



CMPT 413/713: Natural Language Processing

Pretraining Language Models

Spring 2026
2026-02-09

Some slides adapted from Stanford CS224n and Anoop Sarkar

Pretraining and fine-tuning

Pretraining

- Big pile of unlabeled text data!
- Lots of resources to train!



Helps to build

- Useful representations of language
- Provide good initial parameters for downstream tasks
- Probability distributions that can be sampled from

Supervised fine-tuning

- Annotated data specific (usually small)
- Initialize with pre-trained model



Useful for

- Task / domain specific fine-tuning
- Instruction fine-tuning

Brief History of Pre-training

1960 to 2015

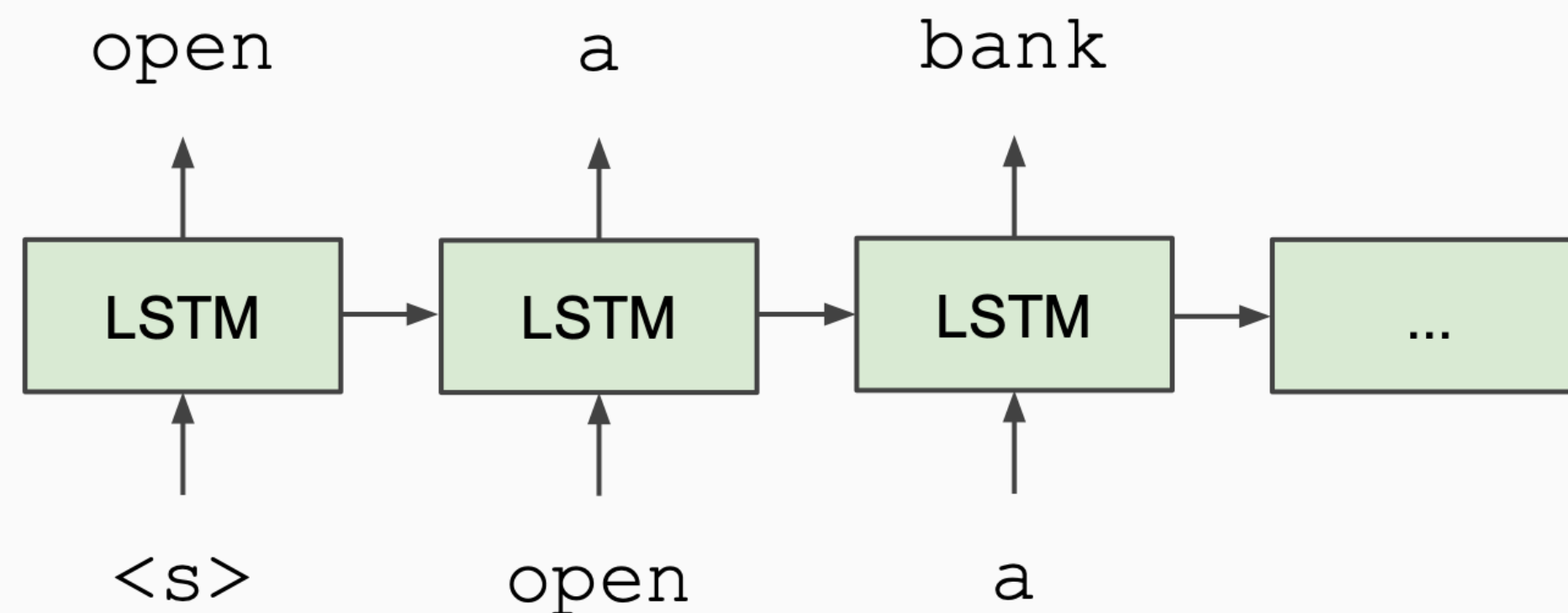
- Singular Value Decomposition (1960s):
 - Take matrix $M \in |V| \times |V|$ of word co-occurrence counts
 - Use SVD to map $M = USV^T$ truncate to $|V| \times k$ initial singular values
 - Use truncated U use as word embeddings.
- Word2Vec/GloVe (2010):
 - Continuous Bag of Words (CBOW) - context words predict target word
 - Skip-gram - target word predicts each context word

Semi-supervised Sequence Learning

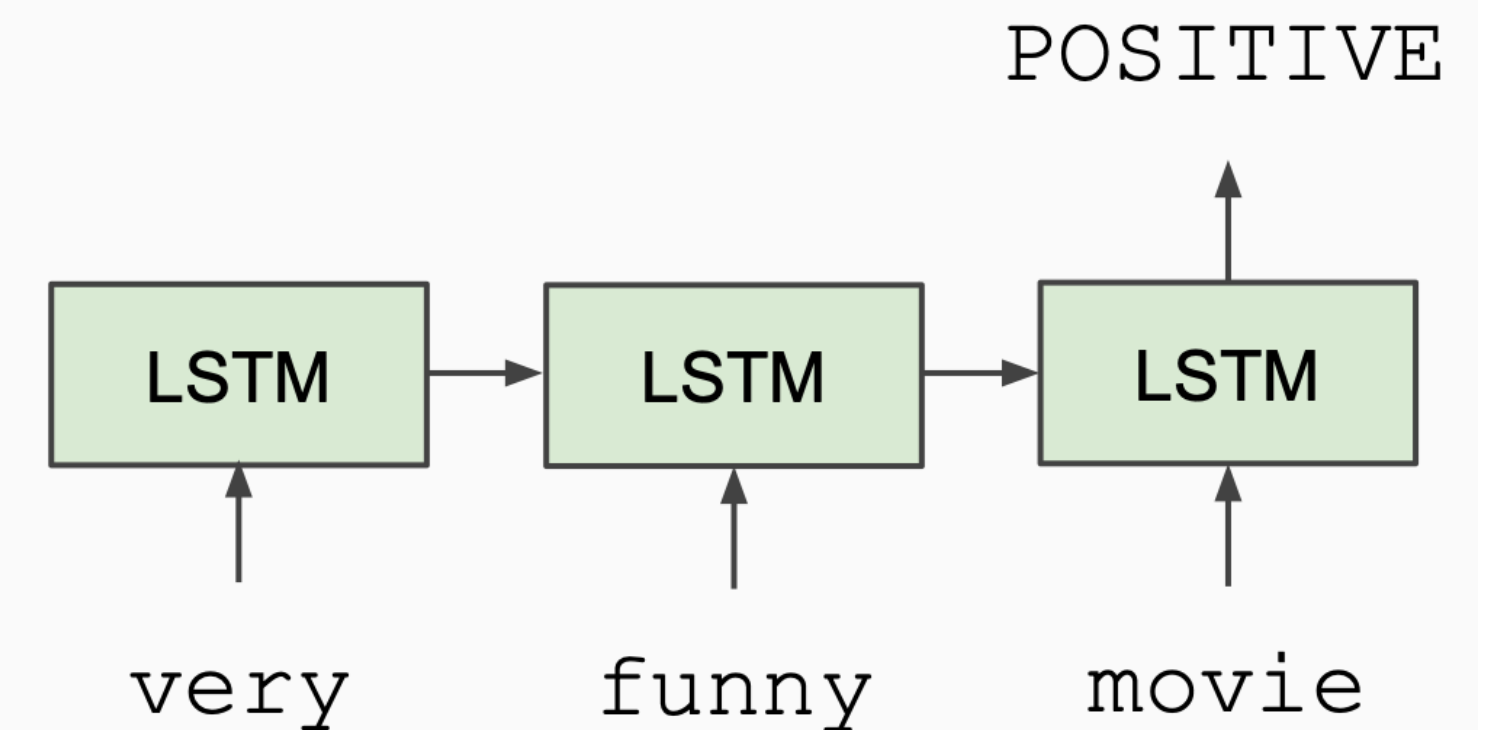
Andrew M. Dai
Google Inc.
adai@google.com

