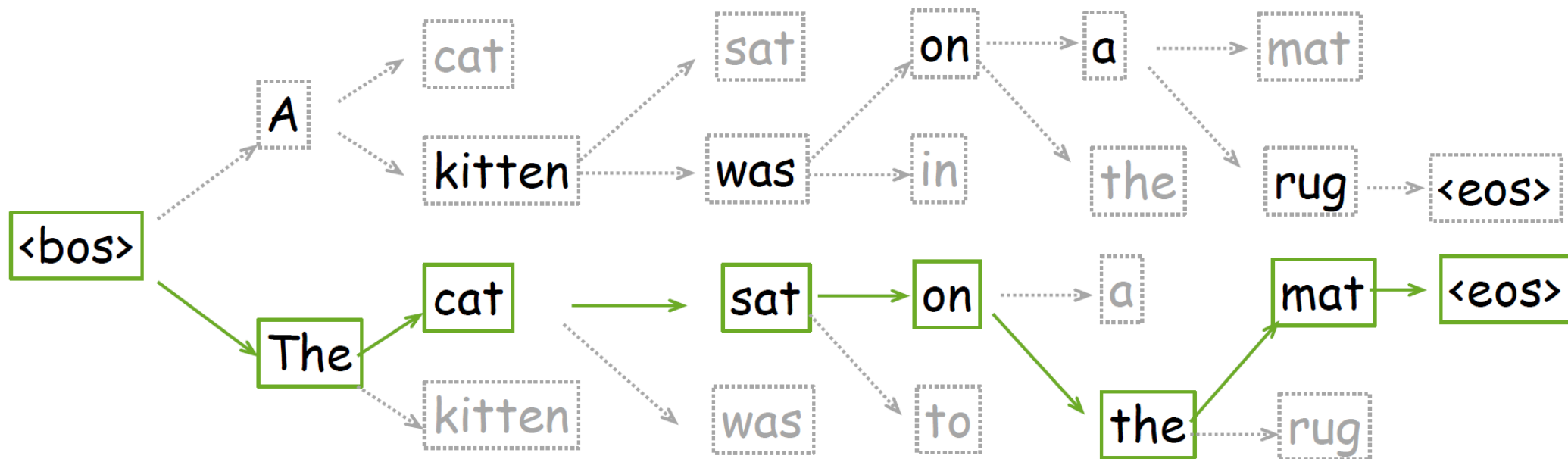
Beam Search (лучевой поиск)

Все гипотезы завершены, генерация окончена



Beam Search (лучевой поиск)

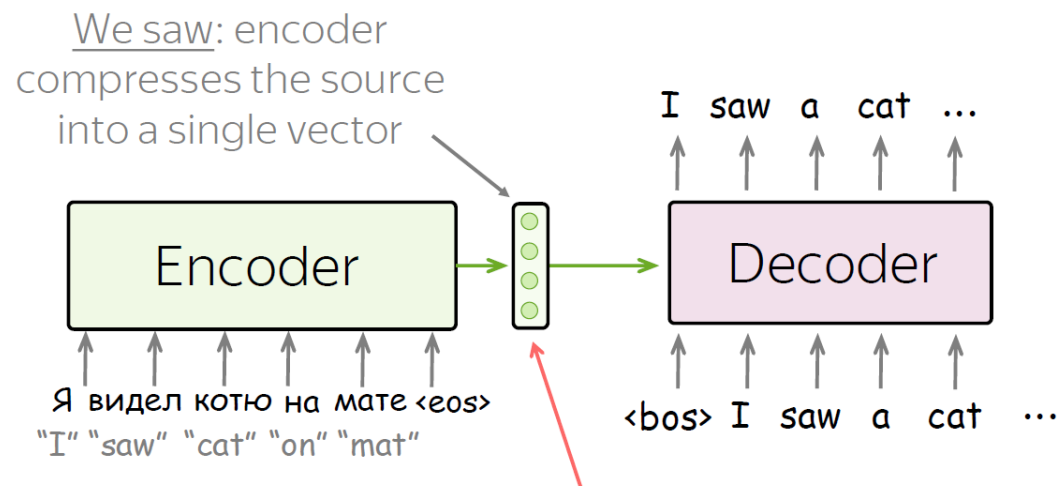
Выберем гипотезу с наибольшей вероятностью.



Attention. Проблема фиксированного представления кодера

Фиксированное представление исходного текста плохо:

- кодеру сложно сжать предложение
- декодеру может потребоваться разная информация на разных этапах



Problem: this is a bottleneck!

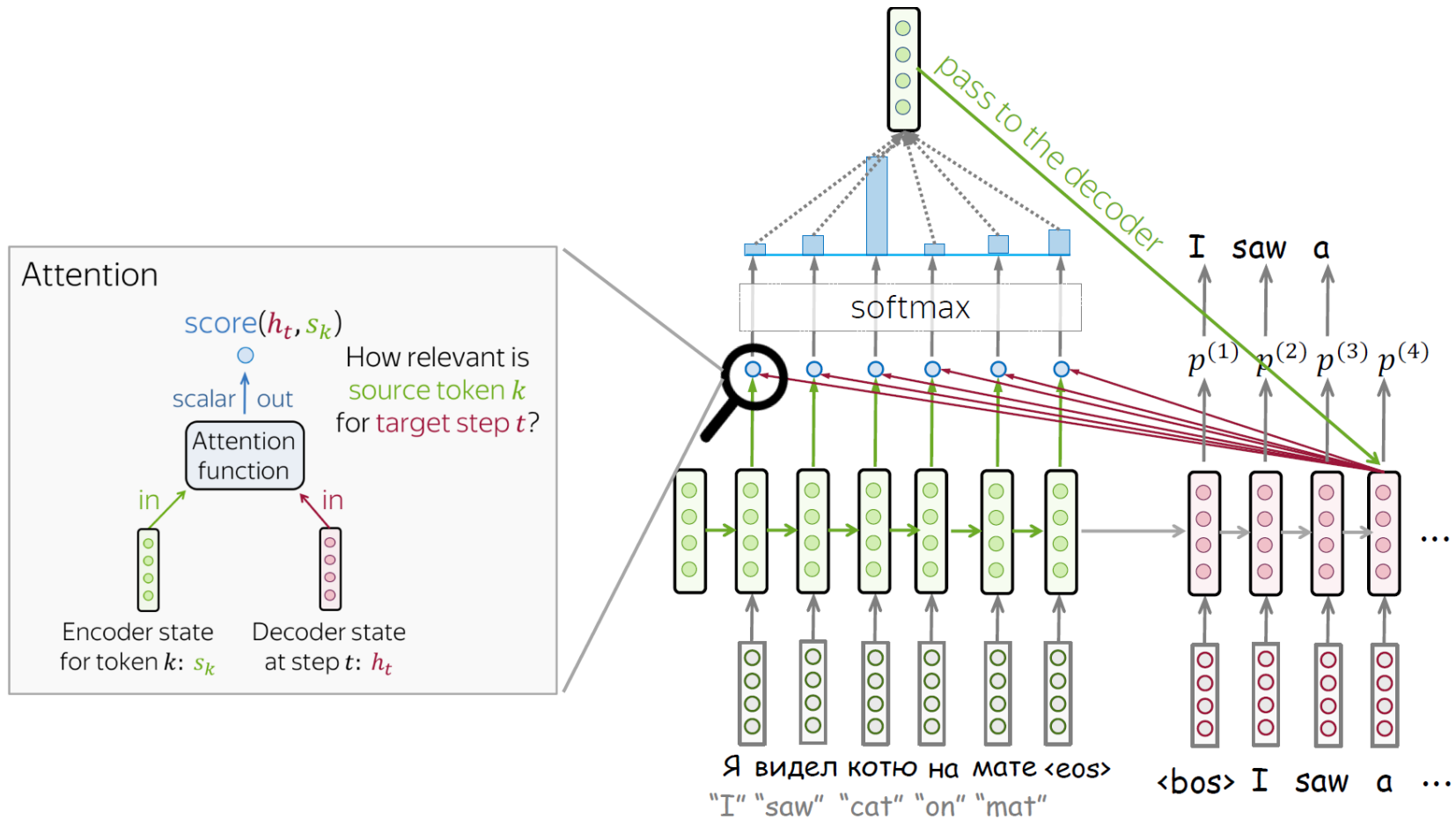
Attention

Attention: на разных этапах позволяем модели «фокусироваться» на разных частях входных данных.

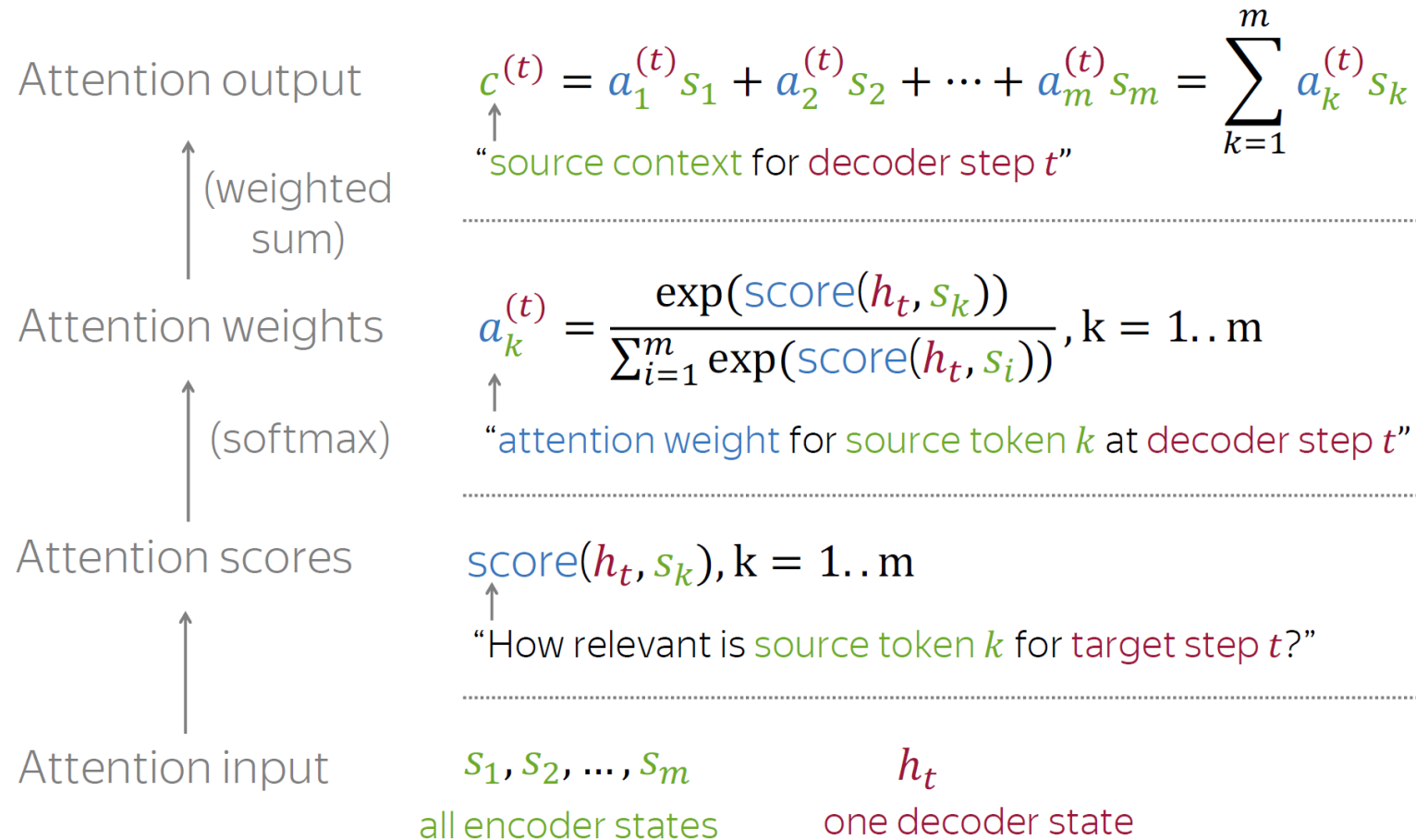
Выход attention: взвешенная сумма состояний кодера с весами attention.

Весы attention: распределение по исходным токенам.