Quoc V. Le
Google Inc.
qvl@google.com

Train LSTM Language Model

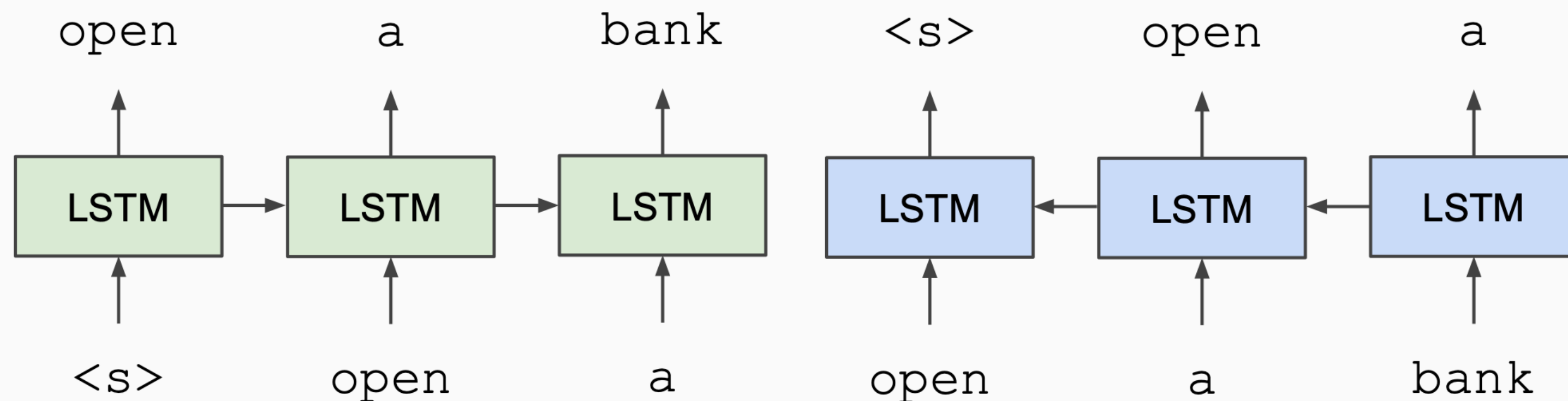


Fine-tune on Classification Task

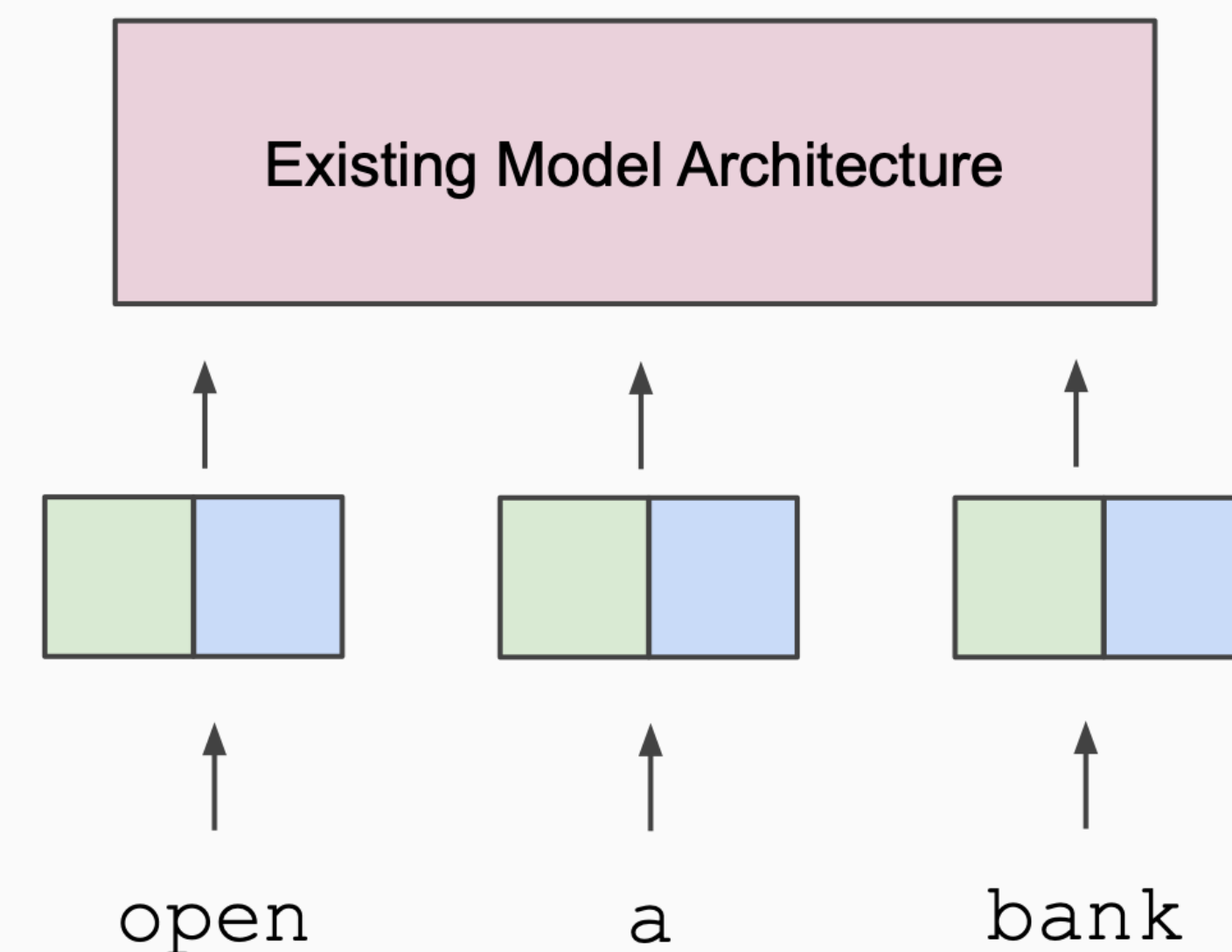


ELMO

Train Separate Left-to-Right and Right-to-Left LMs

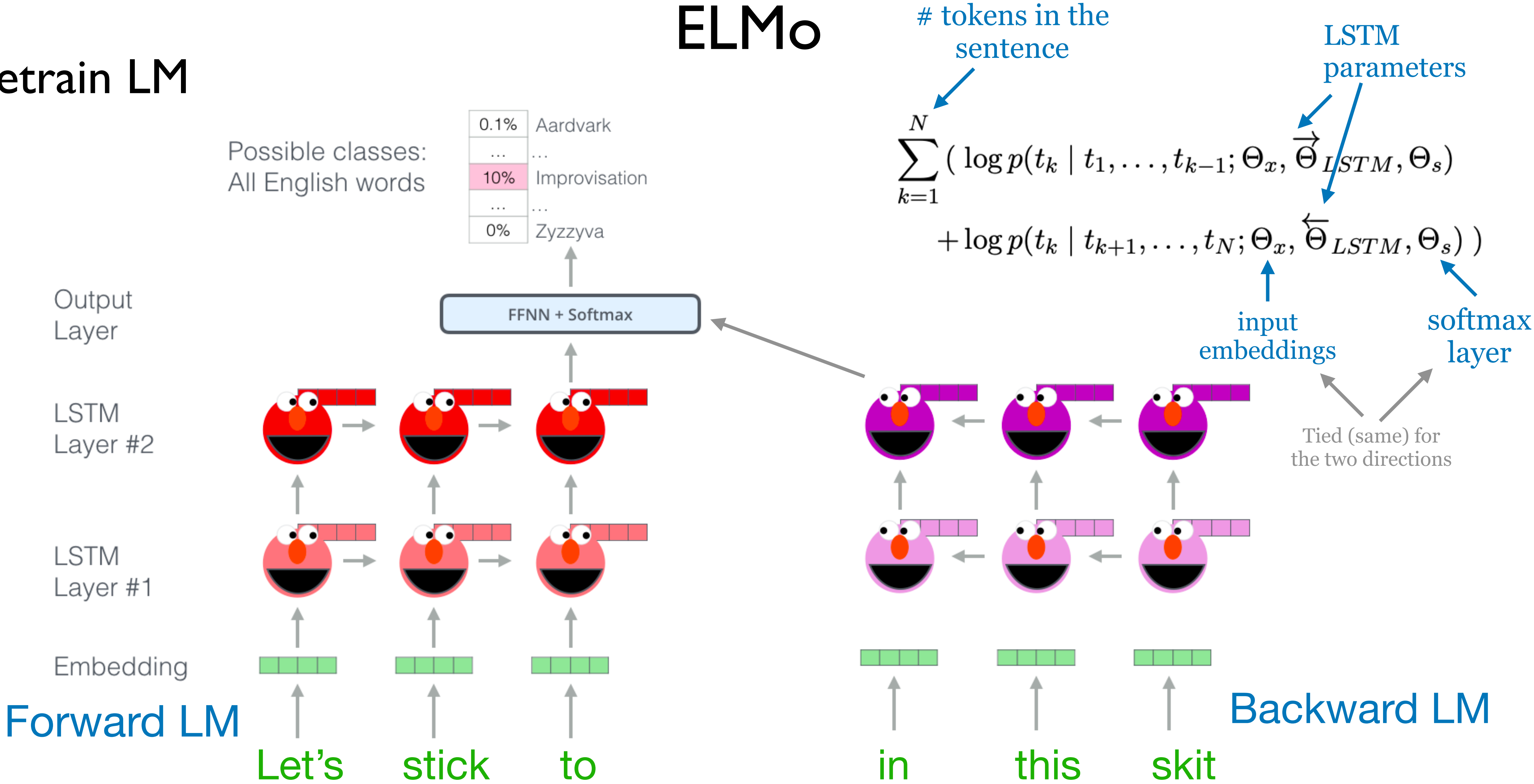


Apply as “Pre-trained Embeddings”



Pretrain LM

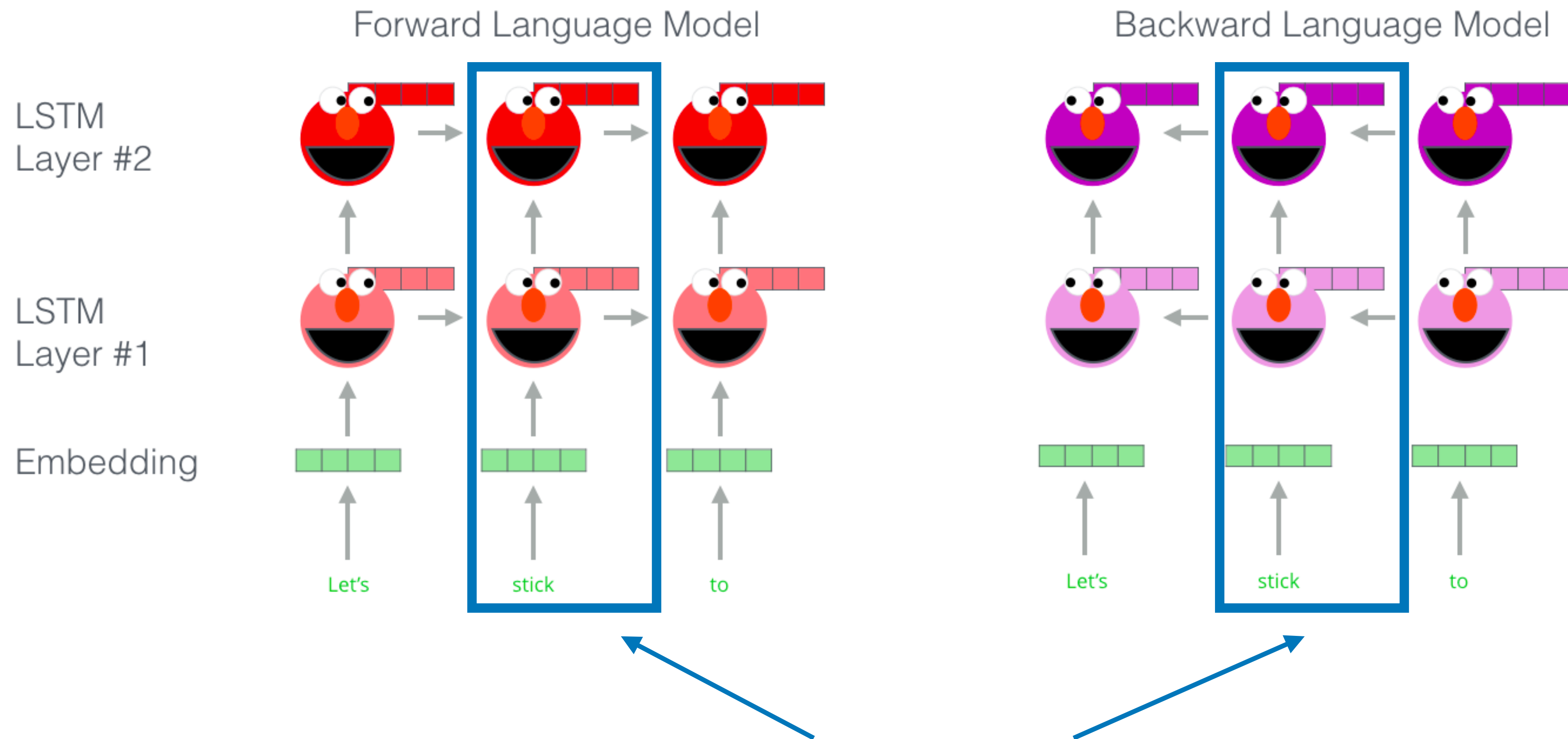
ELMo



(figure credit: Jay Alammam
<http://jalammar.github.io/illustrated-bert/>)

After training LM

ELMo



To get the ELMo embedding of a word ("stick"):