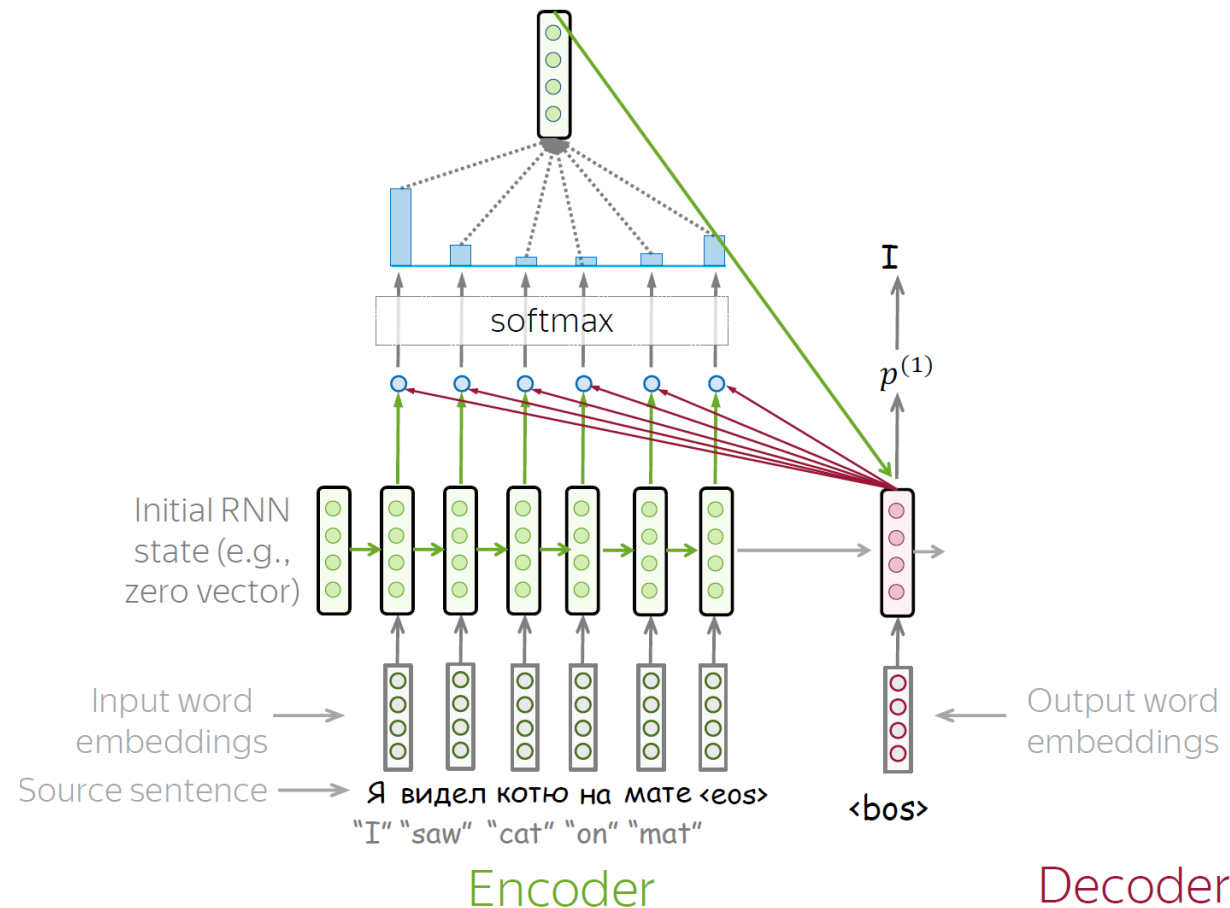
Attention



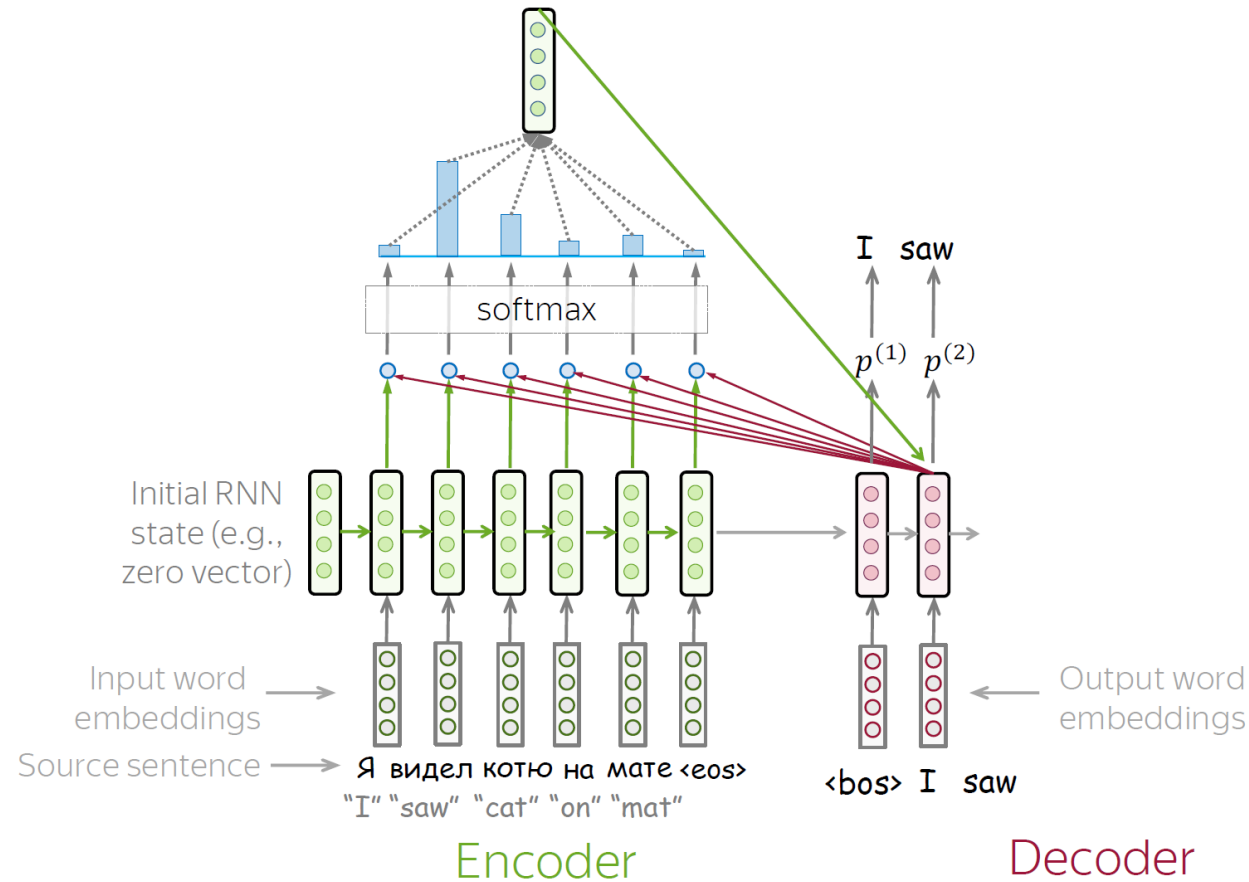
Вычислительная схема



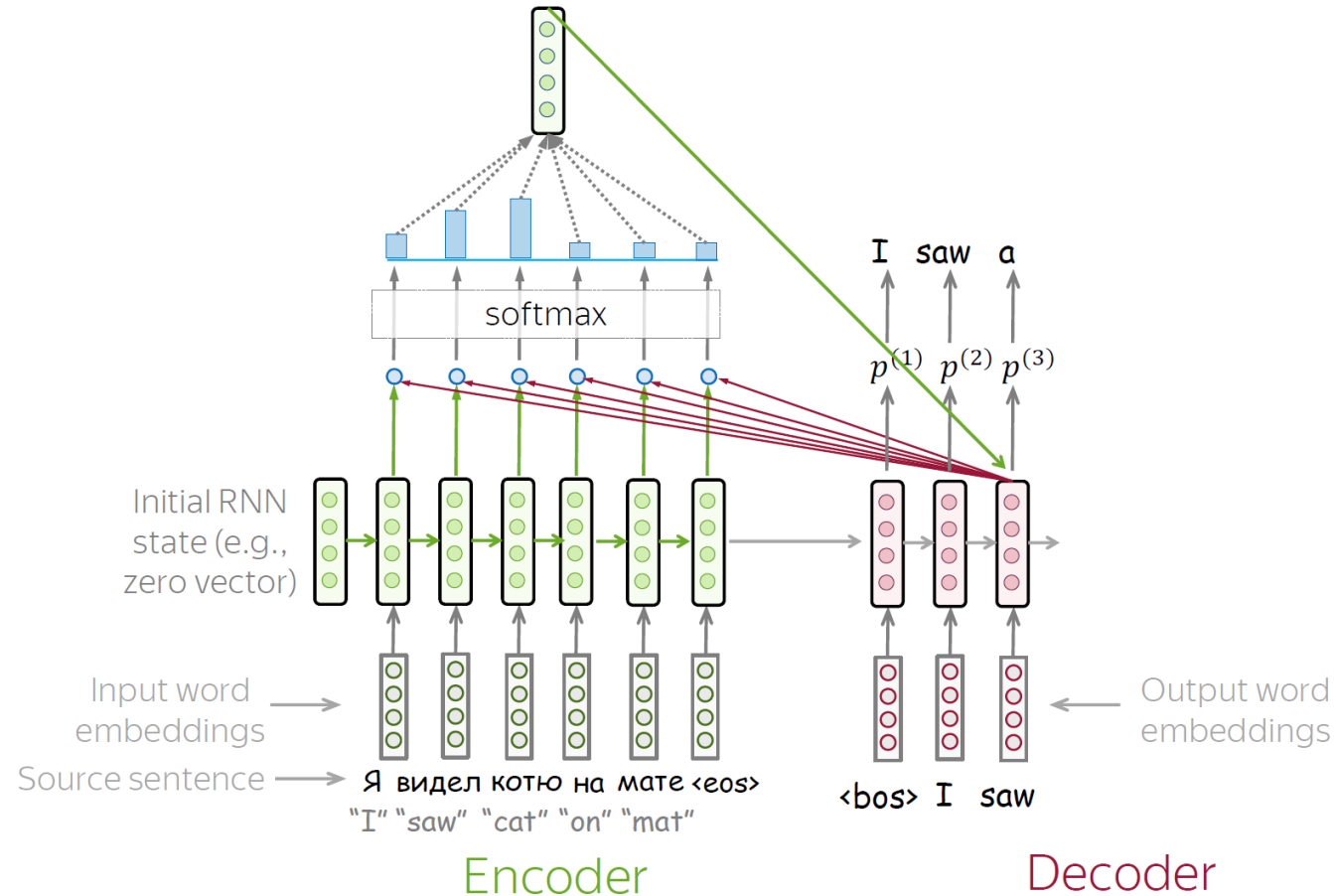
Модель учится выбирать соответствующие токены



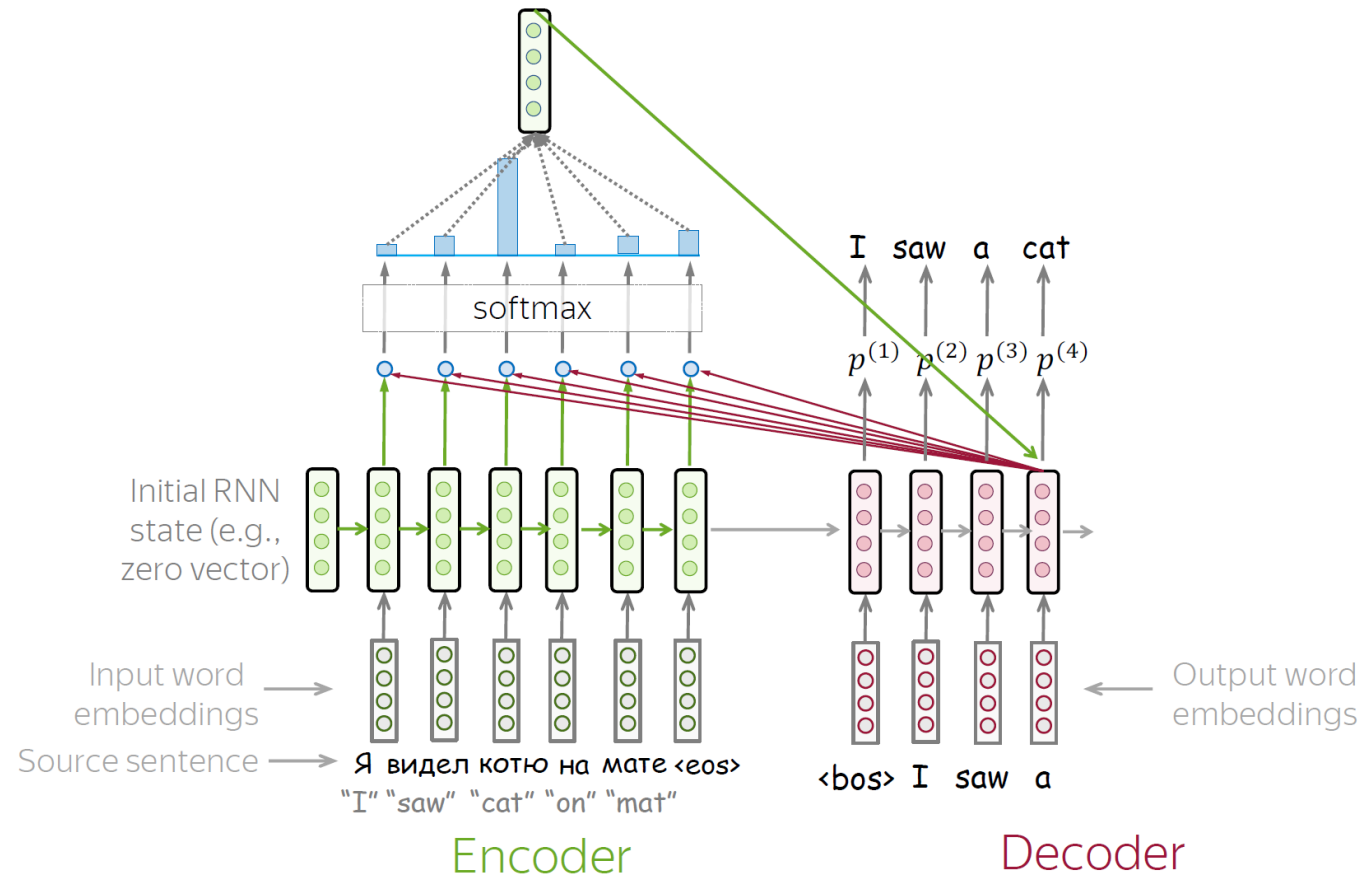
Модель учится выбирать соответствующие токены



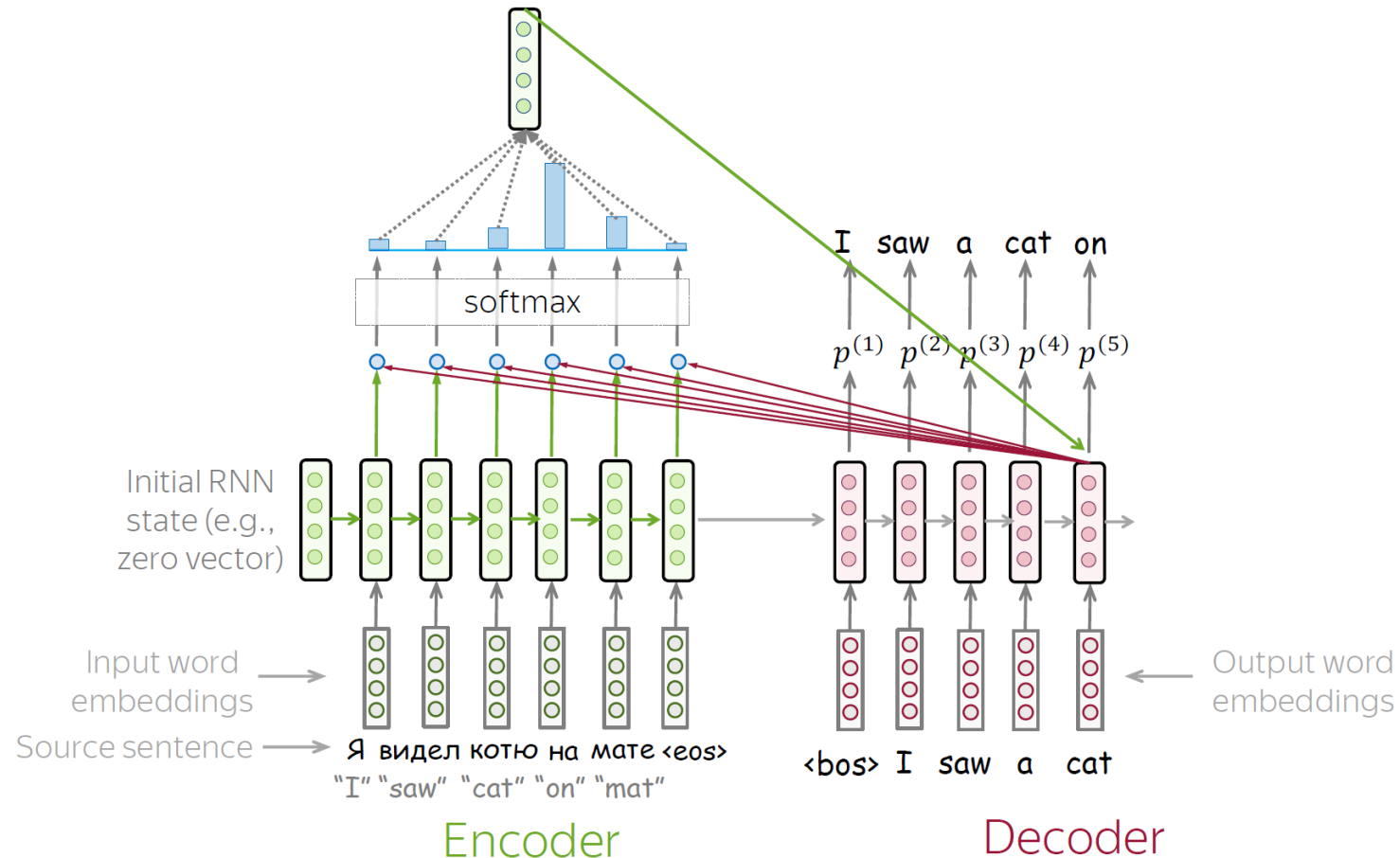
Модель учится выбирать соответствующие токены



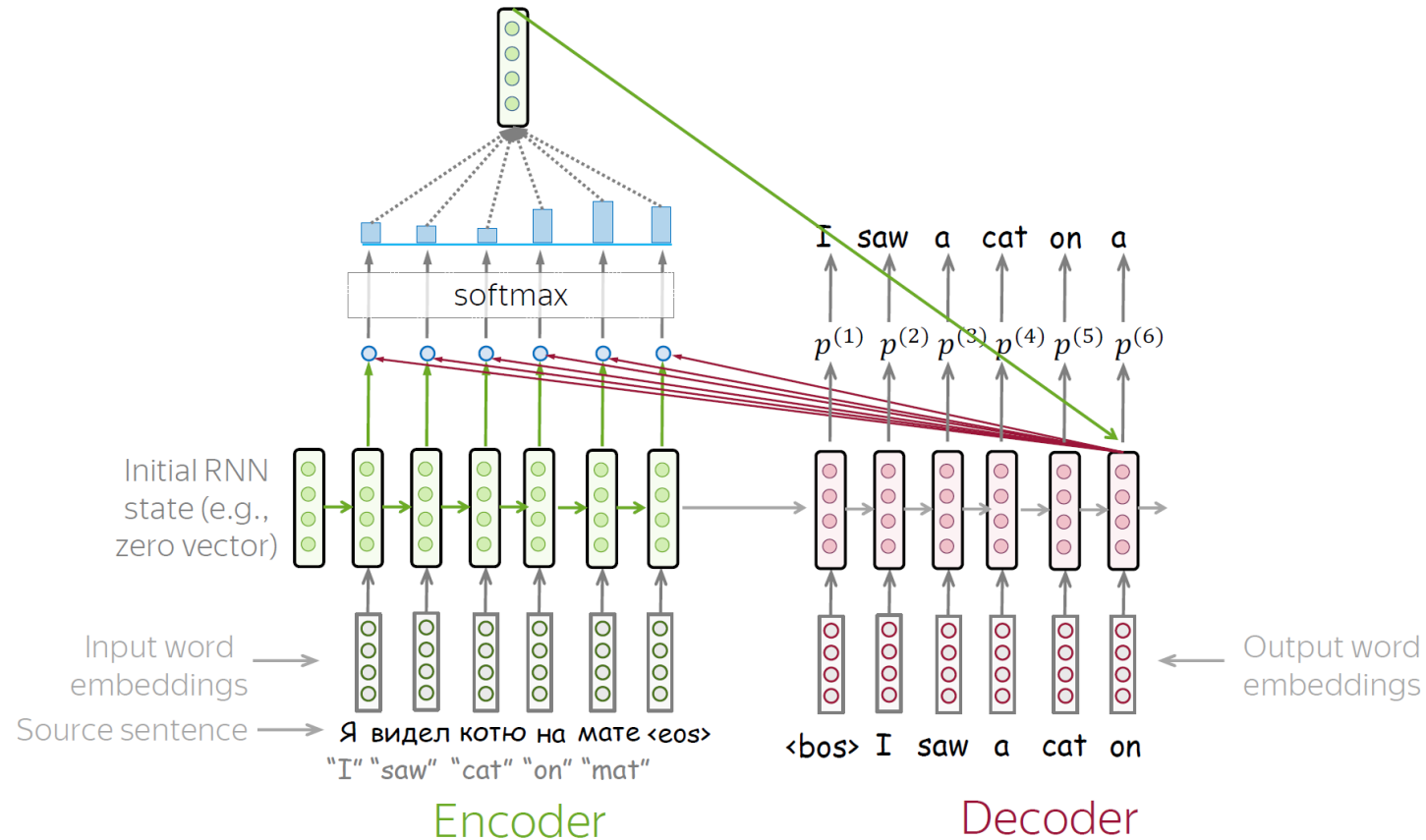
Модель учится выбирать соответствующие токены



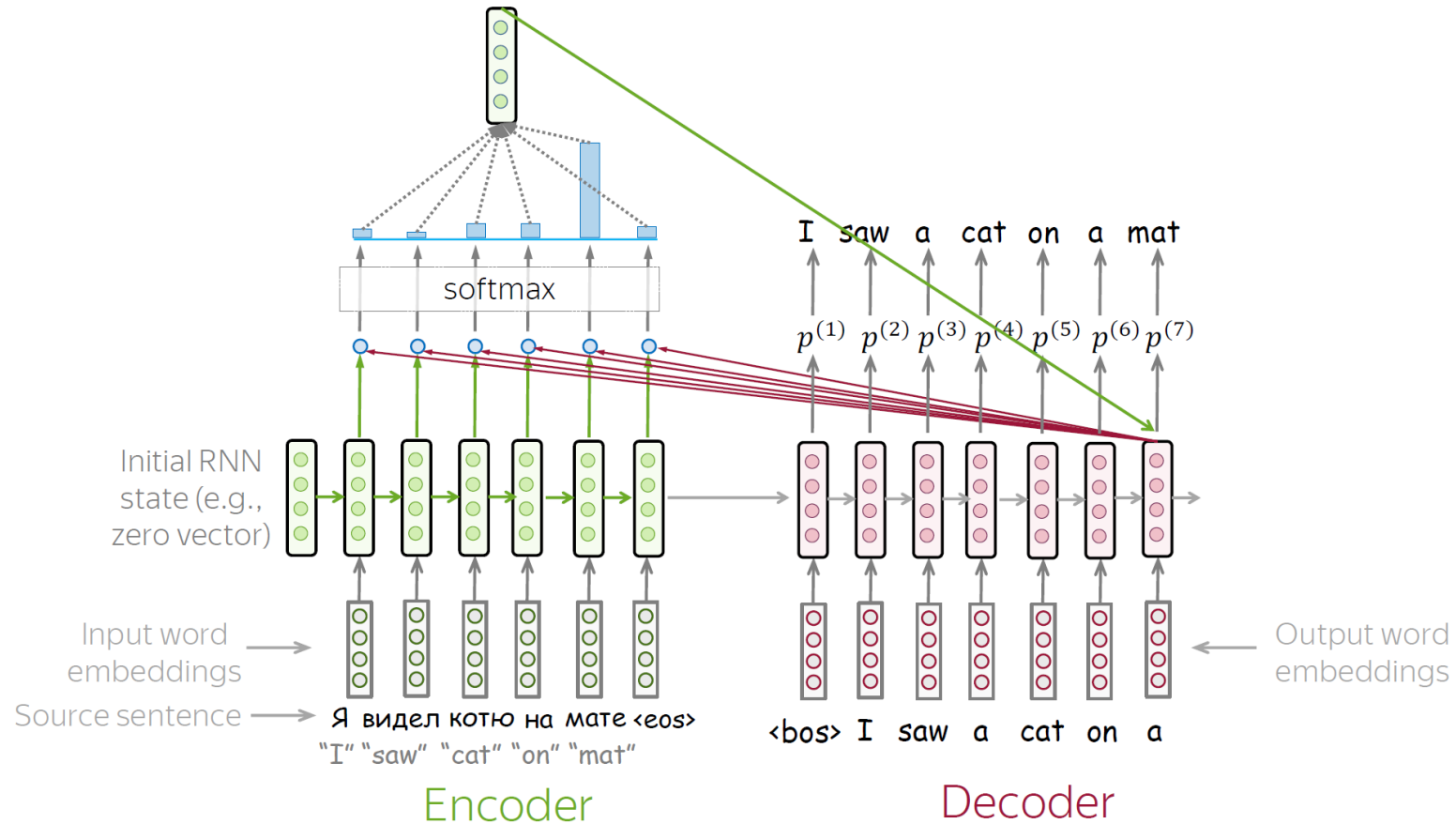
Модель учится выбирать соответствующие токены



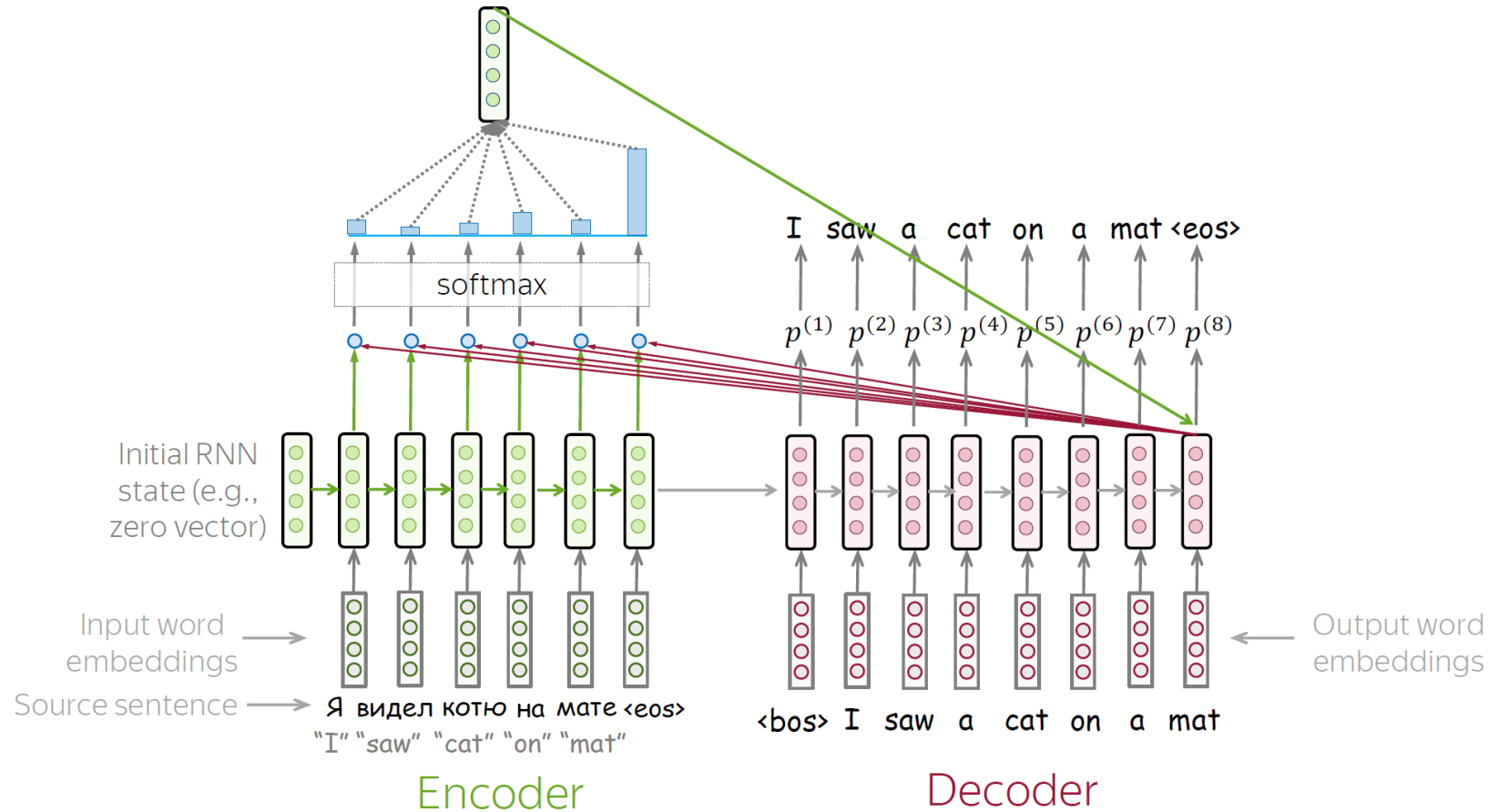
Модель учится выбирать соответствующие токены



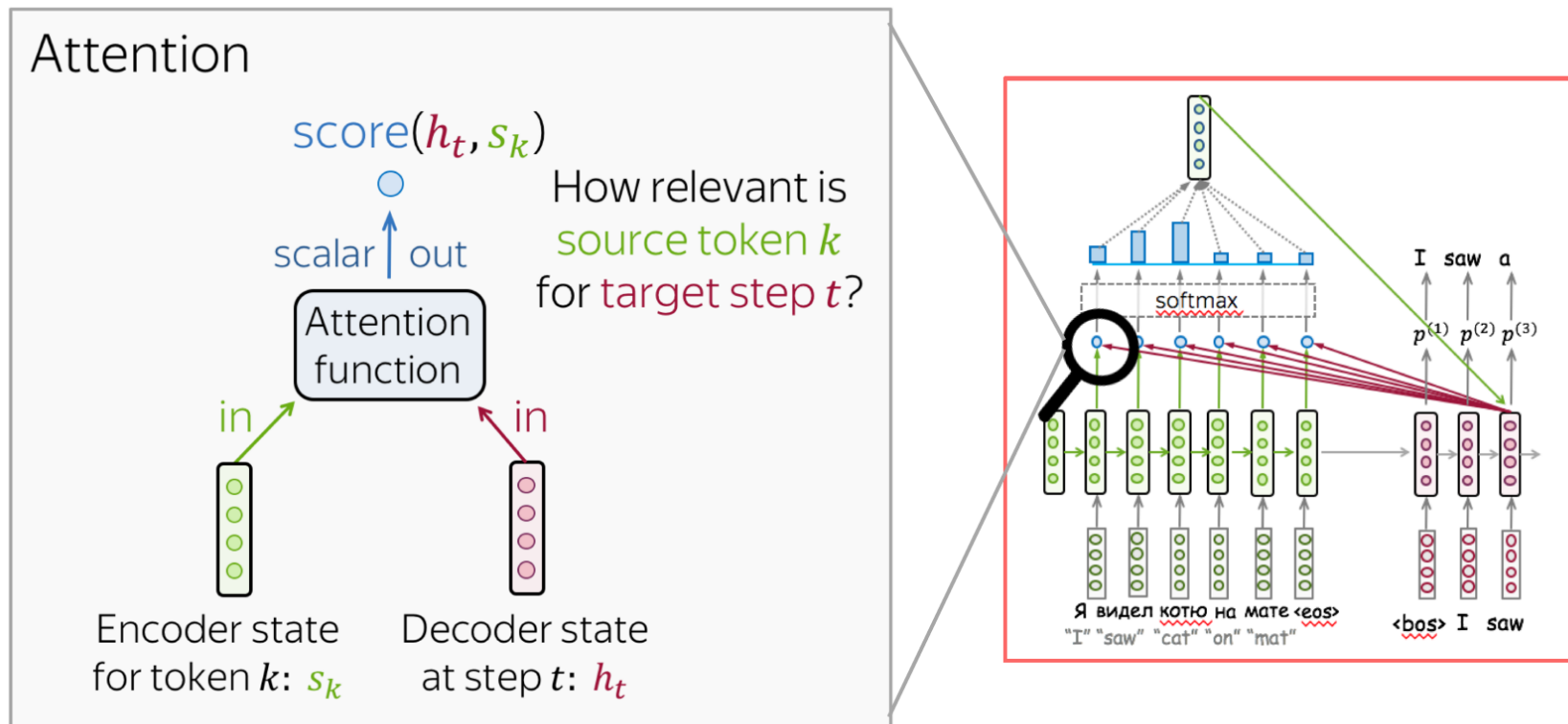
Модель учится выбирать соответствующие токены



Модель учится выбирать соответствующие токены

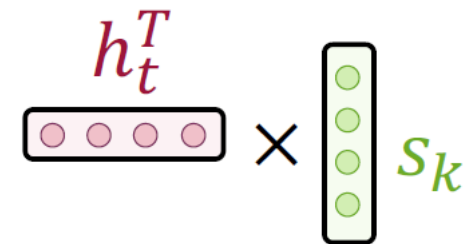


Функции оценки attention



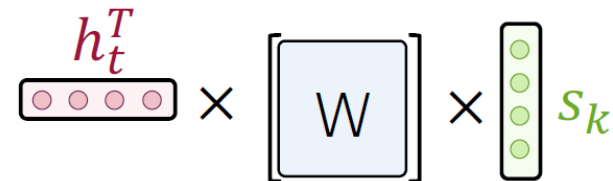
Функции оценки attention

- Скалярное произведение: $\text{score}(h_t, s_k) = h_t^T s_k$



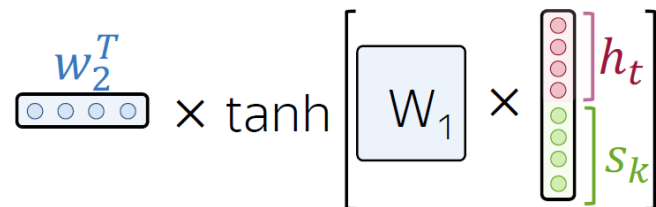
$$h_t^T \times s_k$$

- Билинейная функция: $\text{score}(h_t, s_k) = h_t^T W s_k$



$$h_t^T \times \begin{bmatrix} W \end{bmatrix} \times s_k$$

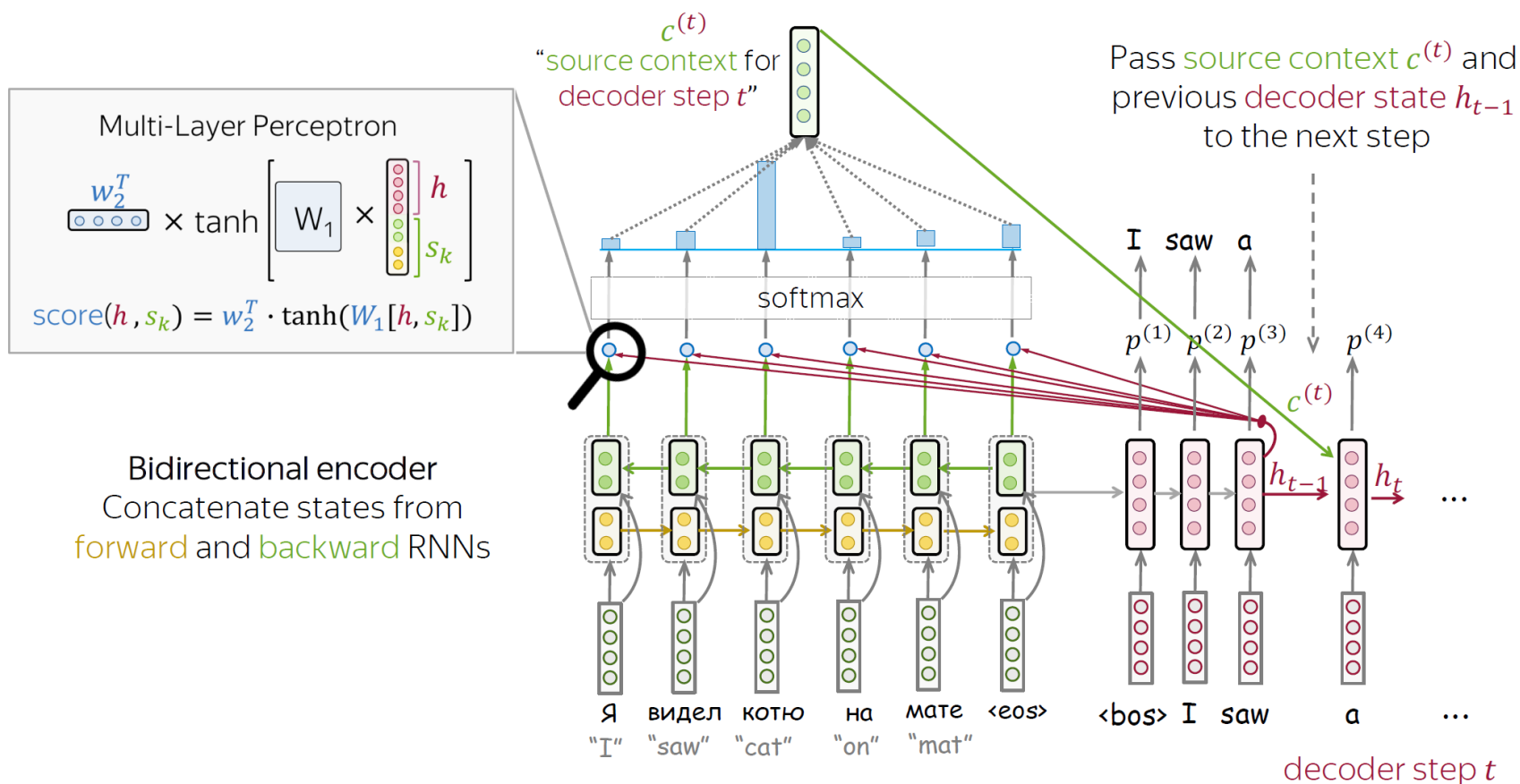
- Многослойный персептрон:



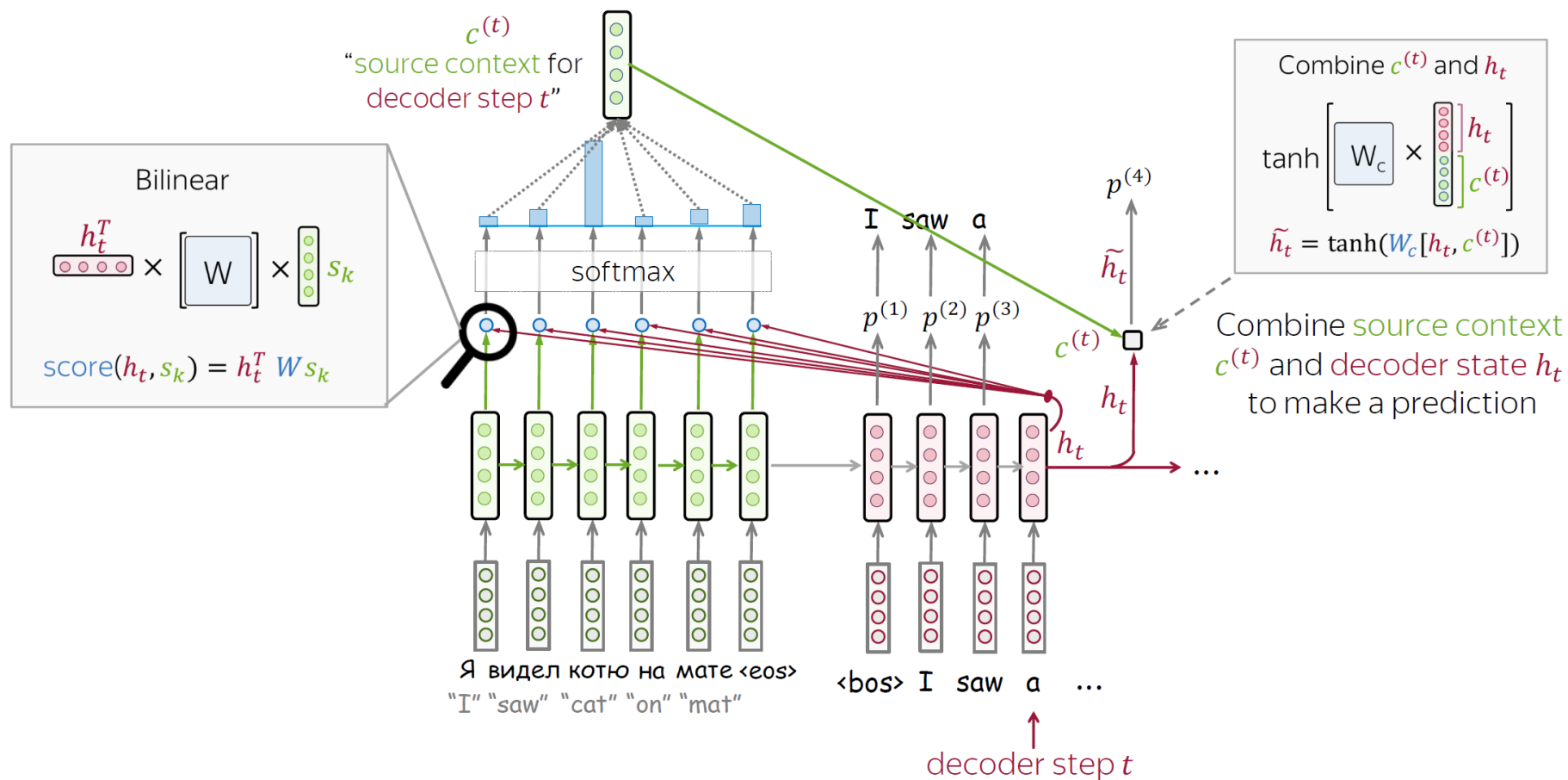
$$w_2^T \times \tanh \left[\begin{bmatrix} W_1 \end{bmatrix} \times \begin{bmatrix} h_t \\ s_k \end{bmatrix} \right]$$

$$\text{score}(h_t, s_k) = w_2^T \cdot \tanh(W_1 [h_t, s_k])$$

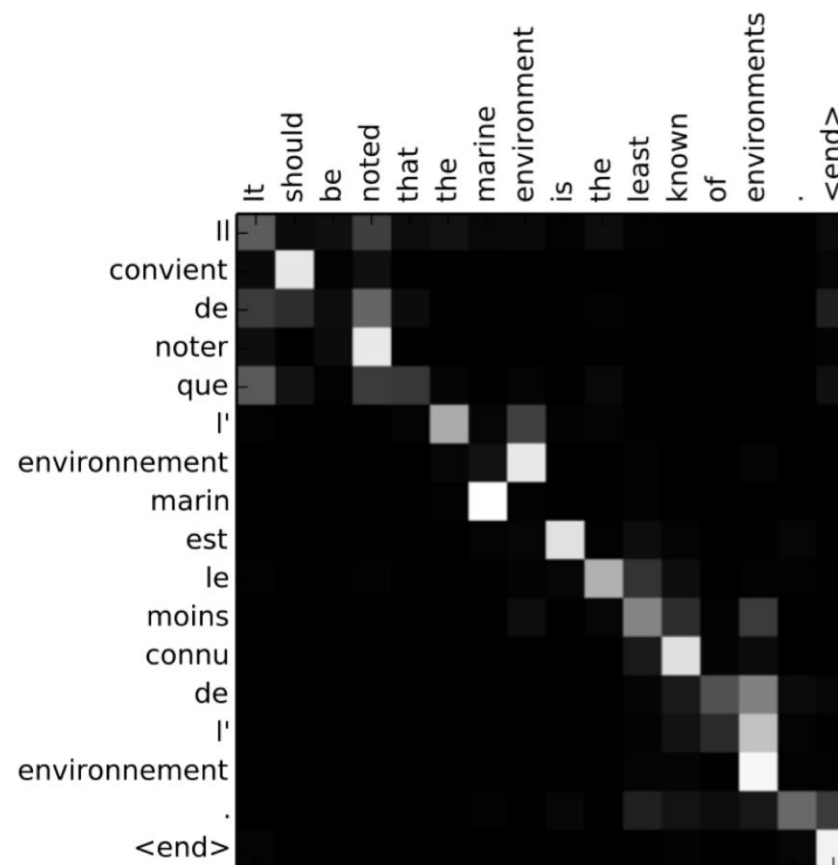
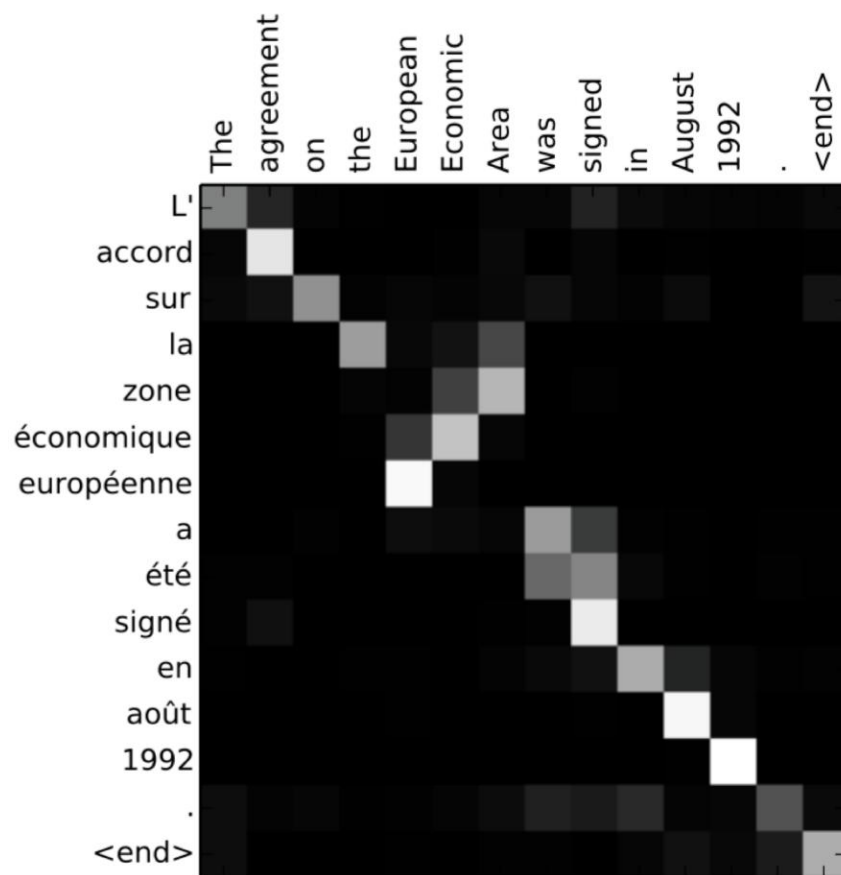
Модель Богданова (исходная модель attention)



Модель Луонга



Выравнивание, выученное с attention



Transformer: основная идея

	Seq2seq without attention	Seq2seq with attention	Transformer
processing within encoder	RNN/CNN	RNN/CNN	attention
processing within decoder	RNN/CNN	RNN/CNN	attention
decoder-encoder interaction	static fixed- sized vector	attention	attention

Attention

Кодировщик

Все исходные токены смотрят друг на друга и обновляют представления (повторить N раз).

Декодировщик

Целевой token на текущем шаге смотрит на предыдущие целевые токены и на исходные представления, обновляет представление (повторить N раз).

Почему это может быть лучше, чем RNN?

I arrived at the **bank** after crossing thestreet? ...river?

What does **bank** mean in this sentence?

Self-Attention

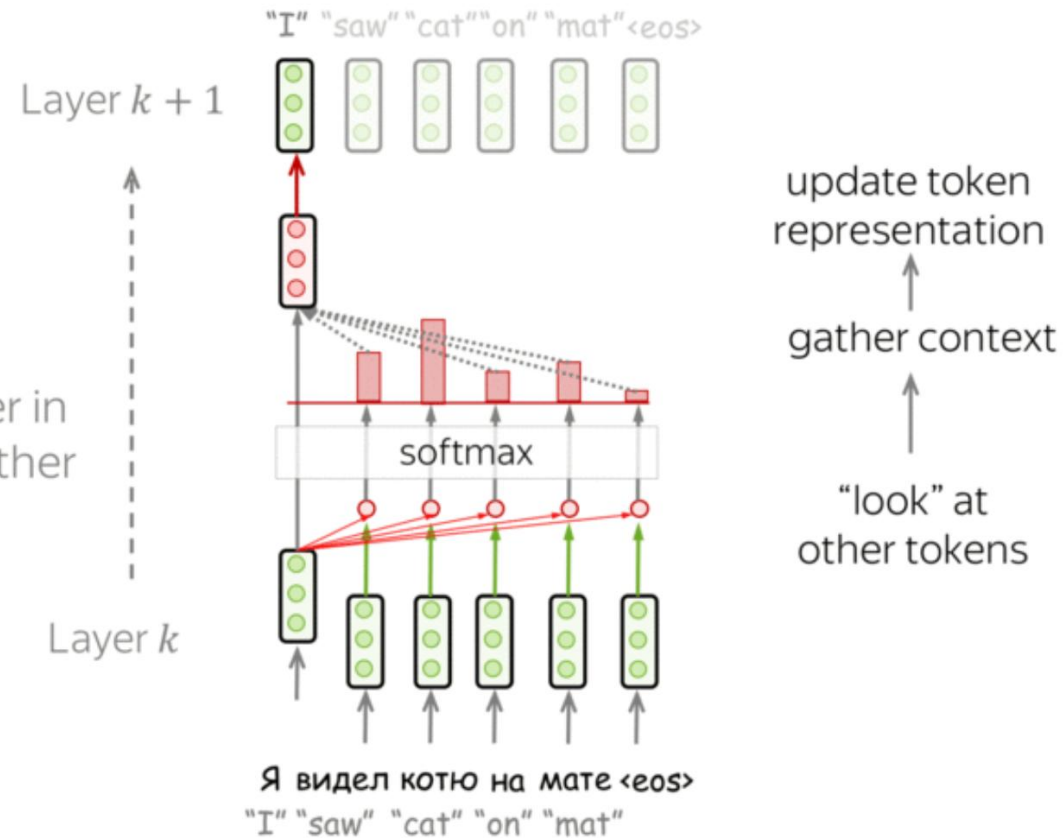
Attention направлено из одного текущего состояния декодера во все состояния кодировщика.

Self-Attention смотрит из каждого состояния из набора состояний во все остальные состояния в том же наборе.

Self-Attention

(in practice, this happens in parallel)

Tokens try to understand themselves better in context of each other



Запрос, Ключ, Значение (Query, Key, Value)

Каждый вектор получает три представления («роли»):

Запрос (query). Вектор, на который направлено attention

«Привет, у тебя есть эта информация?»

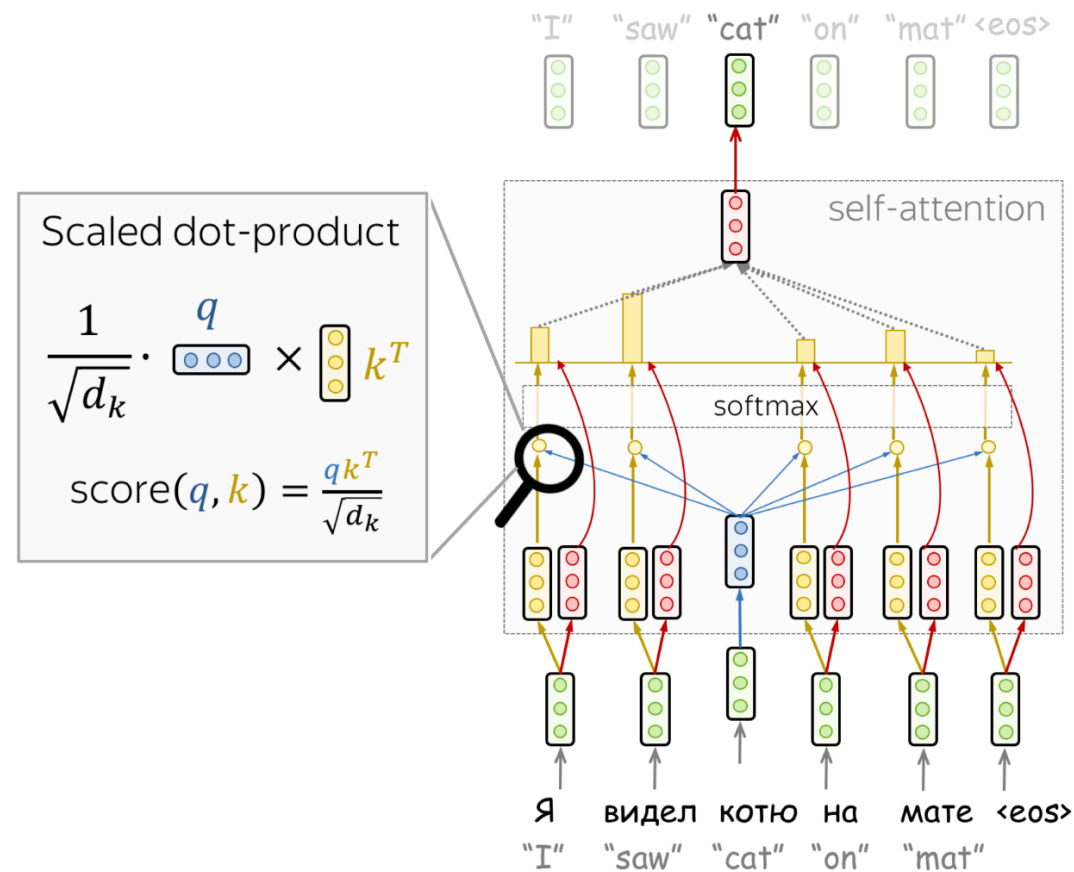
Ключ (key). Вектор, к которому обращается запрос для вычисления весов

«Привет, у меня есть информация — дайте мне большой вес!»

Значение (value). Их взвешенная сумма – это attention output

«Вот информация, которая у меня есть!»

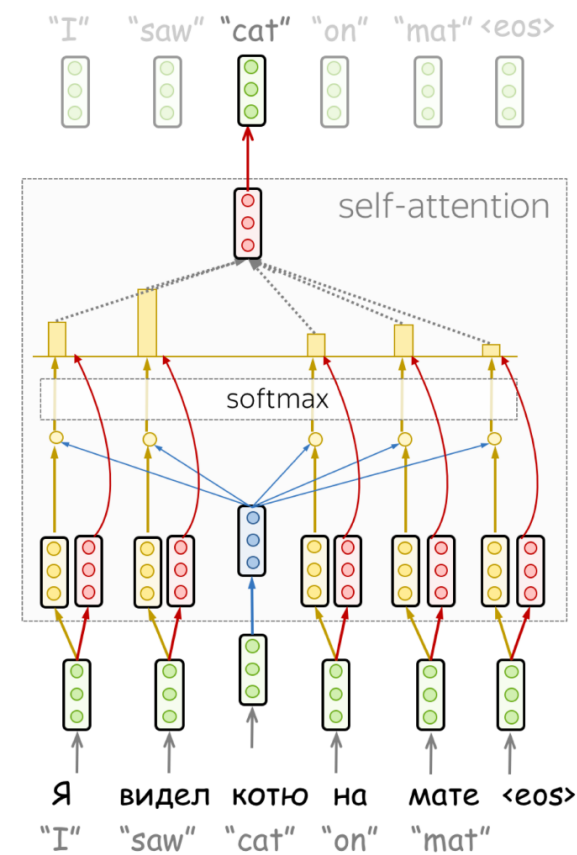
Запрос, Ключ, Значение (Query, Key, Value)