Concatenate forward and backward embeddings
and take weighted sum of layers

ELMo

1- Concatenate hidden layers



2- Multiply each vector by a weight based on the task

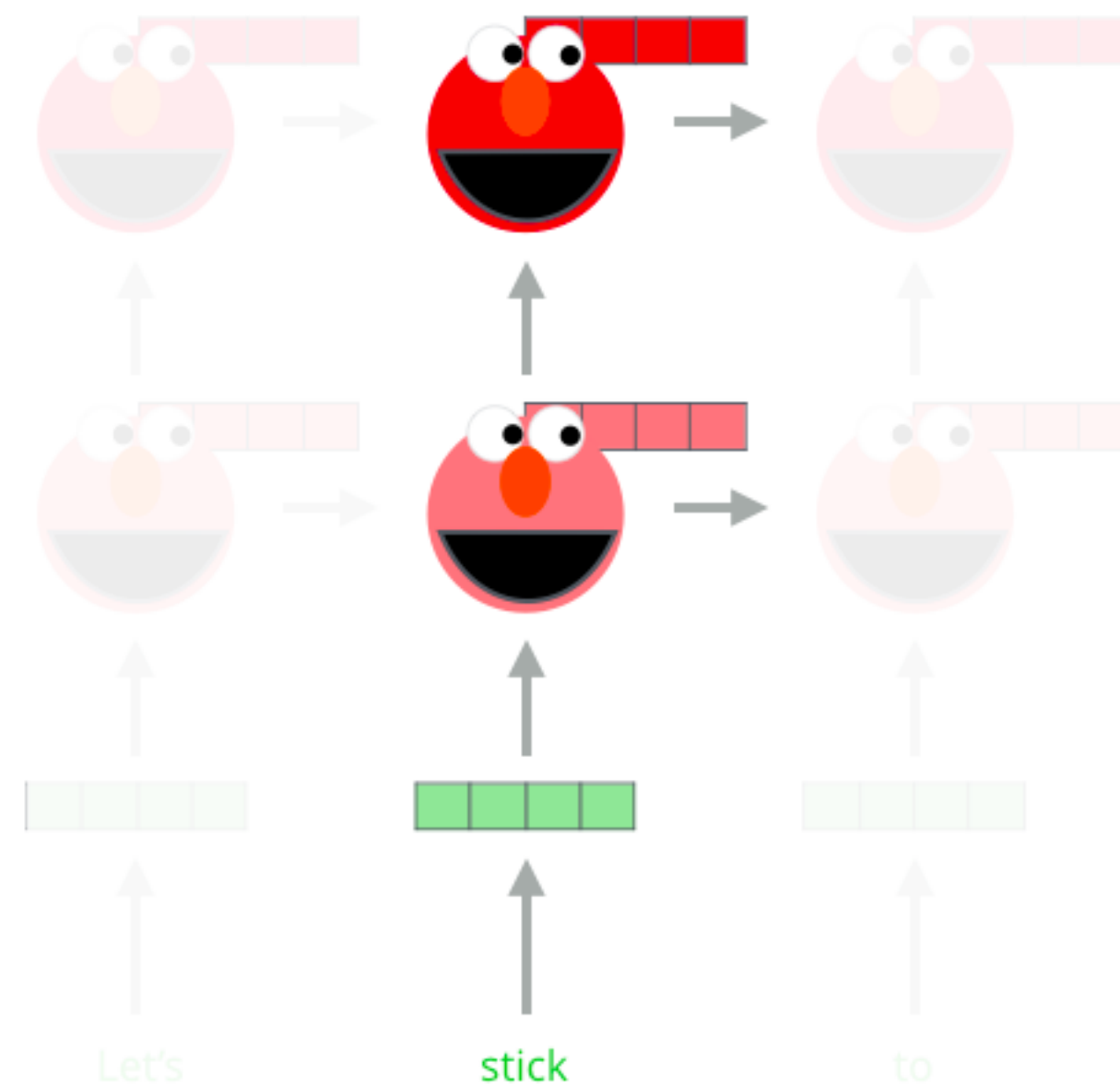


3- Sum the (now weighted) vectors

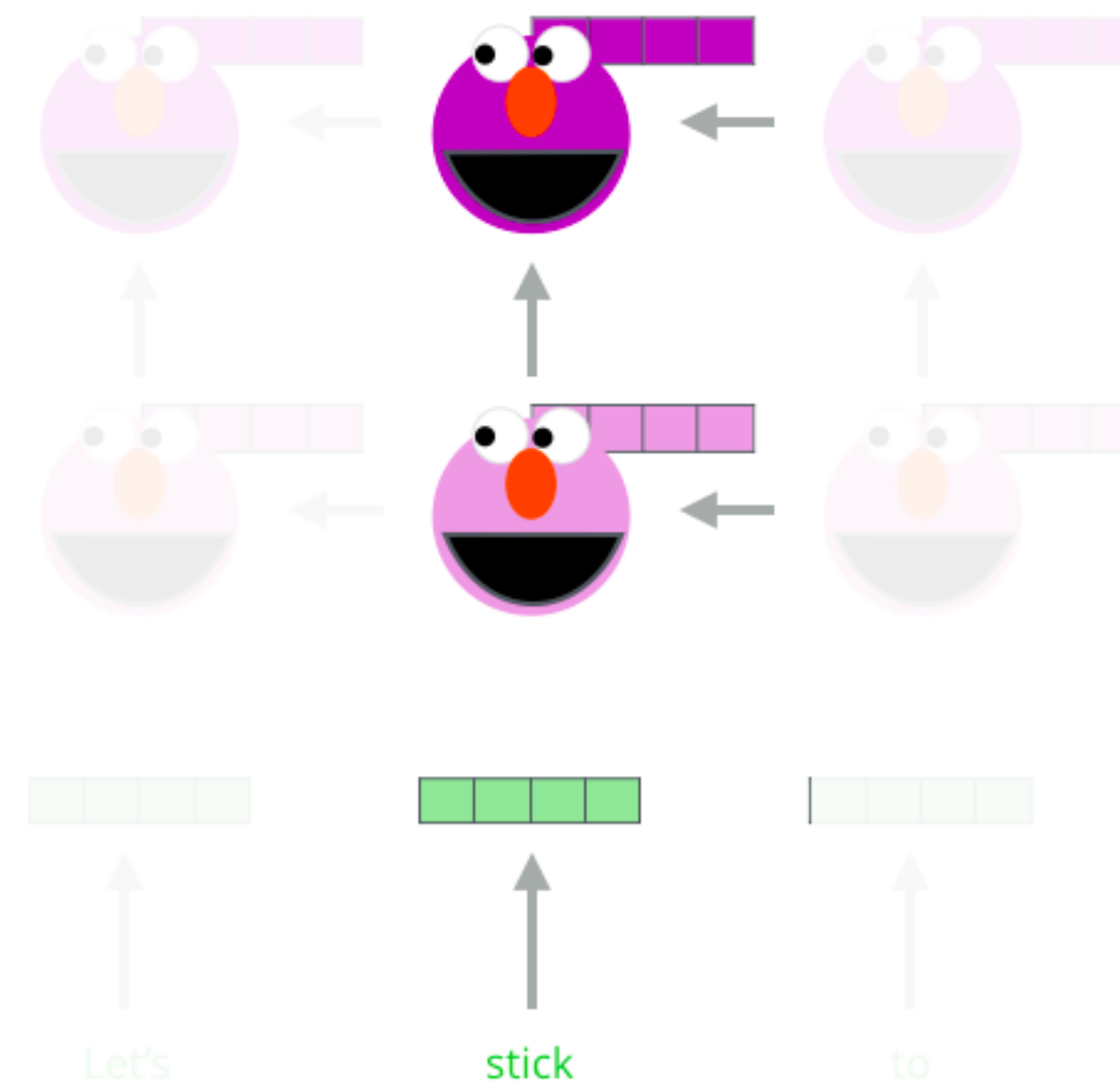


ELMo embedding of "stick" for this task in this context

Forward Language Model



Backward Language Model



LM weights are frozen

Weights s_j are trained on specific task.

To get the ELMo embedding of a word ("stick"):

Concatenate forward and backward embeddings
and take weighted sum of layers

Summary: How to get ELMo embedding?

Input embeddings

Hidden state

$$\begin{aligned} R_k &= \{\mathbf{x}_k^{LM}, \vec{\mathbf{h}}_{k,j}^{LM}, \overleftarrow{\mathbf{h}}_{k,j}^{LM} \mid j = 1, \dots, L\} \leftarrow L \text{ is \# of layers} \\ &= \{\mathbf{h}_{k,j}^{LM} \mid j = 0, \dots, L\}, \end{aligned}$$

Token representation $\rightarrow \mathbf{h}_{k,0}^{LM} = \mathbf{x}_k^{LM}, \mathbf{h}_{k,j}^{LM} = [\vec{\mathbf{h}}_{k,j}^{LM}; \overleftarrow{\mathbf{h}}_{k,j}^{LM}] \leftarrow \text{hidden states}$

$$\mathbf{ELMo}_k^{task} = E(R_k; \Theta^{task}) = \gamma^{task} \sum_{j=0}^L s_j^{task} \mathbf{h}_{k,j}^{LM}$$

Task specific learnable parameters $\rightarrow$

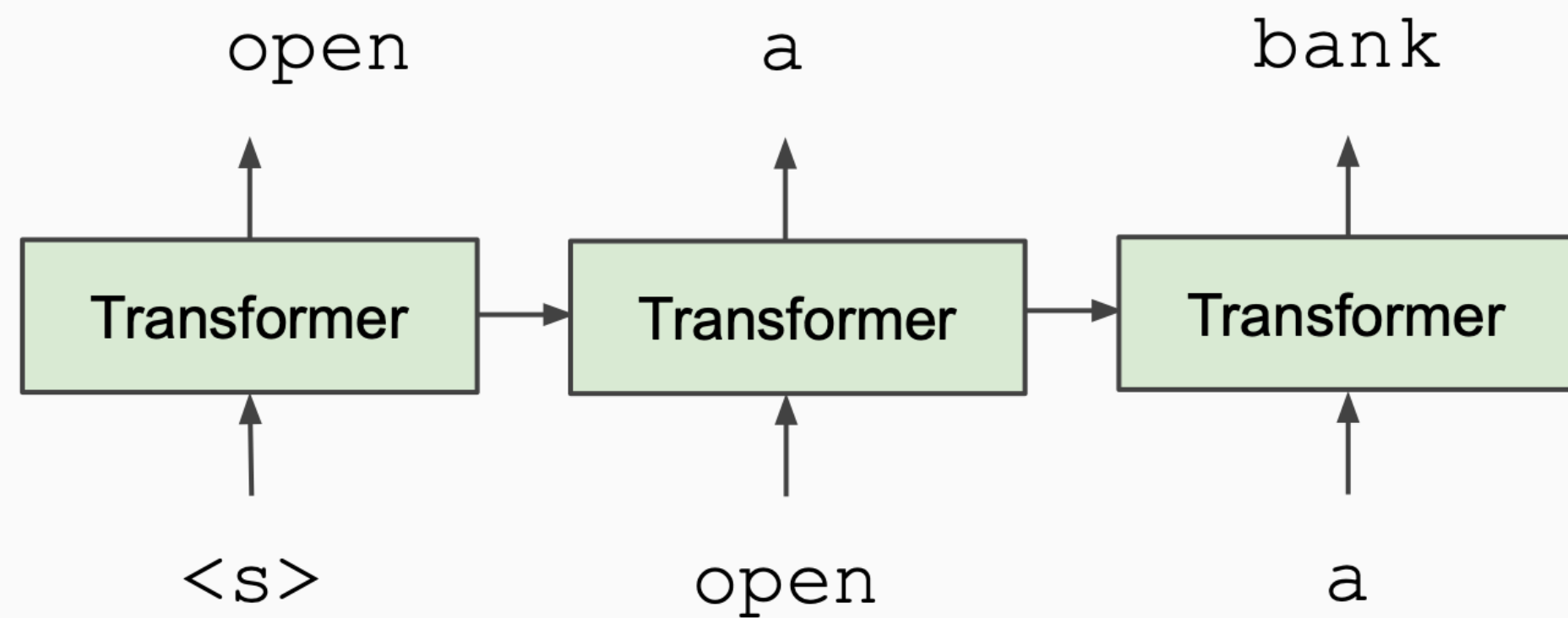
- γ^{task} : allows the task model to scale the entire ELMo vector
- s_j^{task} : softmax-normalized weights across layers
- **To use:** plug ELMo into any (neural) NLP model: freeze all the LMs weights and change the input representation to:

$$[\mathbf{x}_k; \mathbf{ELMo}_k^{task}]$$

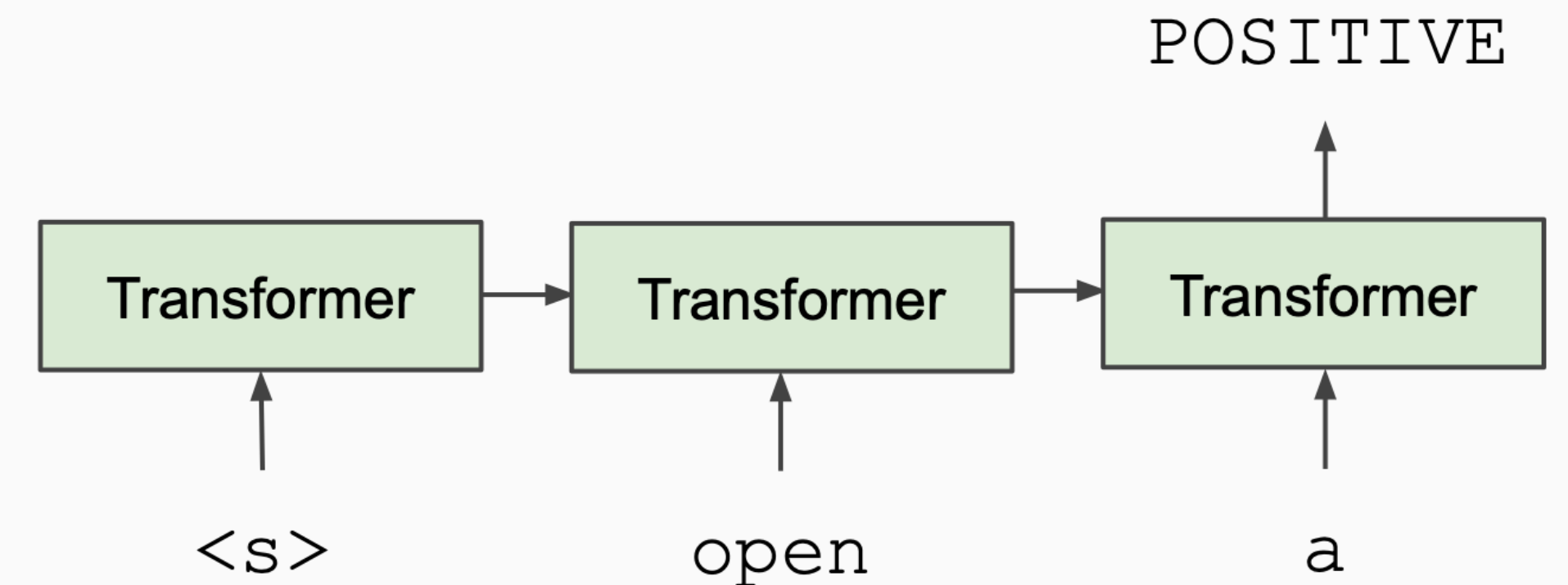
(could also insert into higher layers)

GPT1

Train Deep (12-layer) Transformer LM



Fine-tune on Classification Task



GPT models



GPT

- Improving language understanding by generative pre-training [Radford et al, 2018]
- Large language model with transformers with supervised fine-tuning
 - different model for each task
- Trained on BooksCorpus (800M words), 117M parameters (12 layers)

GPT-2

- Language Models are Unsupervised Multitask Learner [Radford et al, 2019]
- Model all tasks as sequence completion with special tokens indicating task
- Trained on WebText (40B words), 1.5B parameters (48 layers)
- No fine-tuning, demonstrated few-shot learning

GPT-3

- Language Models are Few-Shot Learners [Brown et al, 2020]
- Trained on Web+Books+Wikipedia (300B words), 175B parameters (96 layers)
- Demonstrated zero-shot and few-shot prompting abilities

GPT models (after GPT-3)

InstructGPT and GPT-3.5 [2022]

- Align responses to human feedback
- Instruction fine-tuning
- Reinforcement learning from human feedback
- Used in initial ChatGPT

- Supervised fine-tuning on human conversations
- Data where human will pretend to be user or AI assistant

GPT-4 [March 2023]

- Multimodal with **images** and text (GPT-4V)
- Larger, better model (estimated 1.7 trillion parameters)
- Turbo [Nov 2023] - longer context (128K)

- Human rank generated output
- Use reinforcement learning to improve generation

GPT-4o (omni) [May 2024]

- Multimodal with **audio**, images and text (GPT-4V)
- Real-time processing and generation

o1 [September 2024], o3 [mini - January 2025] - Reasoning

Training stages for modern LLMs

Piles of
unlabeled
text!



Self-supervised training
LM objective

$$p(w_t | w_{<t})$$

Pre-training

LM training on large, large
amount of data

Pre-training can be
broken into stages (mid-
training)

Text with
“instructions”
and
“responses”

List three fruit
Apple, orange,
banana

Supervised training
LM objective

$$p(w_t | w_{<t}; \text{prompt})$$

Instruction-tuning

Supervised fine-tuning for
instructions

Human
preference data
Data about what
people prefer



Reinforcement
learning

Preference optimization

Align to human
preferences

Post-training
(Can have more iterations)

LLM performance depends on

- Model architecture
- Training strategy
- Training objective
- Training data

Pretraining language models

- Model (Neural Architecture)
 - Does it use FFN, RNN (LSTM, GRU), or Transformer?
 - Is it an **encoder**-based, **decoder**-based, or **encoder-decoder** model?
 - Specifics of the neural architecture (number of layers, embedding size, etc)
- Dataset
 - What is the data that is used to pretrain the model?
- Training objective
 - What is the training objective?
- Other details
 - Tokenization: what tokenization is applied?
 - Implementation and training details?

Development of Open LLMs

Closed LLMs

- **GPT (OpenAI)**
- Claude (Anthropic)
- Gemini (Google)

Open weights

- **LLaMa (Meta)**
- DeepSeek
- Mistral (Mistral AI)
- Qwen (Alibaba)
- Gemma (Google)

Open weights + data

- **OLMo (AI2)**
- DCLM
- Amber
- BLOOM
- Pythia

Open weights + partial data

- StableLM
- Zamba
- Falcon

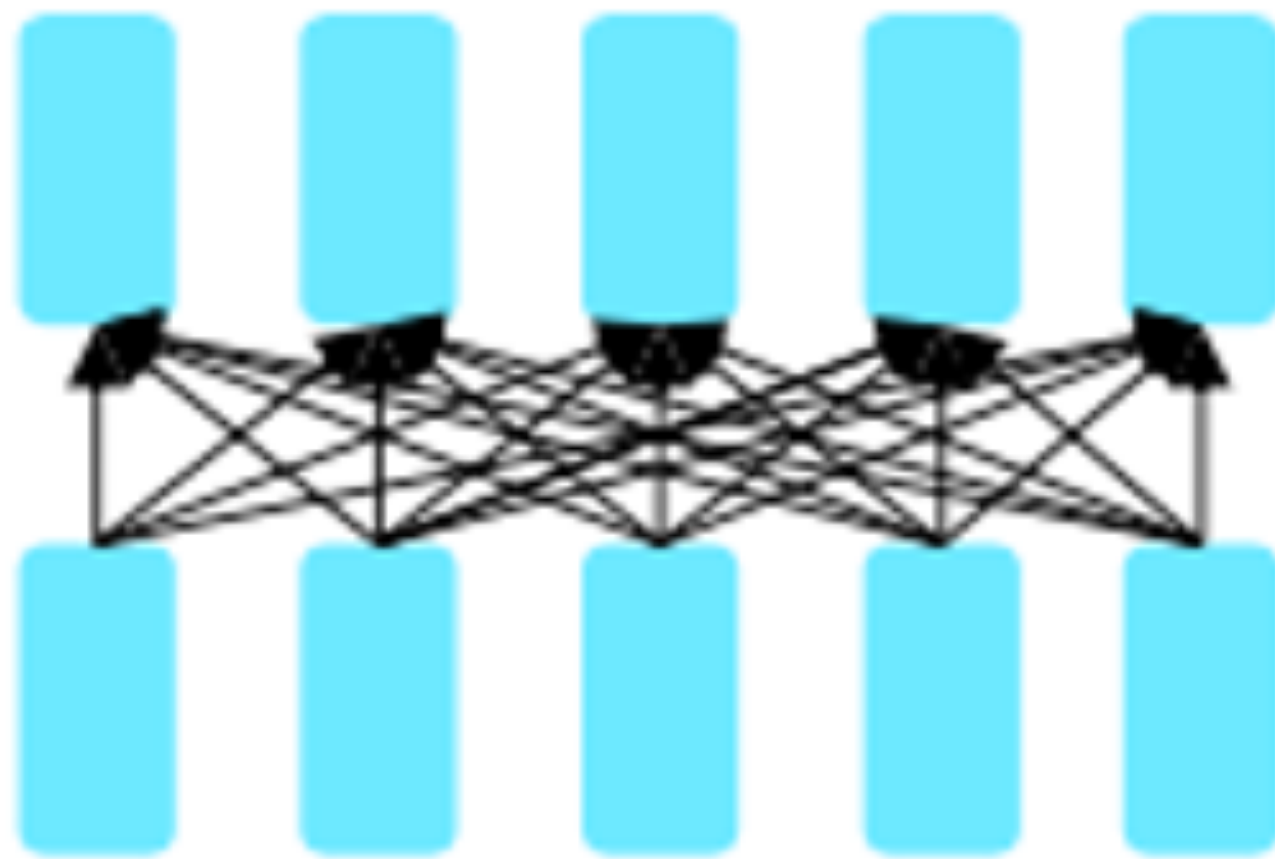
Pre-training Transformers

Representation Learning

Transformers for pretraining

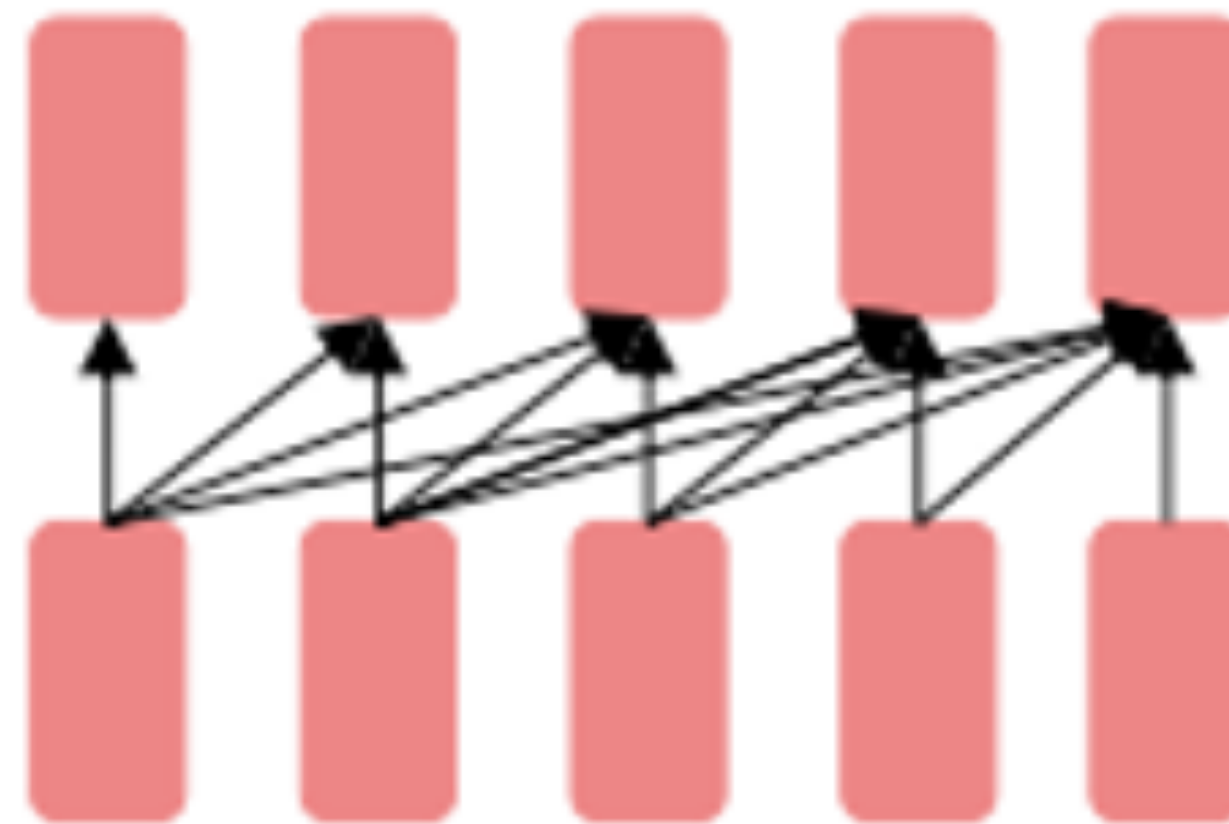
- Self-supervised Transformer based models shattered language understanding benchmarks in NLP in 2018.
- Trained on large text corpus with self-supervised objectives and then transferred.