Запрос, Ключ, Значение (Query, Key, Value)

$$\text{Attention}(\underset{\substack{\text{from} \\ \text{blue arrow}}}{q}, \underset{\substack{\text{to} \\ \text{red arrow}}}{k}, v) = \overbrace{\text{softmax}\left(\frac{qk^T}{\sqrt{d_k}}\right)}^{\text{Attention weights}} v$$

vector dimensionality of K, V

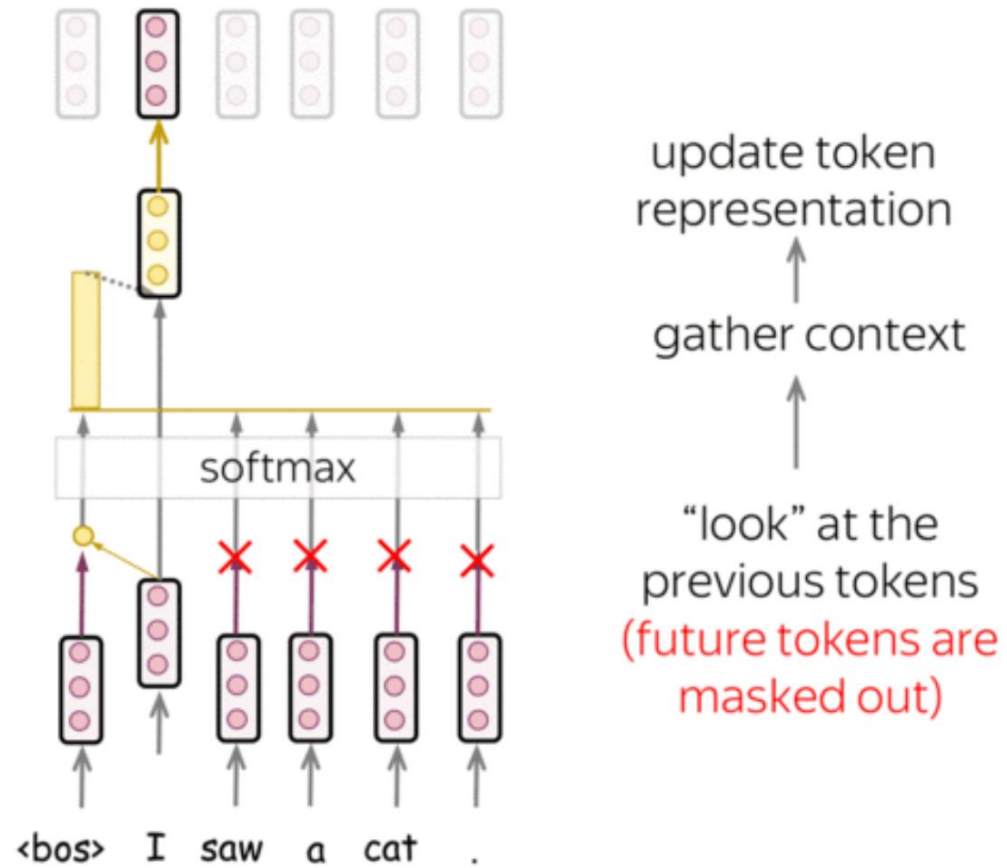


Скрытое Self-Attention

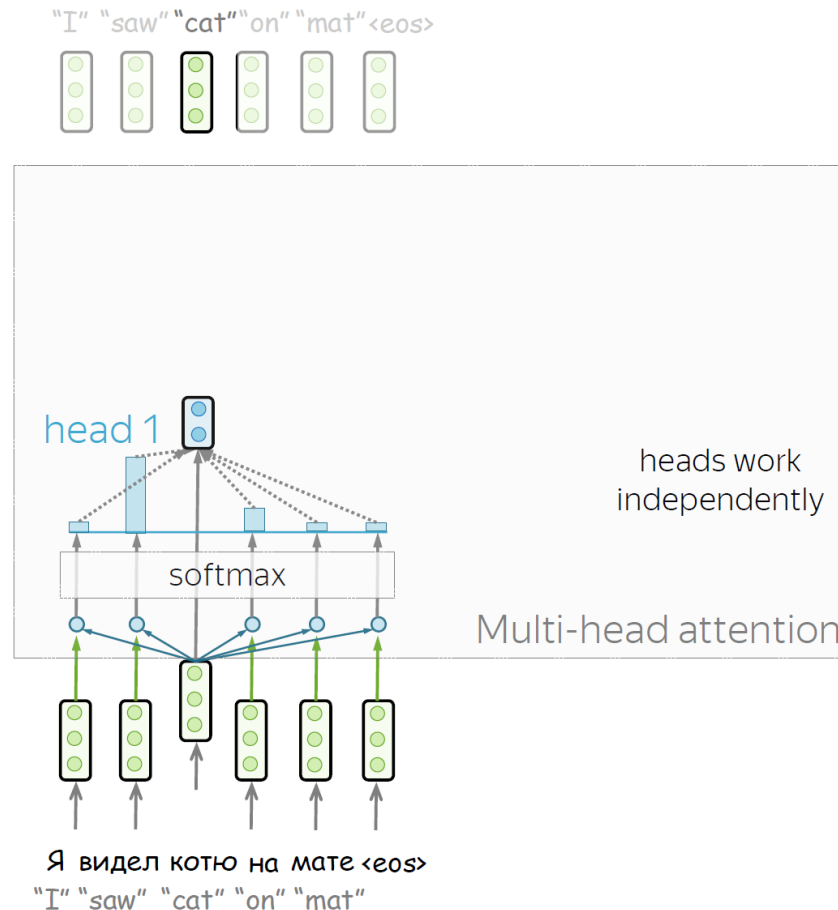
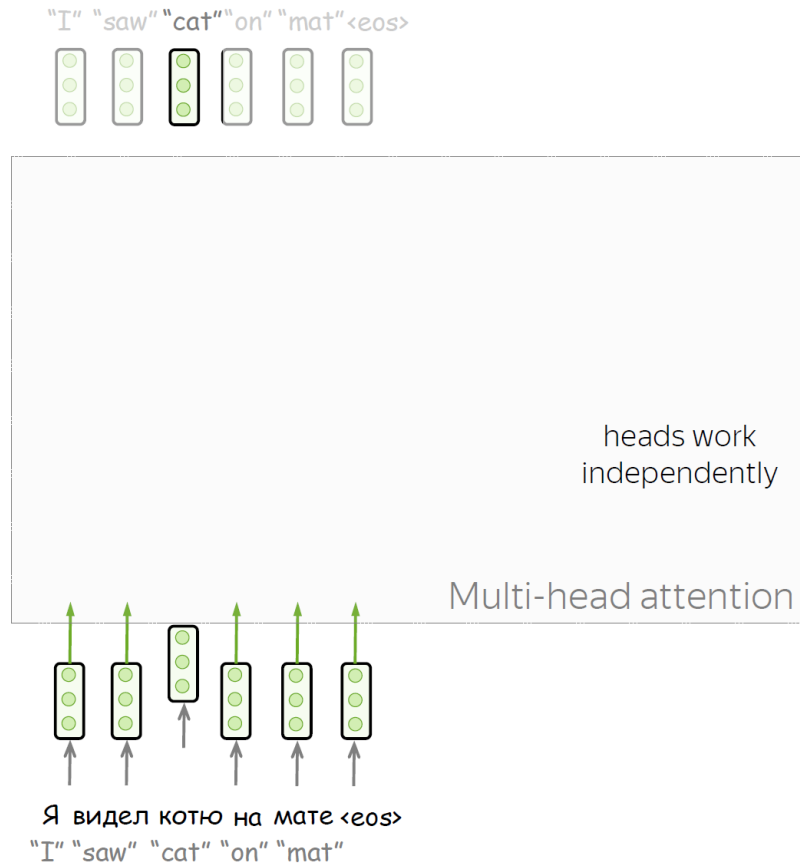
В декодере мы запрещаем просмотр будущих токенов – мы их не знаем.

Примечание: при обучении декодер обрабатывает все целевые токены одновременно – без масок он увидел бы будущее.

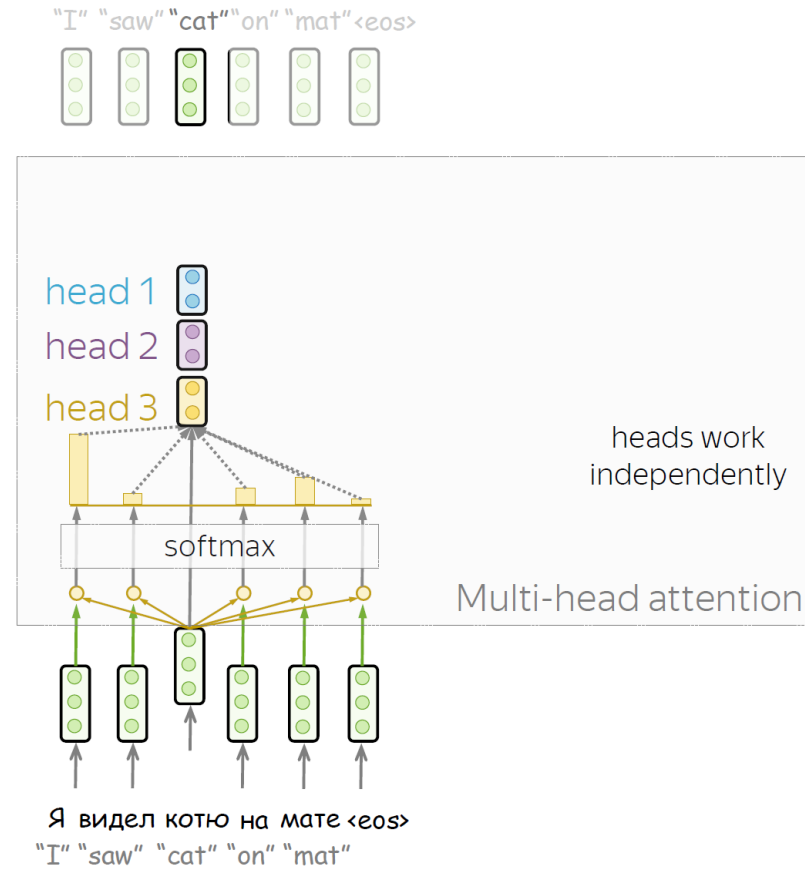
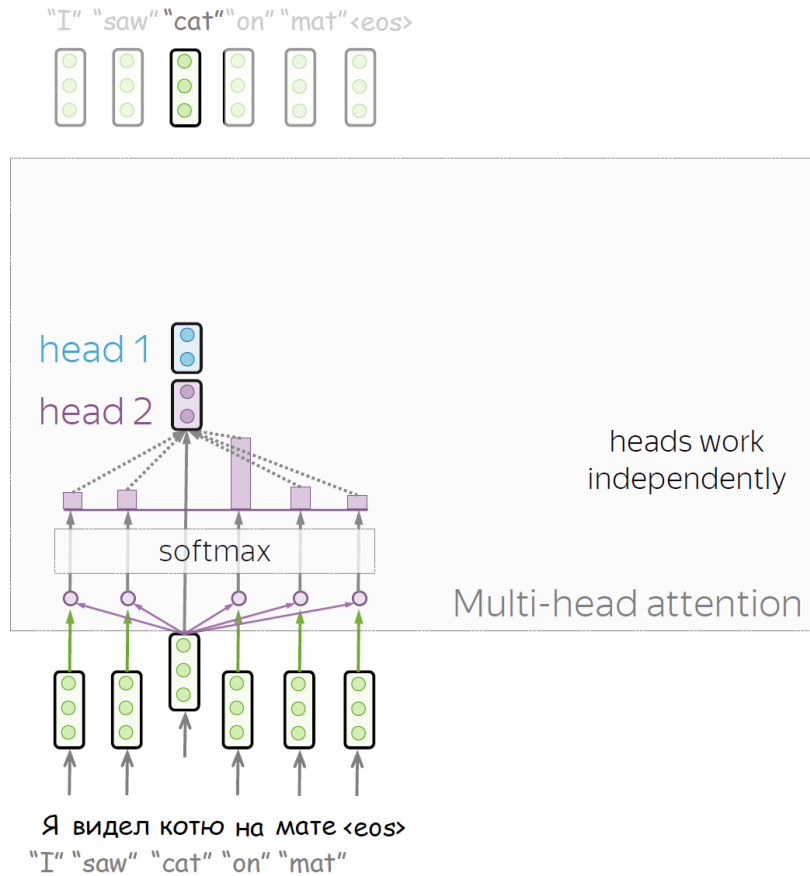
Скрытое Self-Attention



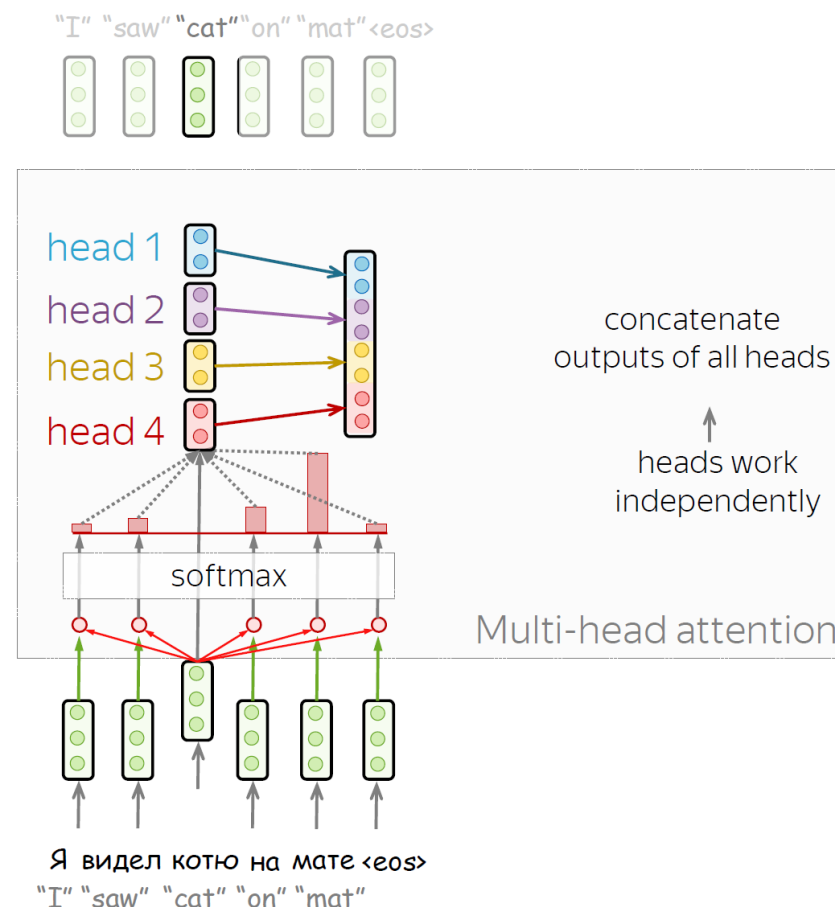
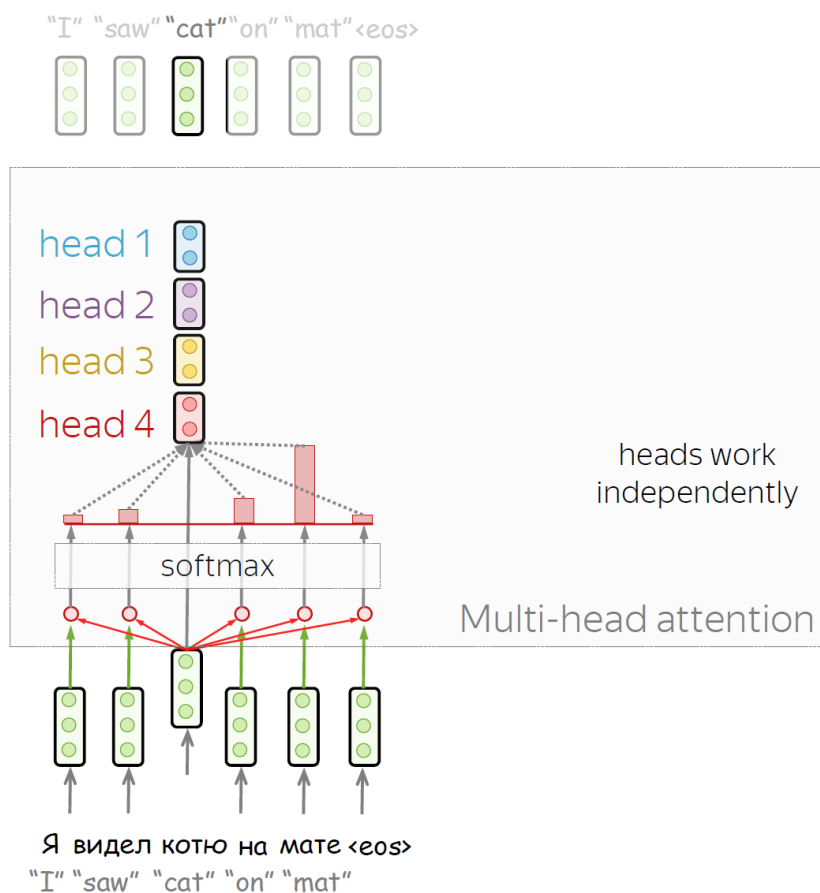
Multi-Head Attention



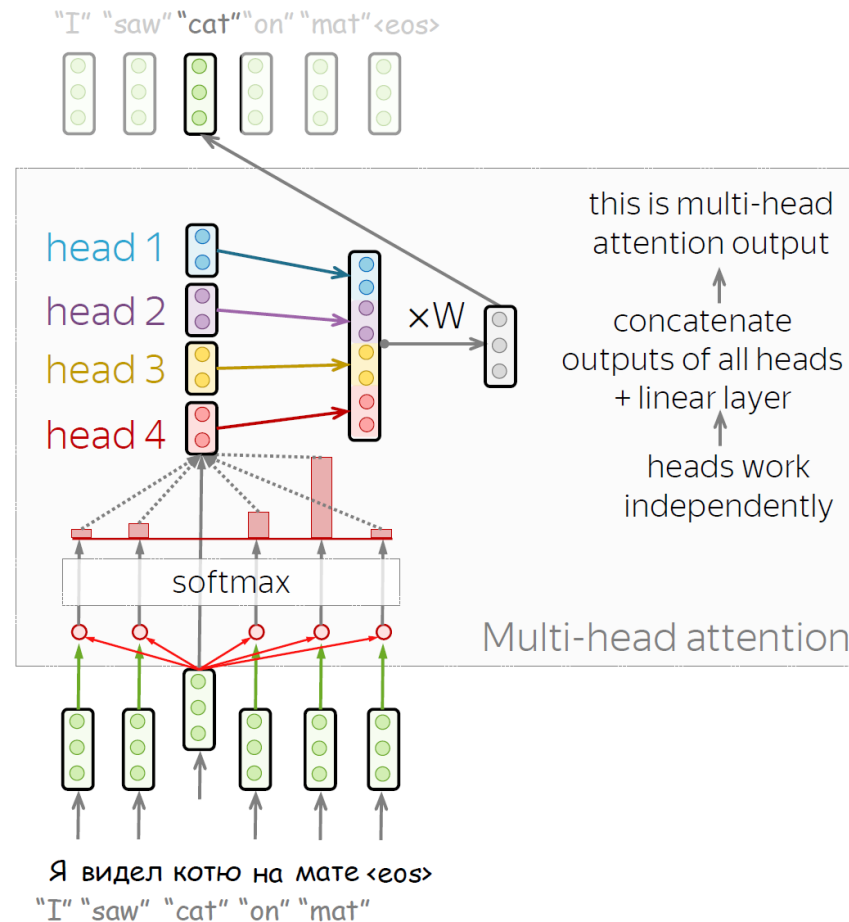
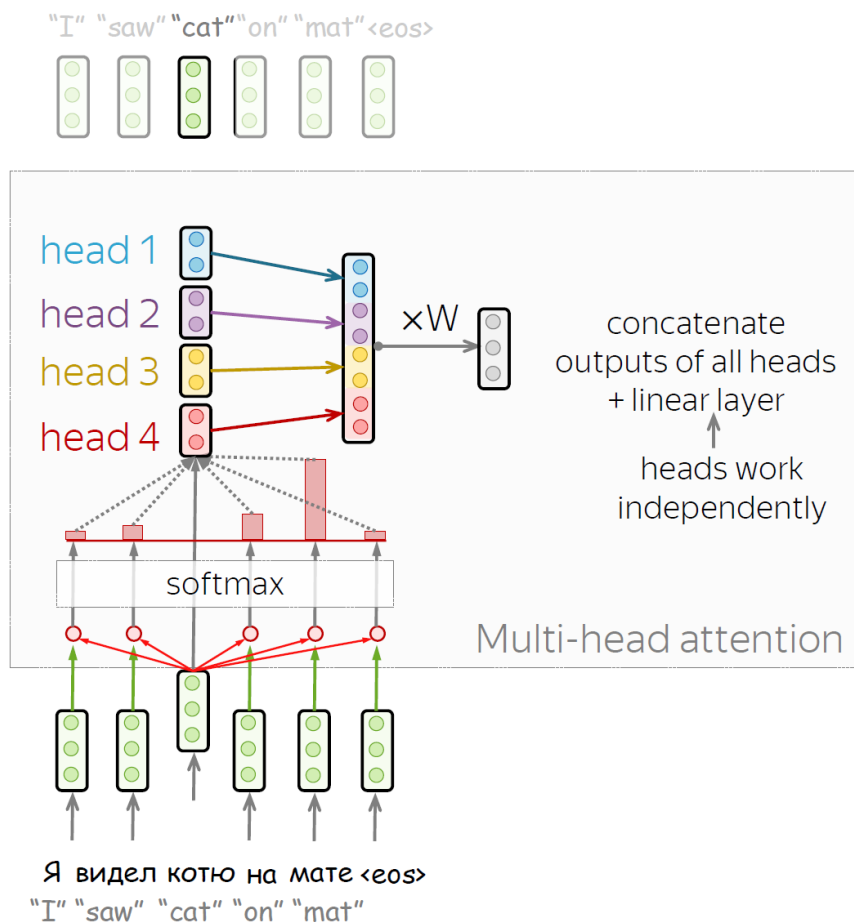
Multi-Head Attention



Multi-Head Attention



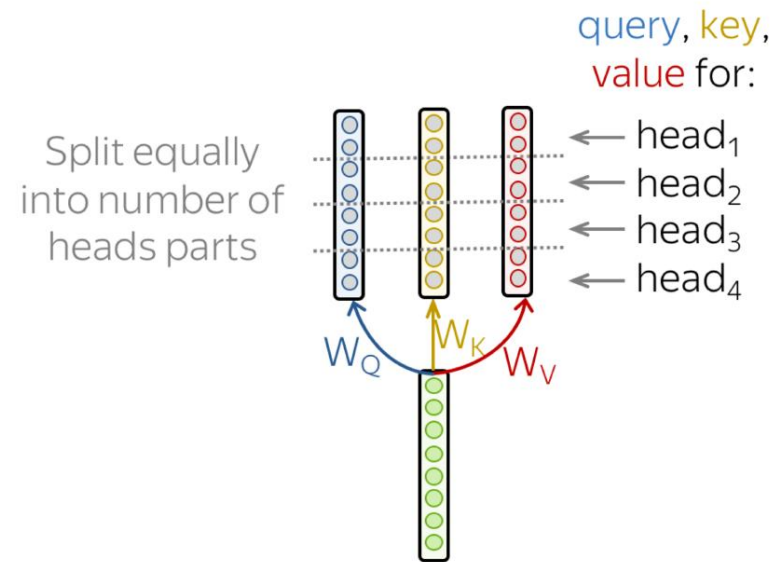
Multi-Head Attention



Multi-Head Attention

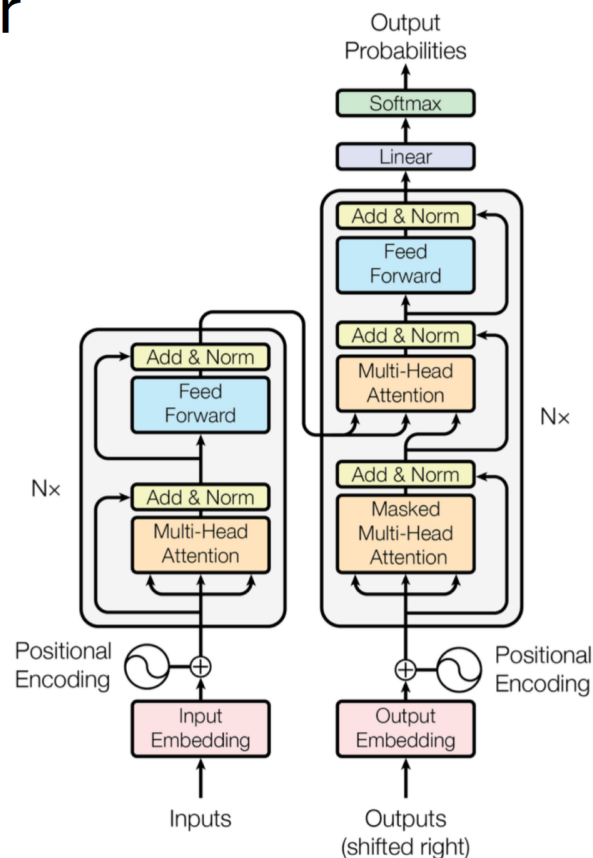
$$\text{MultiHead}(Q, K, V) = \text{Concat}(\text{head}_1, \dots, \text{head}_n)W_o,$$

$$\text{head}_i = \text{Attention}(QW_Q^i, KW_K^i, VW_V^i)$$



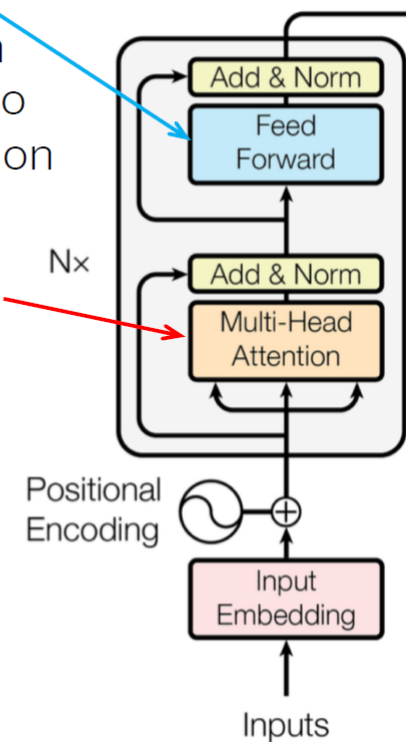
Transformer

Transformer

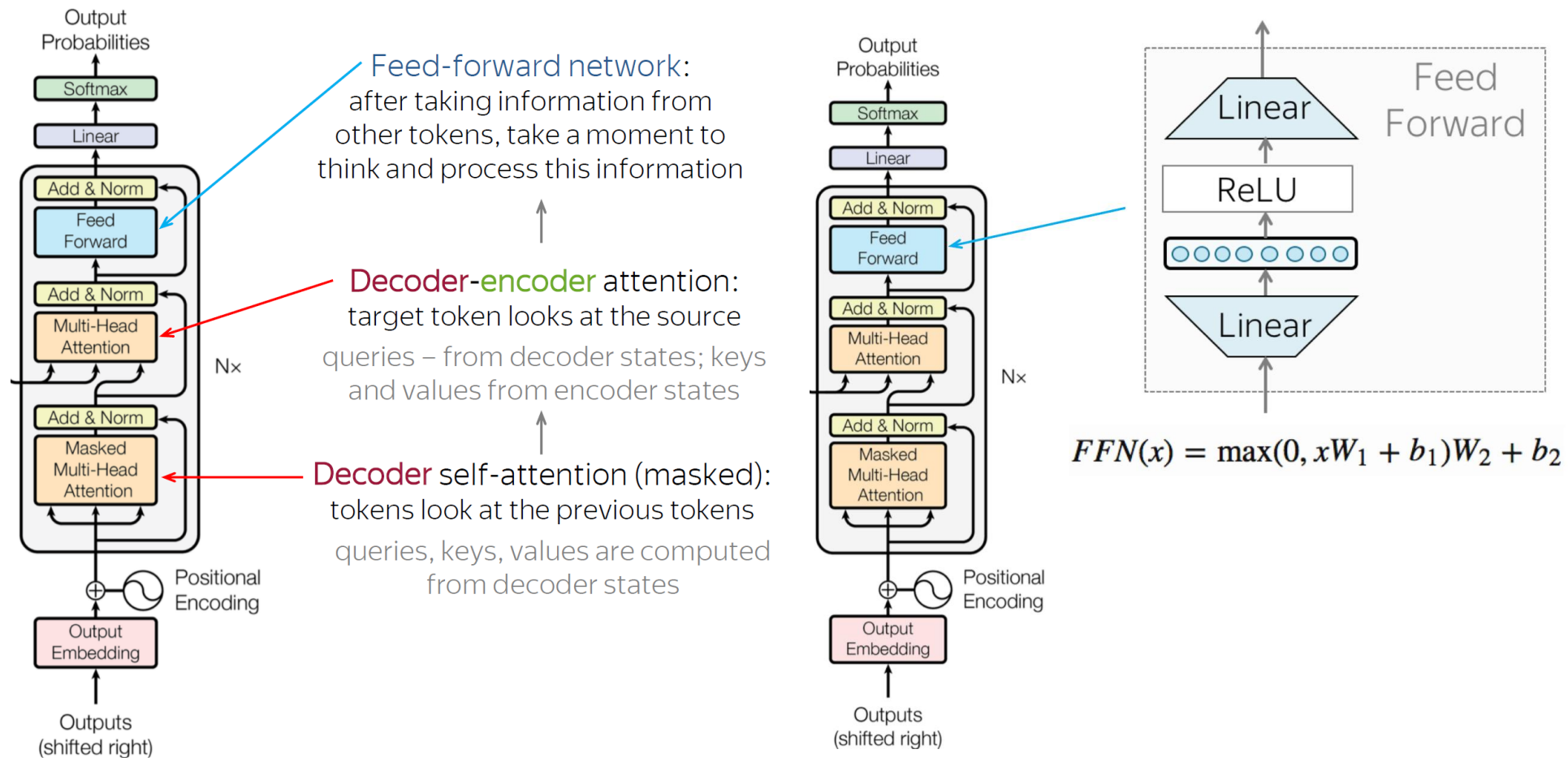


Feed-forward network:
after taking information from
other tokens, take a moment to
think and process this information

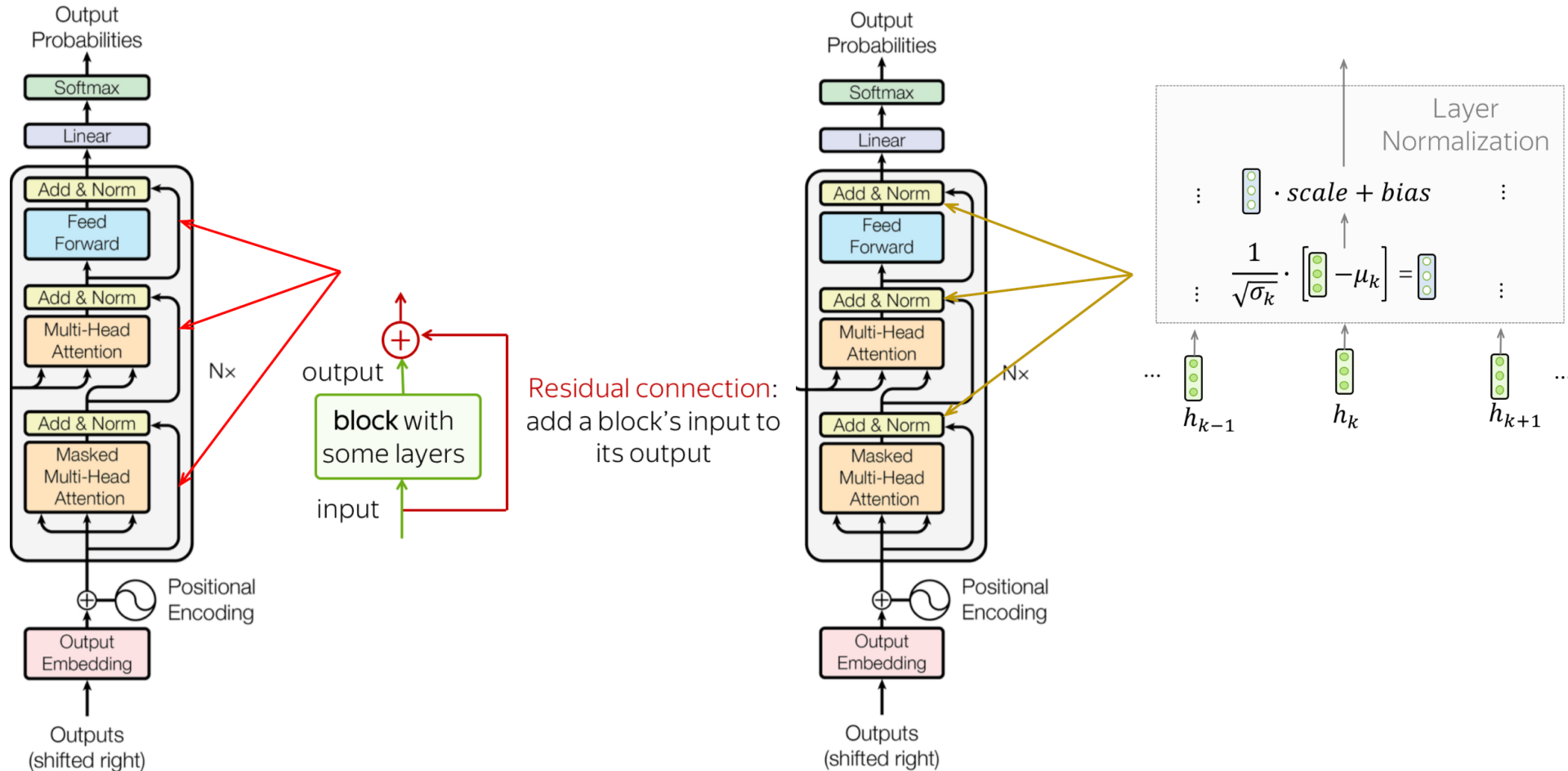
Encoder self-attention:
tokens look at each other
queries, keys, values
are computed from
encoder states



Transformer



Transformer



Позиционное кодирование

Проблема: без повторения или свертки модель ничего не знает о положении

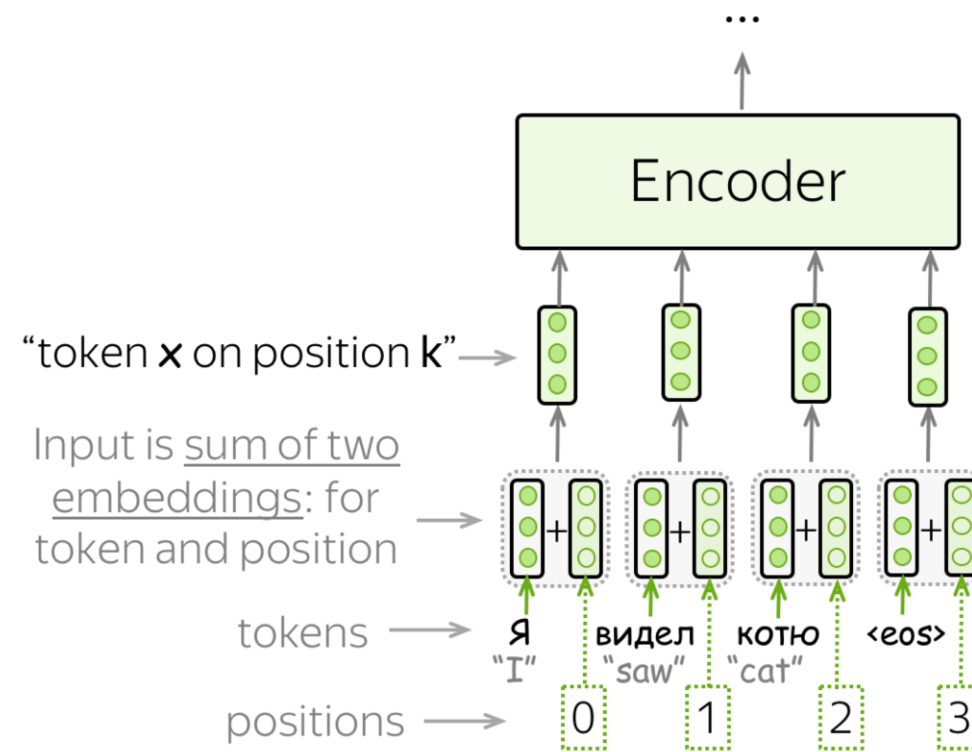
Решение: явно закодировать позицию

Фиксированные кодировки:

$$PE_{pos,2i} = \sin(pos/10000^{2i/d_{model}}),$$

$$PE_{pos,2i+1} = \cos(pos/10000^{2i/d_{model}})$$

pos – позиция, i – размер

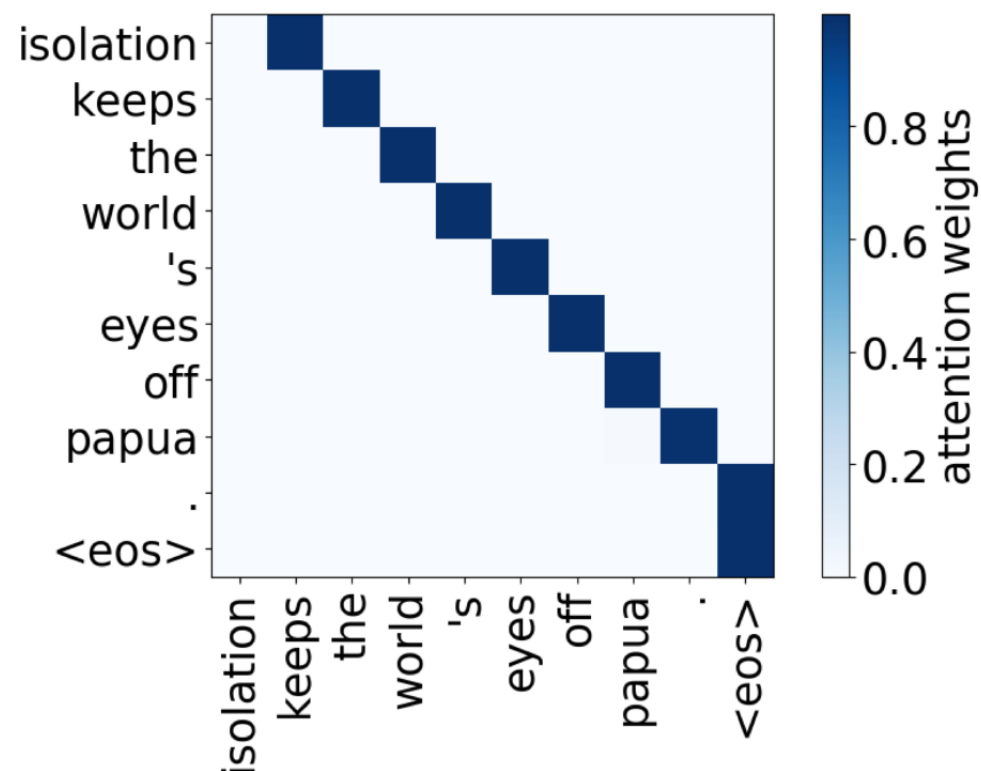
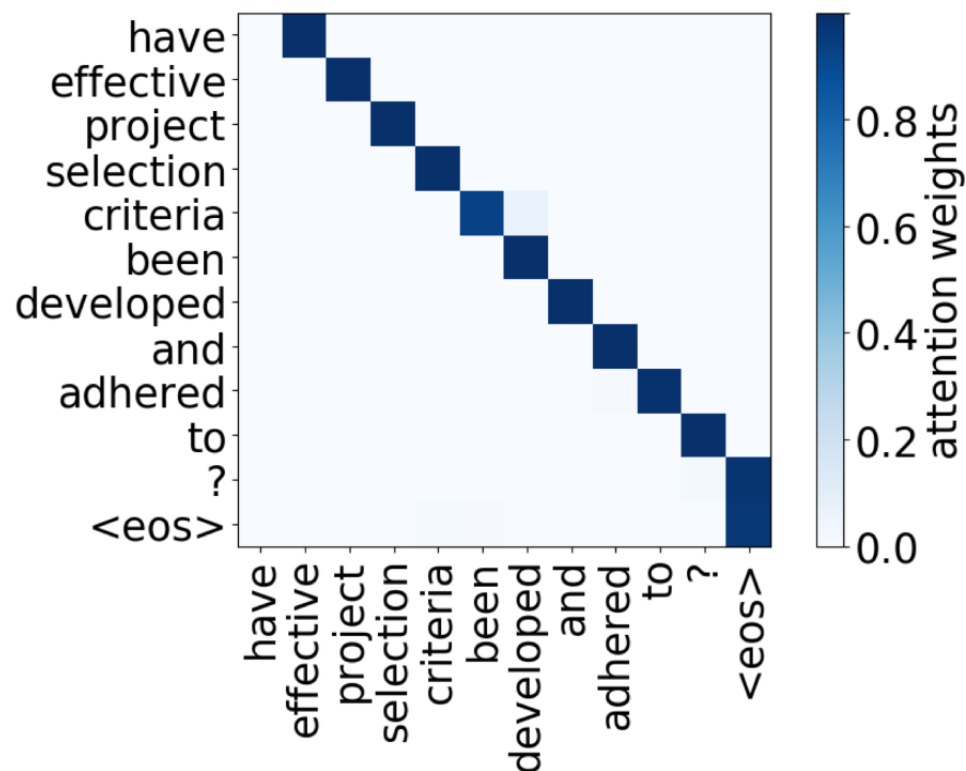


Анализ. Heads в Multi-Head Attention

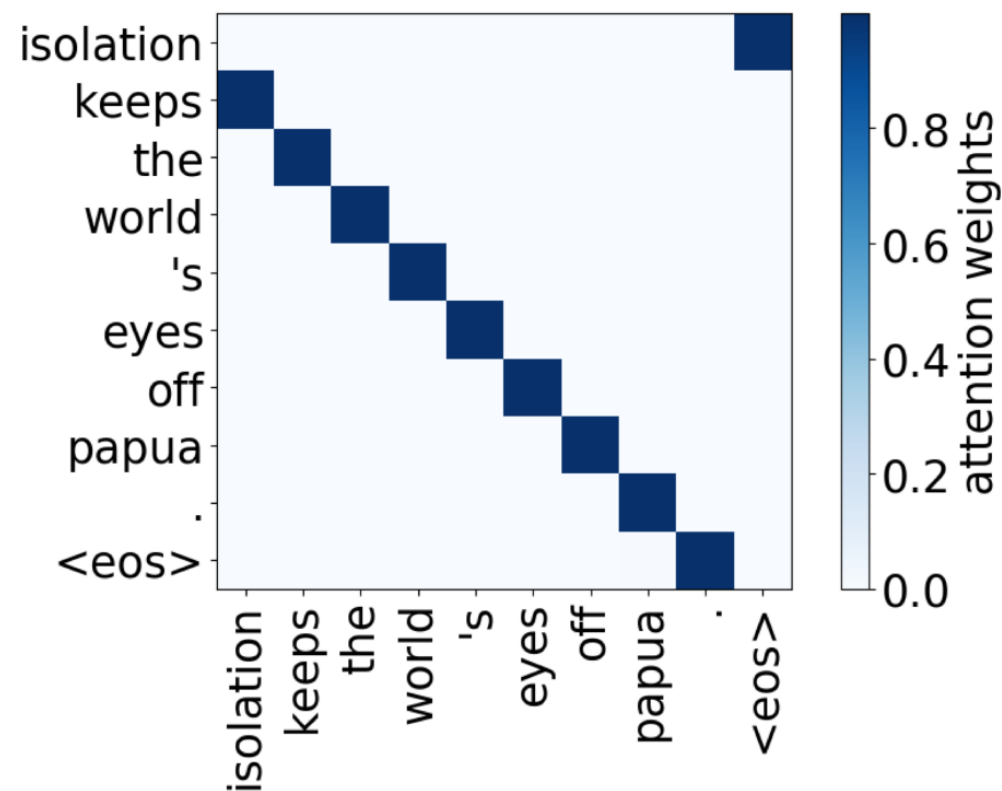
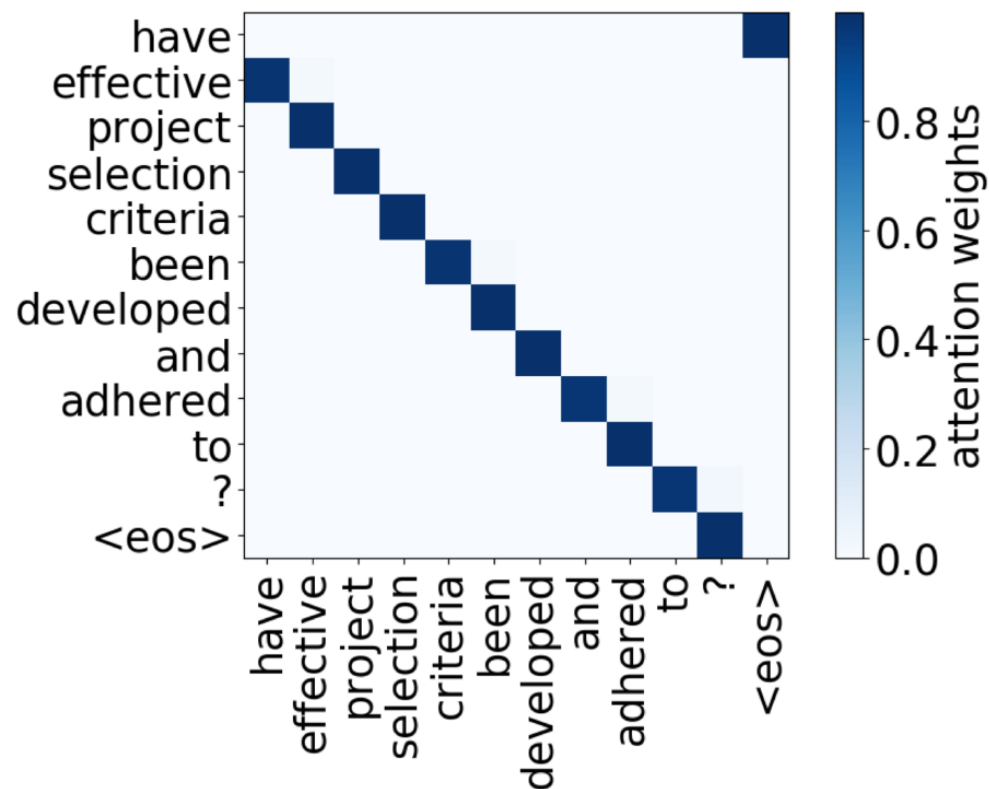
Интерпретируемые роли «голов»:

- Позиционные
- Синтаксические
- Внимание к редким токенам

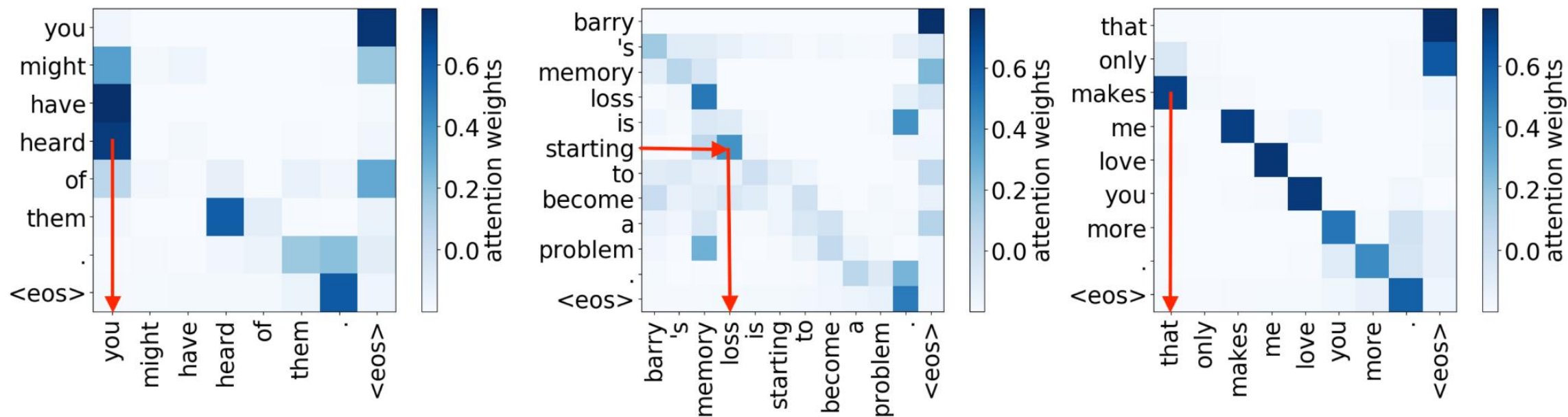
Позиционные головы: внимание к соседям



Позиционные головы: внимание к соседям

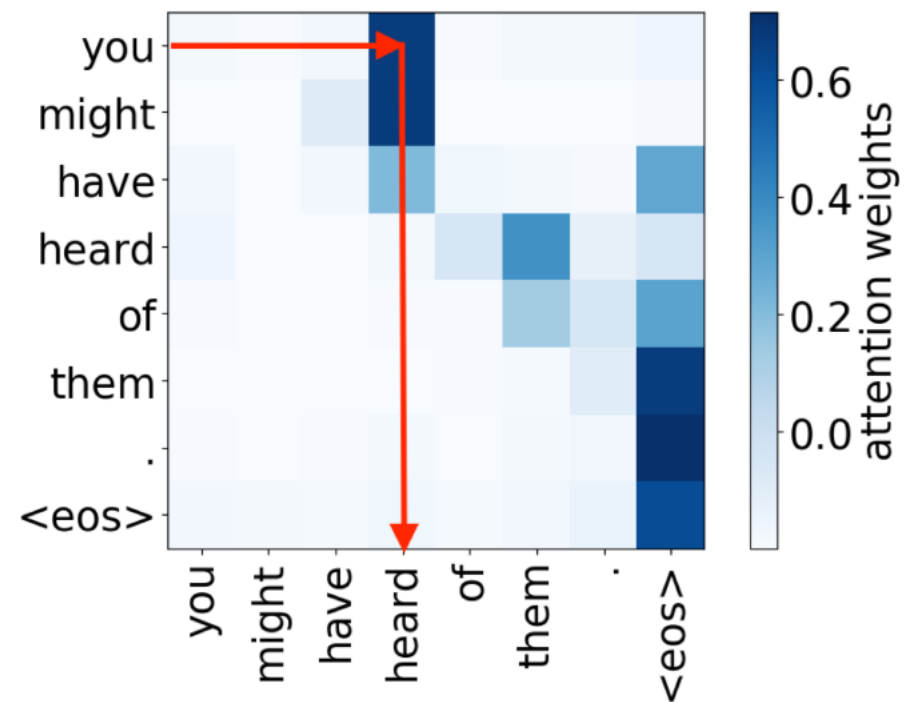
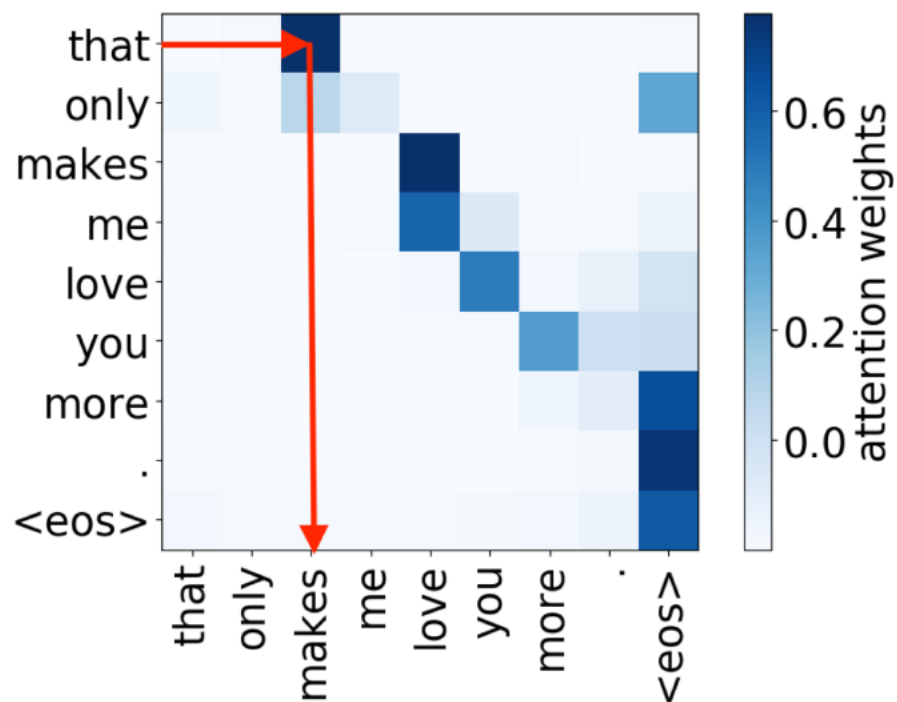