Encoder only



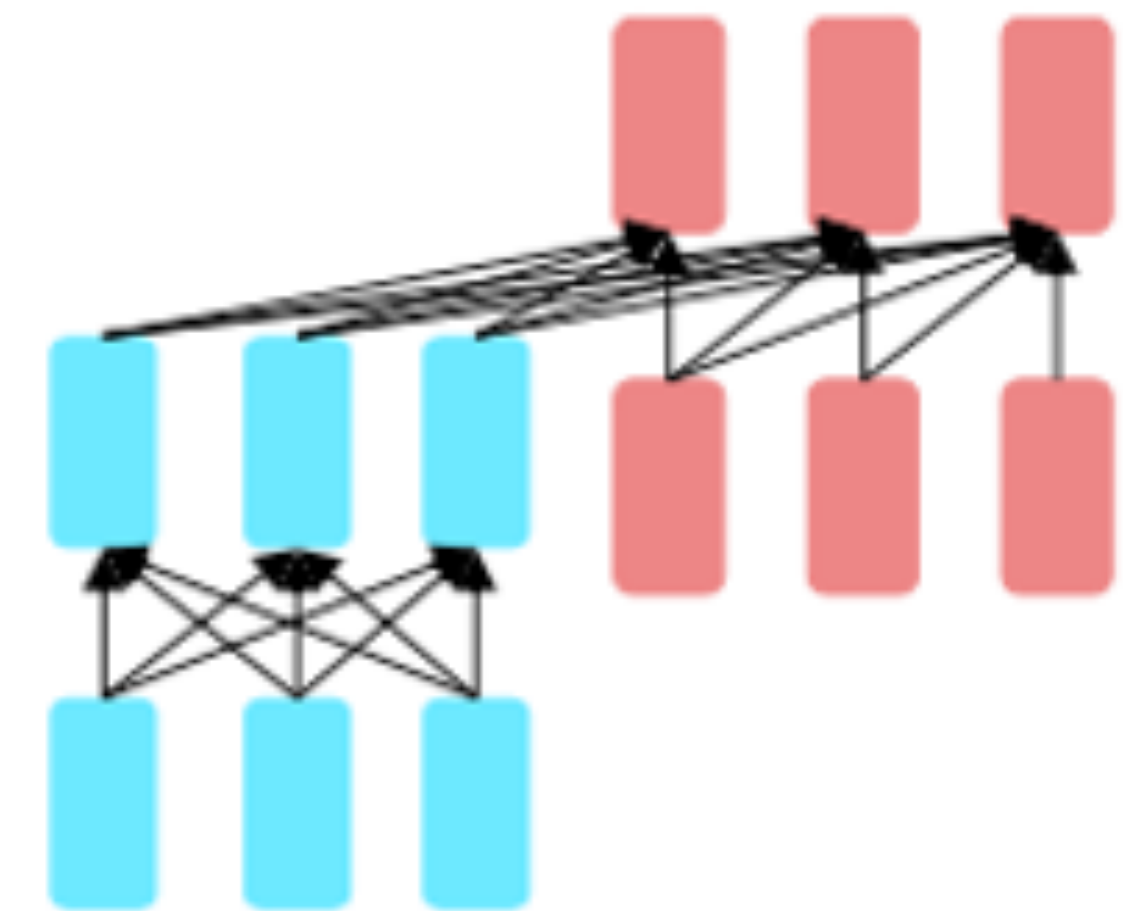
- Masked language models
- Bidirectional context
- BERT + variants (e.g. RoBERTa)
-

Decoder only



- Language models
- Can't condition on future words, good for generation
- GPT, LLaMa, PaLM

Encoder-Decoder

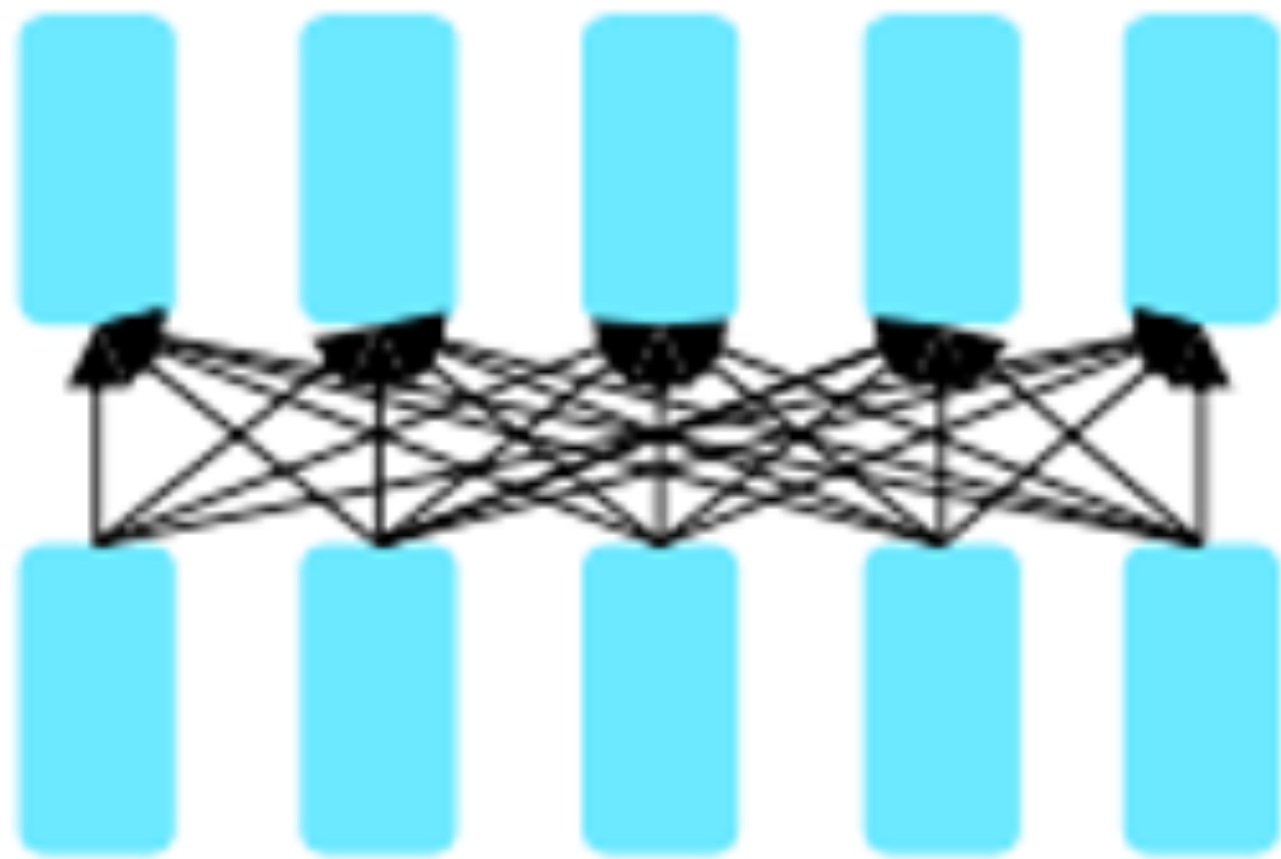


- Combine benefits of both
- Original Transformer, UniLM, BART, T5

Transformers for pretraining

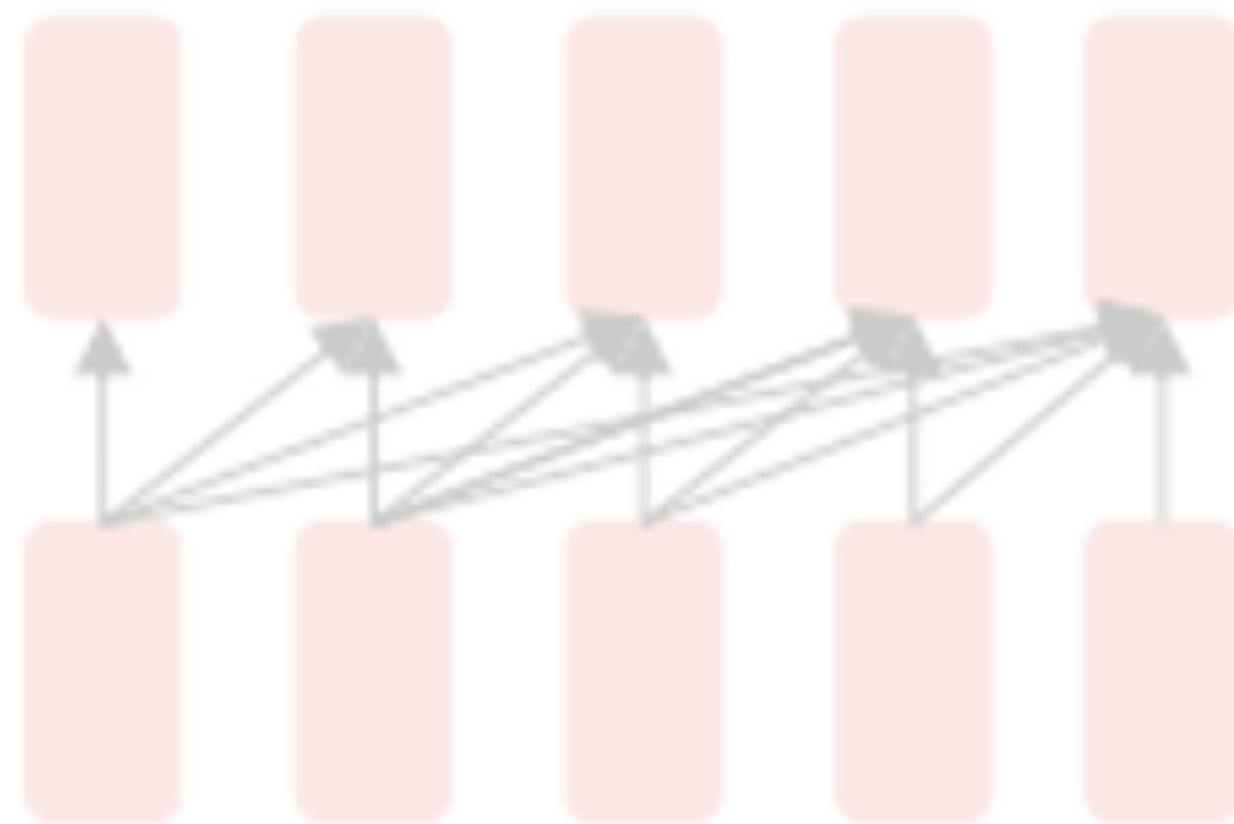
- Self-supervised Transformer based models shattered language understanding benchmarks in NLP in 2018.
- Trained on large text corpus with self-supervised objectives and then transferred.