Синтаксические заголовки: отслеживание зависимостей



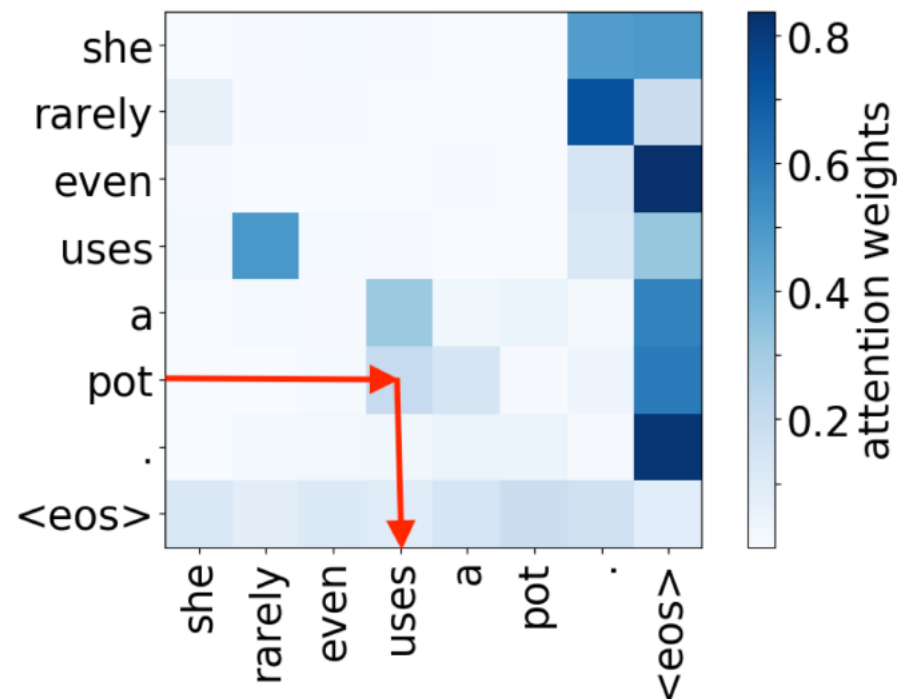
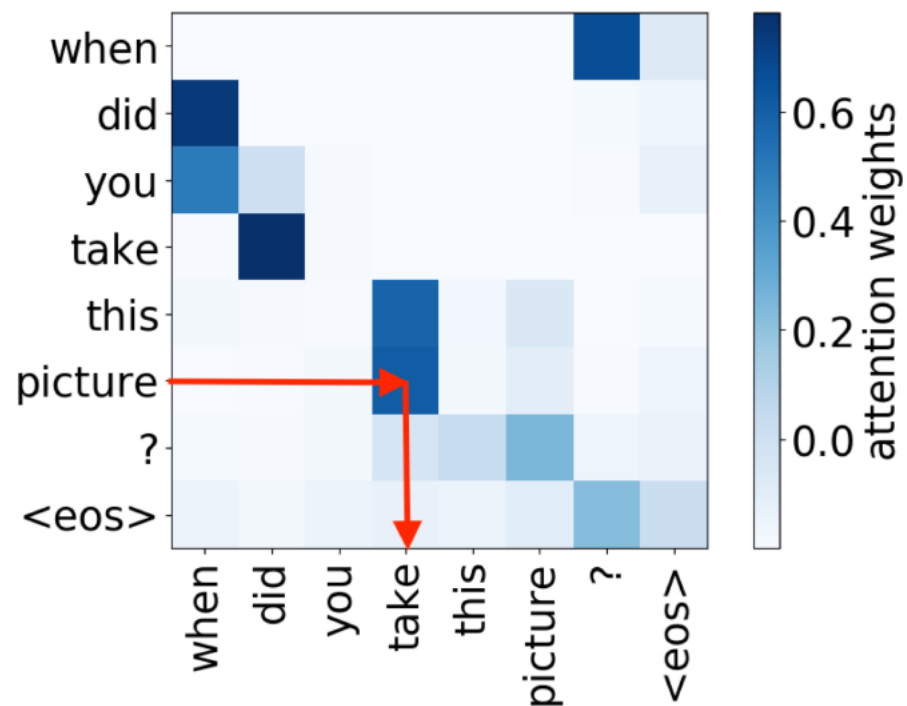
глагол->существительное

Синтаксические заголовки: отслеживание зависимостей



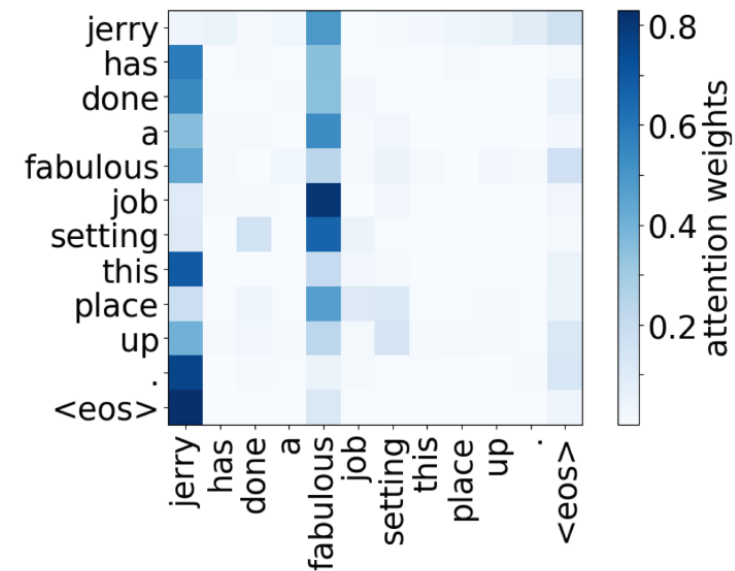
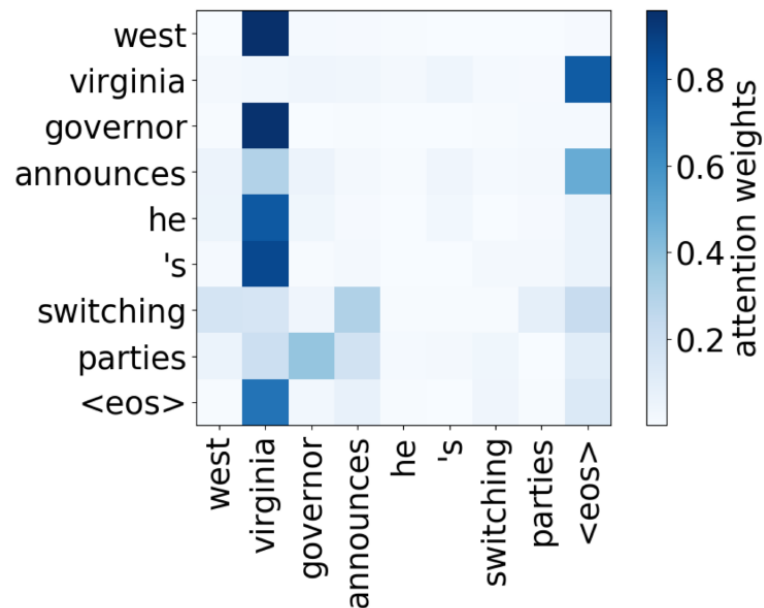
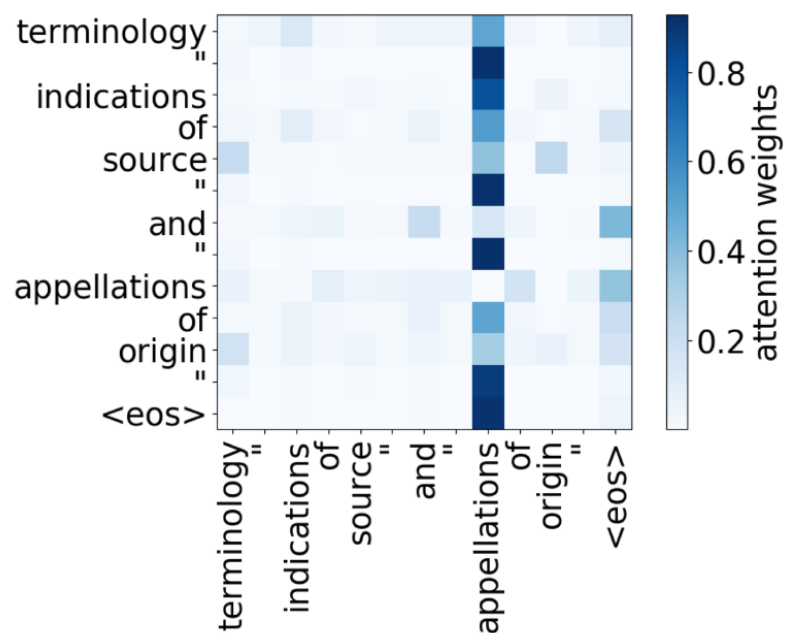
Существительное->глагол

Синтаксические заголовки: отслеживание зависимостей

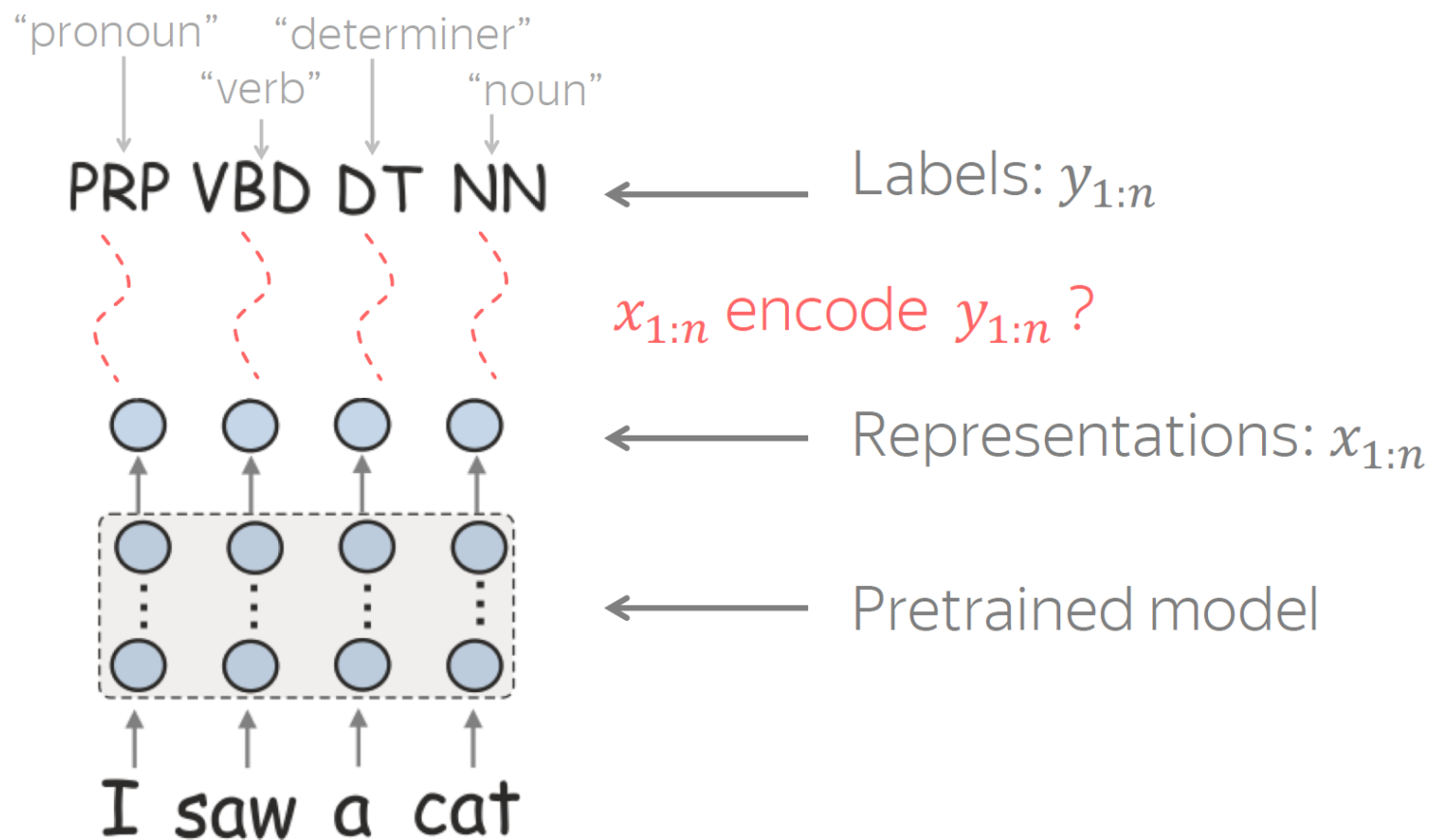


объект->глагол

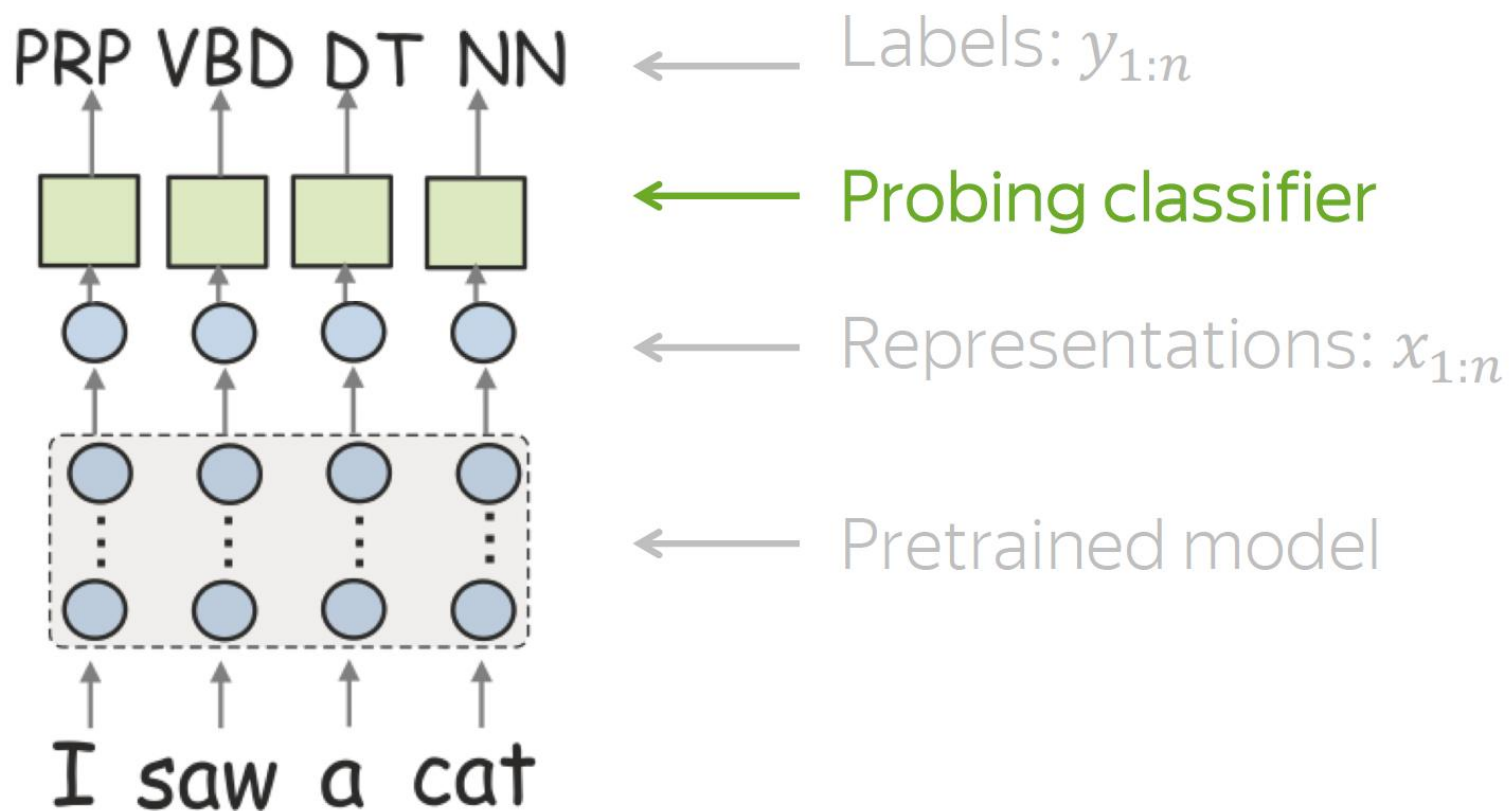
Внимание редким токенам



Как понять, отражает ли модель лингвистическое свойство?

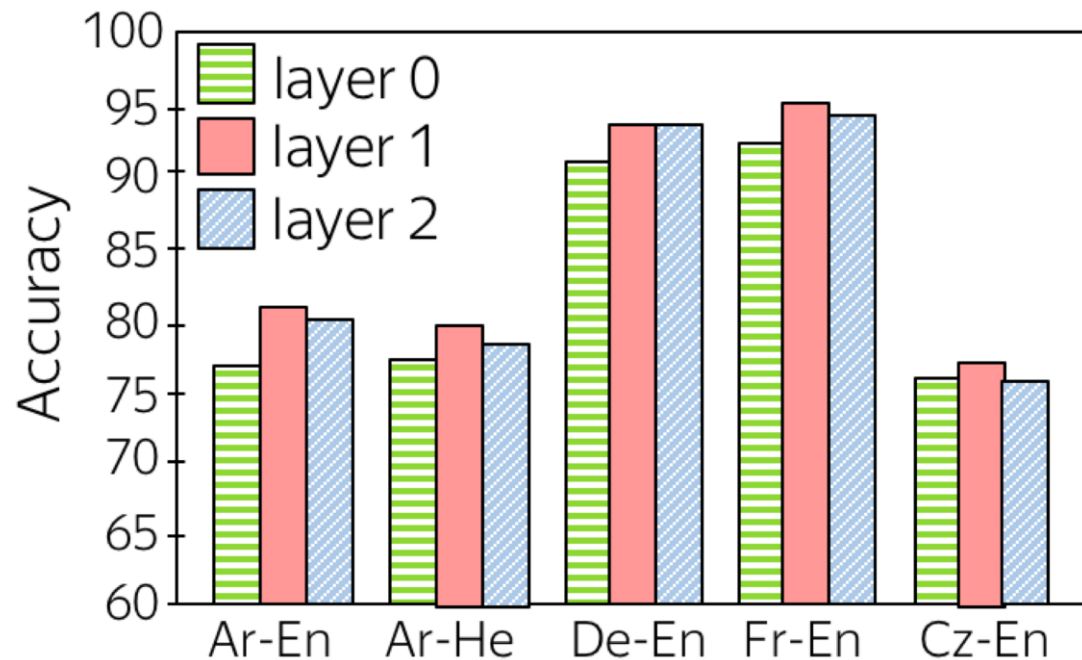


Standard Probing: обучение классификатора, использование его точности



Что модели узнают о морфологии?

POS accuracy by Representation Layer



Effect of Target Language on POS/Morph