Encoder only



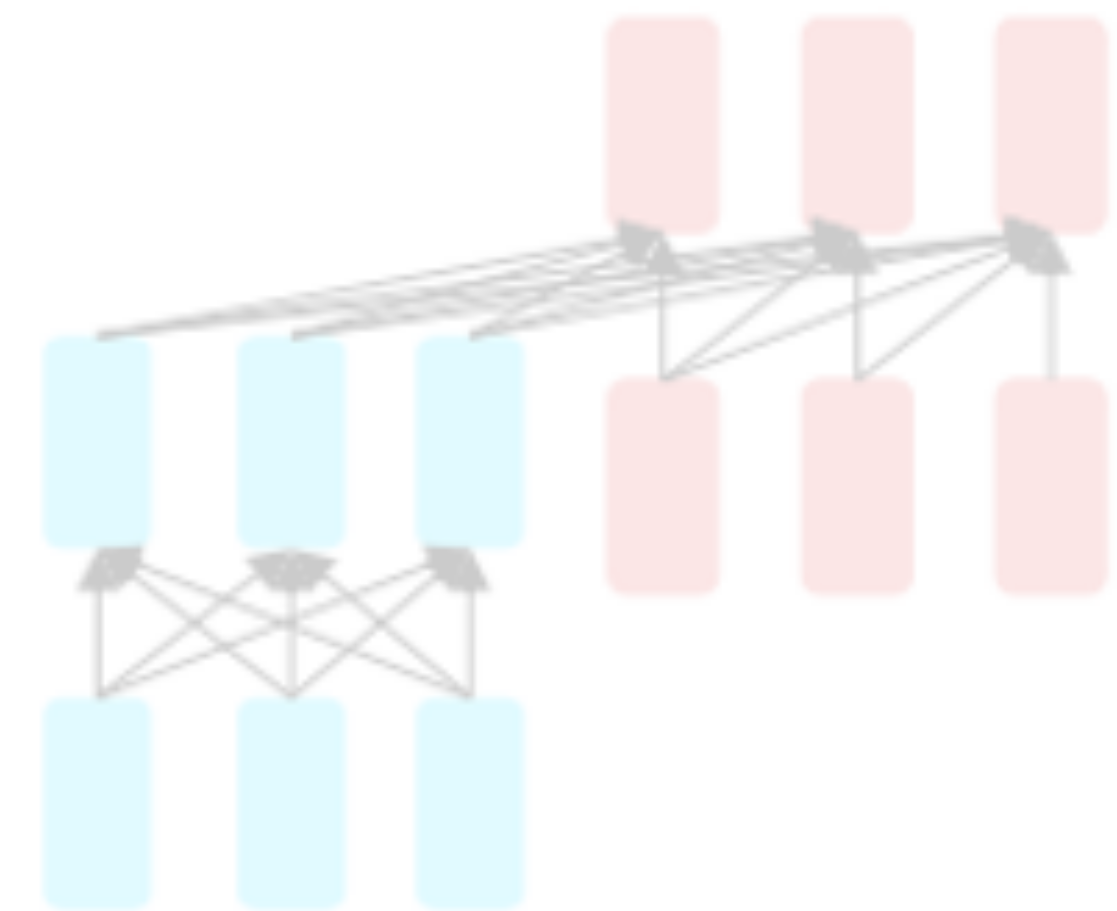
- Masked language models
- Bidirectional context
- BERT + variants (e.g. RoBERTa)
-

Decoder only



- Language models
- Can't condition on future words, good for generation
- GPT, LLaMa, PaLM

Encoder-Decoder

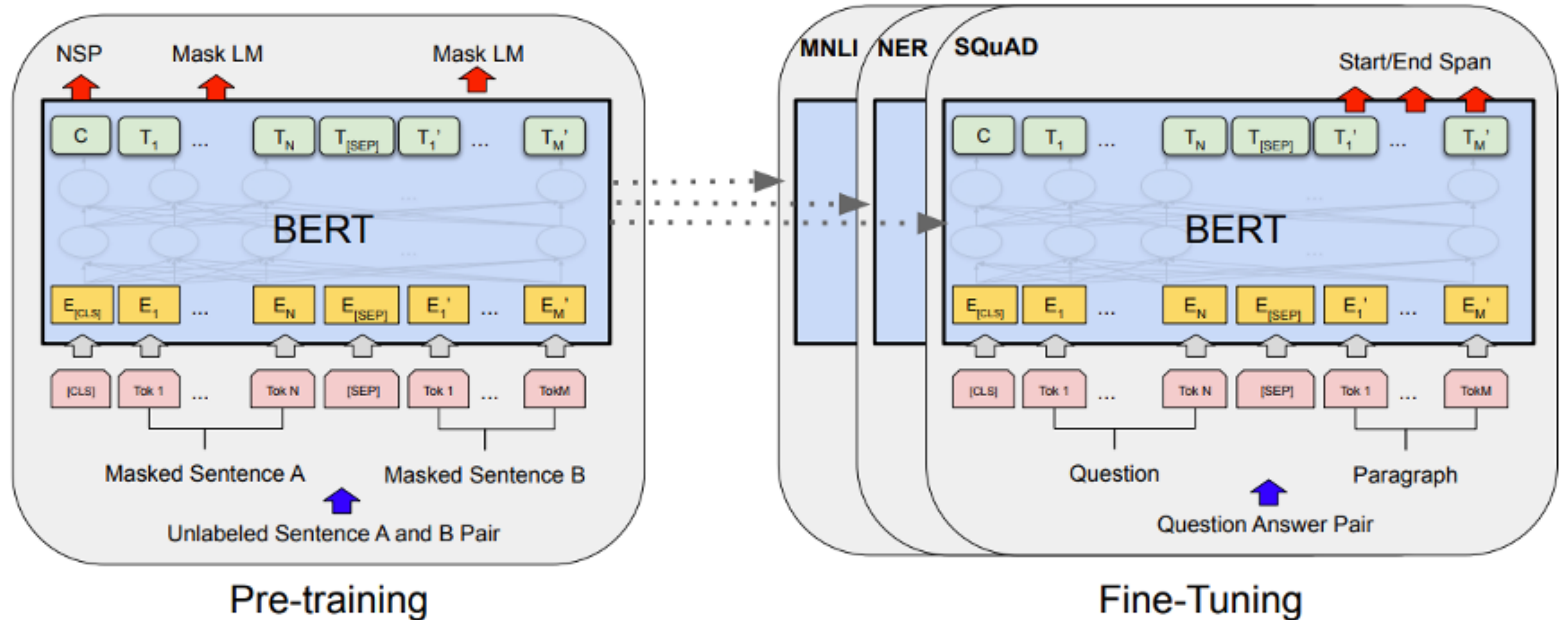


- Combine benefits of both
- Original Transformer, UniLM, BART, T5

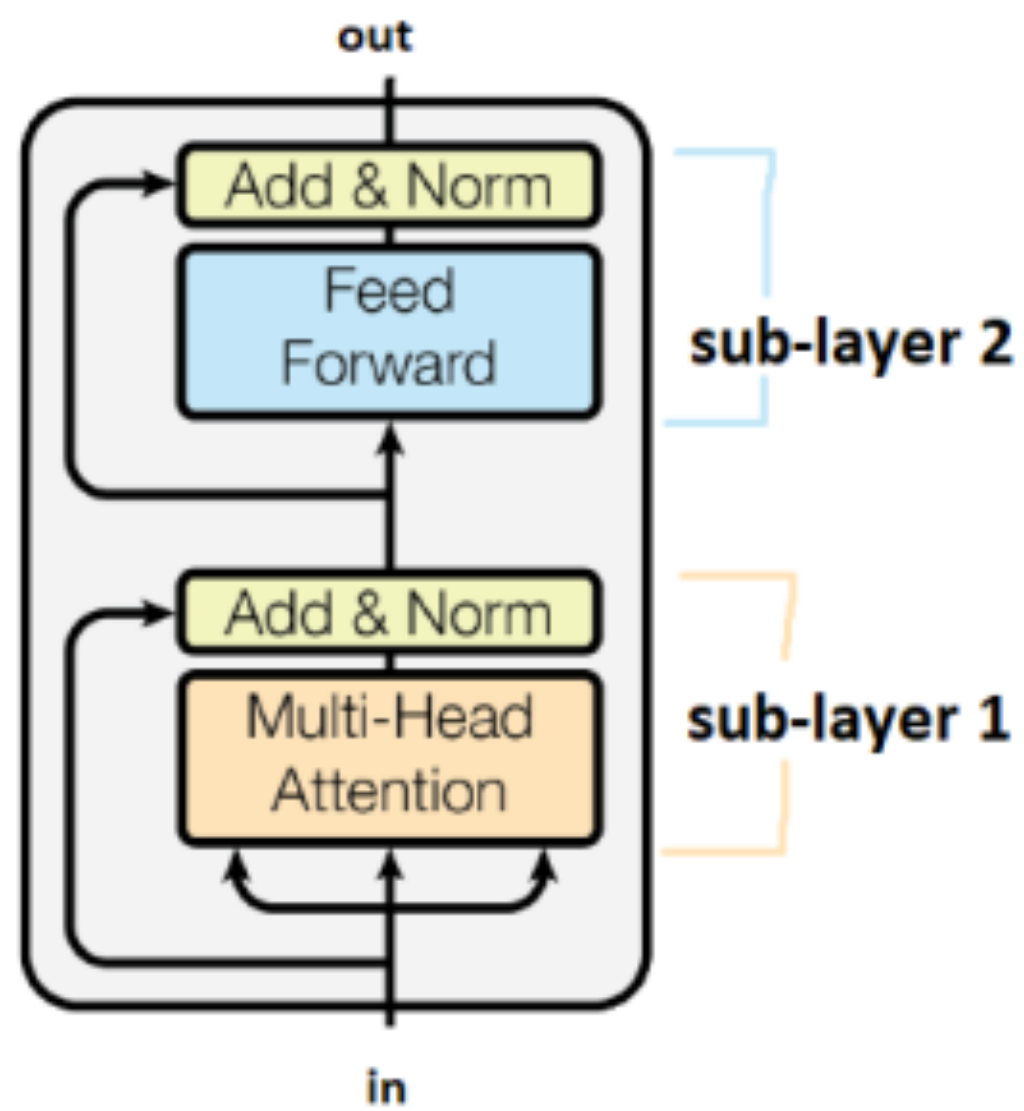
Masked language models: Word2Vec to BERT

Paper	Model	Dataset	Training Objective
W2V CBOW [Miklov et al, 2013]	FFN	Google News (100B words)	Masked LM (within window)
ELMo [Peters et al, 2018]	Bi-LSTM	1B Word benchmark (800M words)	Bidirectional LM
BERT [Devlin et al, 2018]	Transformer (encoder block)	BookCorpus + English Wikipedia (3.3B words)	Masked LM Next sentence prediction

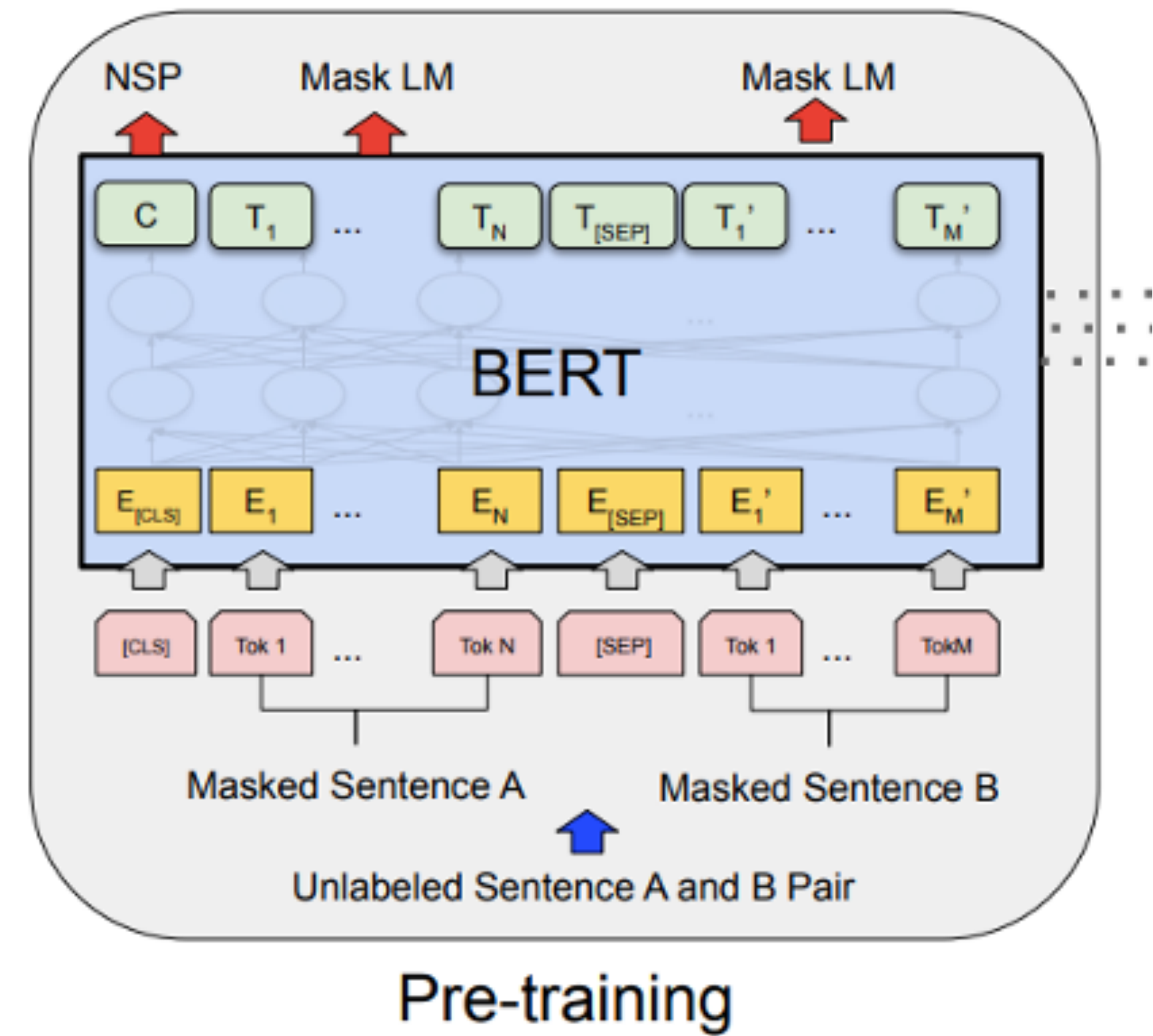
Pre-training and fine-tuning



BERT

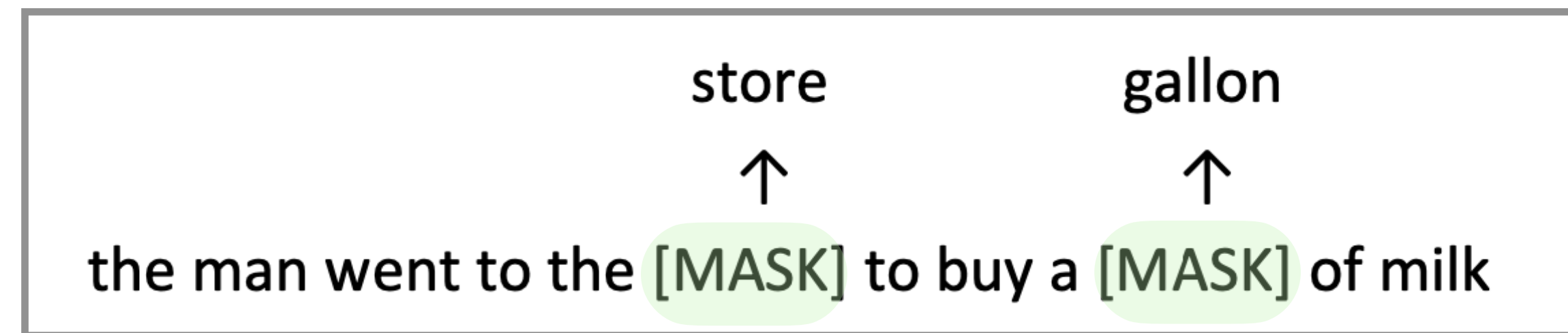


- Transformer Encoder
- Two training objectives
 - Masked Language Modeling
 - Next Sentence Prediction



Masked language models (MLMs)

- Solution: Mask out 15% of the input words, and then predict the masked words



- Too little masking: too expensive to train
- Too much masking: not enough context

Masked language models (MLMs)

A little more complex

(don't always replace with [MASK]):

Example: `my dog is hairy`, we replace the word `hairy`

- 80% of time: replace word with [MASK] token

`my dog is [ MASK ]`

- 10% of time: replace word with random word

`my dog is apple`

- 10% of time: keep word unchanged to bias representation toward actual observed word

`my dog is hairy`

Because [MASK] is never seen when BERT is used...

Next sentence prediction (NSP)

Always sample two sentences, predict whether the second sentence is followed after the first one.

Input = [CLS] the man went to [MASK] store [SEP]
he bought a gallon [MASK] milk [SEP]

Label = IsNext

Input = [CLS] the man [MASK] to the store [SEP]
penguin [MASK] are flight ##less birds [SEP]

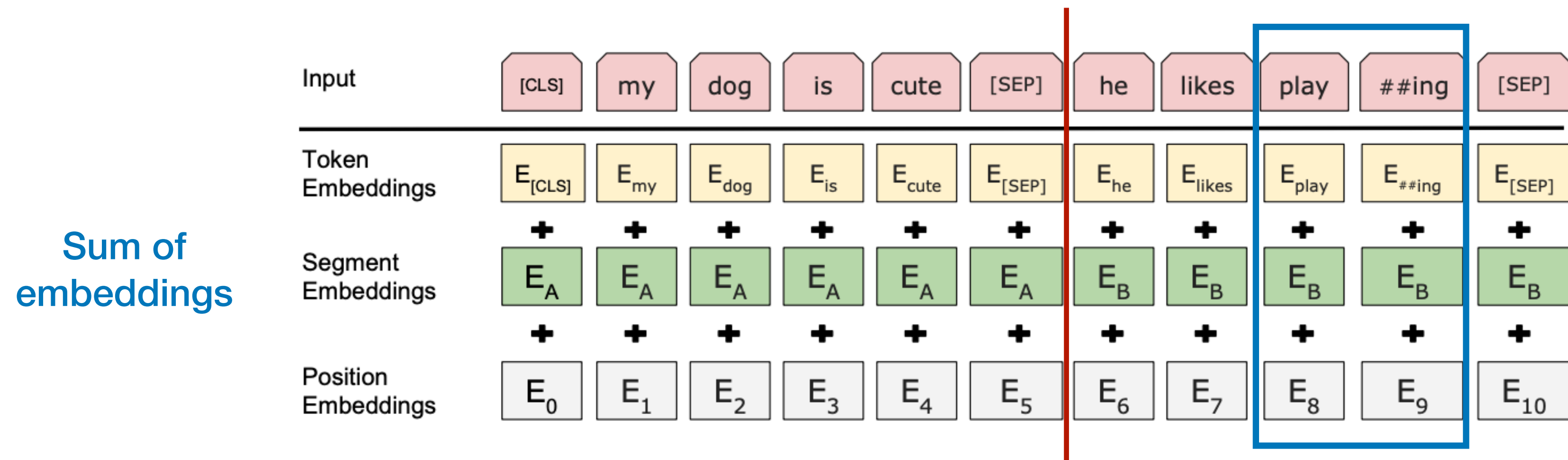
Label = NotNext

Recent papers show that NSP is not necessary...

(Joshi*, Chen* et al, 2019) :SpanBERT: Improving Pre-training by Representing and Predicting Spans
(Liu et al, 2019): RoBERTa: A Robustly Optimized BERT Pretraining Approach

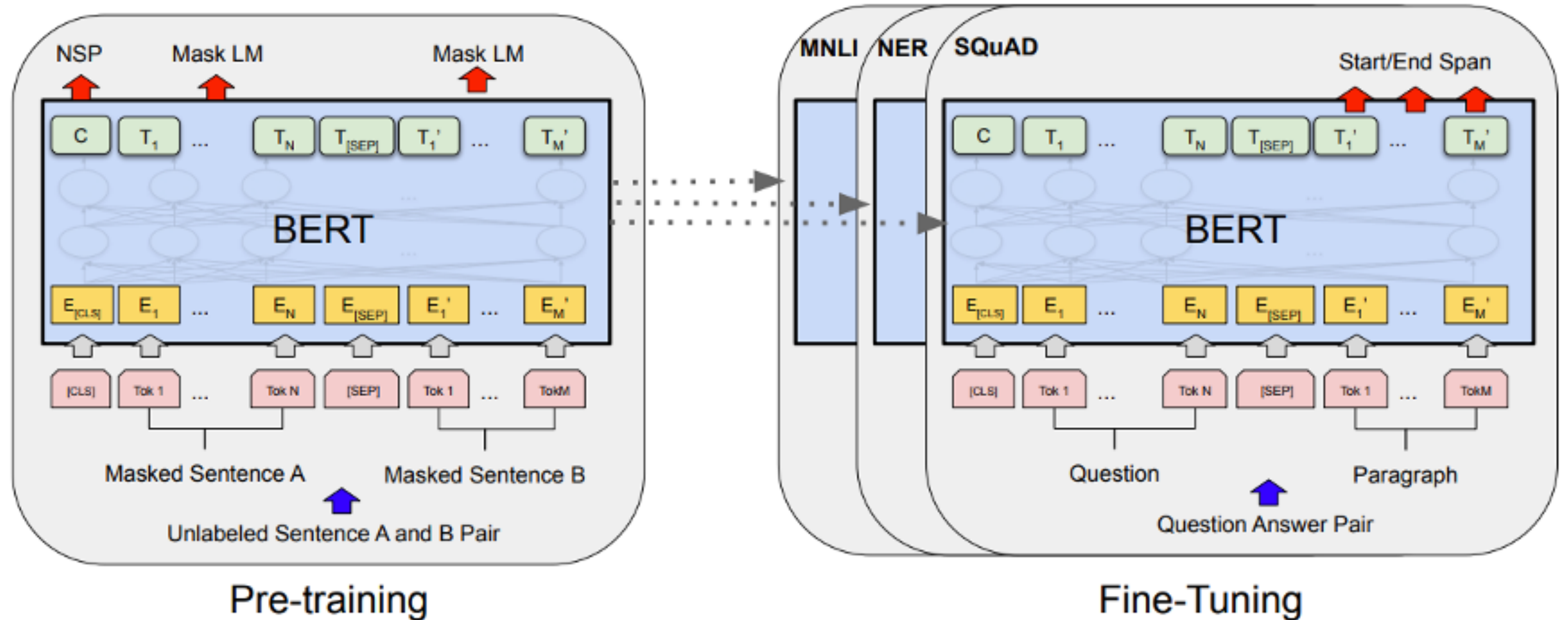
More details

- Input representations



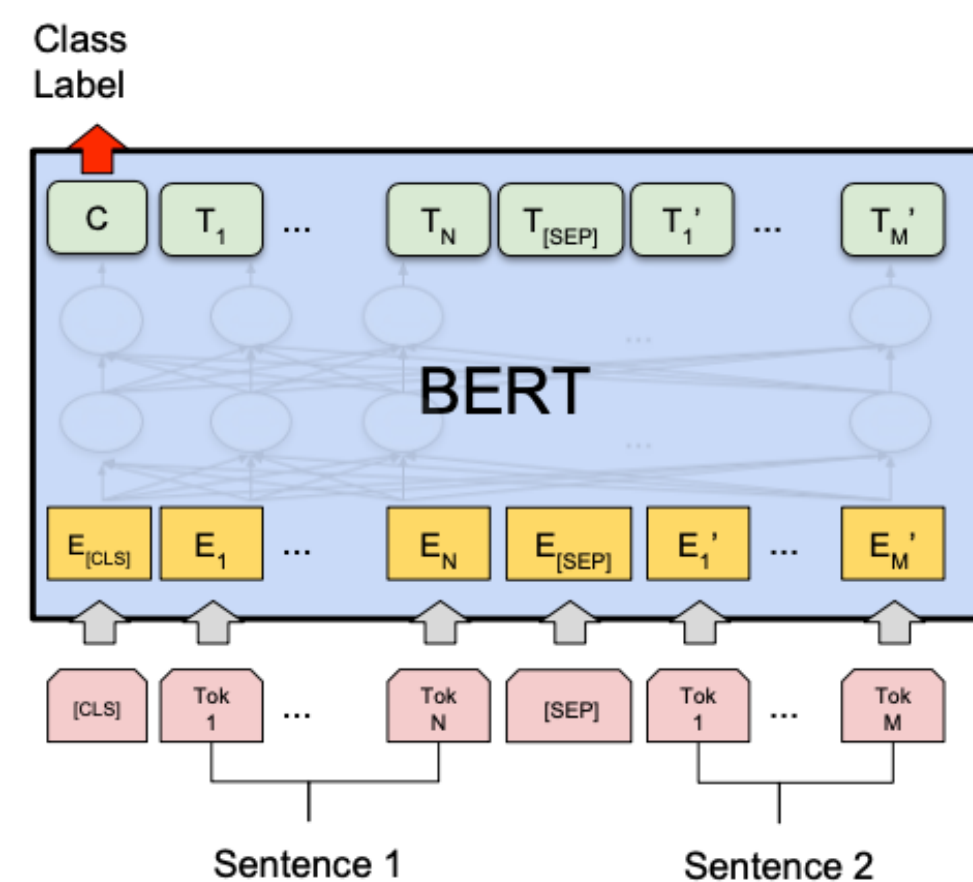
- Use word pieces instead of words: **playing** => **play ##ing** (30K token vocabulary)
- Segment length: 512 tokens
- Trained 40 epochs on Wikipedia (2.5B tokens) + BookCorpus (0.8B tokens)
- Released two model sizes: BERT_base, BERT_large

Pre-training and fine-tuning

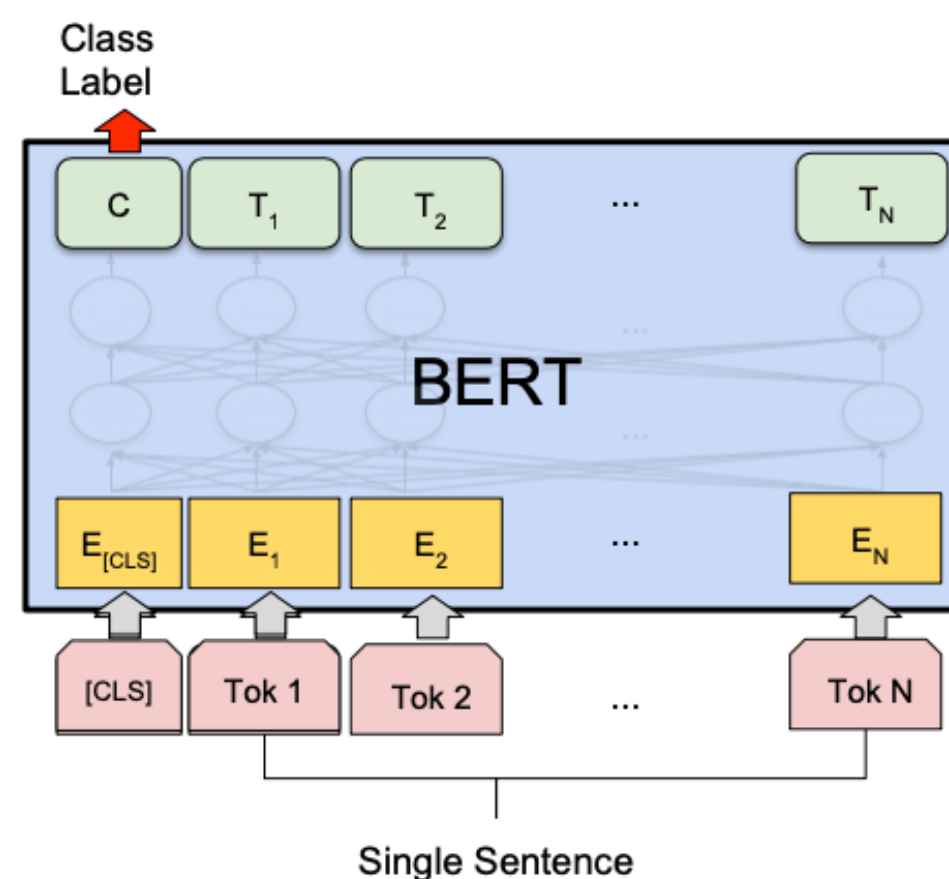


Key idea: **all** the weights are fine-tuned on downstream tasks

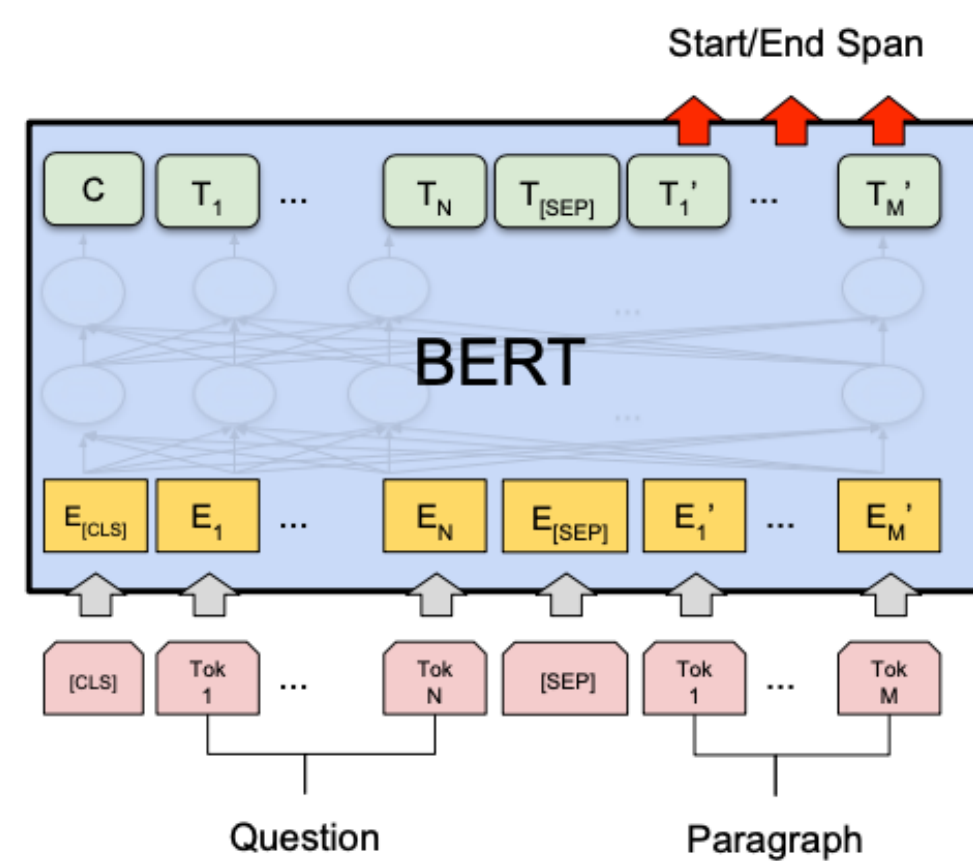
Applications



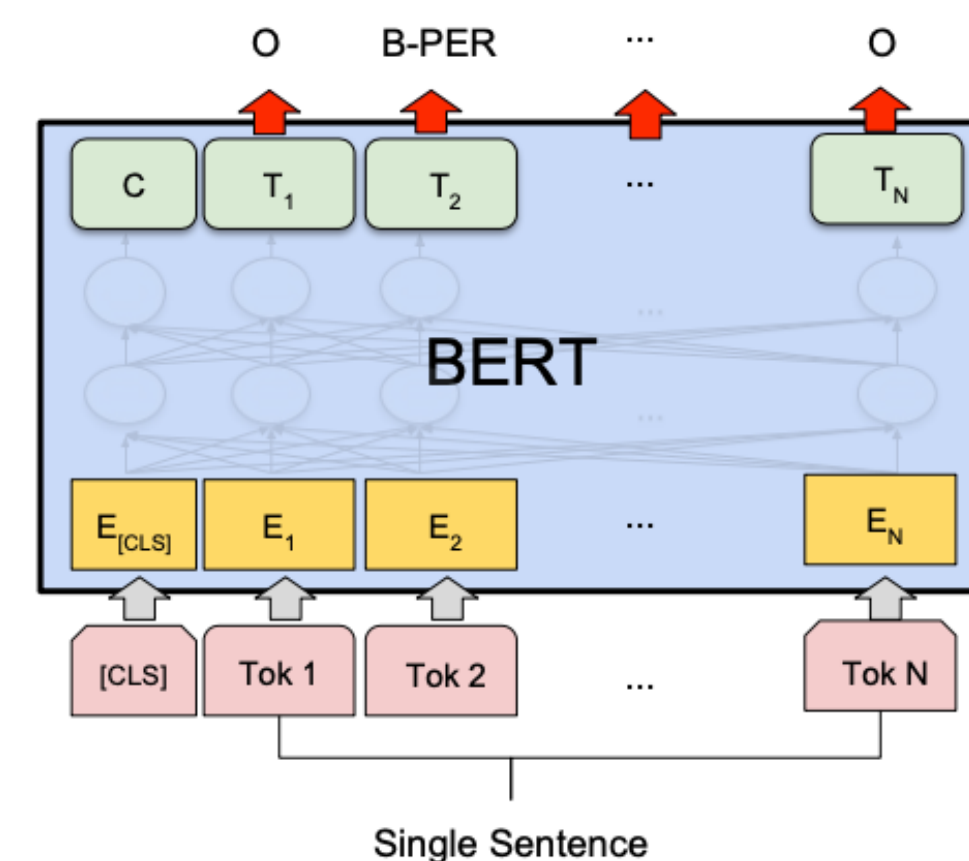
(a) Sentence Pair Classification Tasks:
MNLI, QQP, QNLI, STS-B, MRPC,
RTE, SWAG



(b) Single Sentence Classification Tasks:
SST-2, CoLA

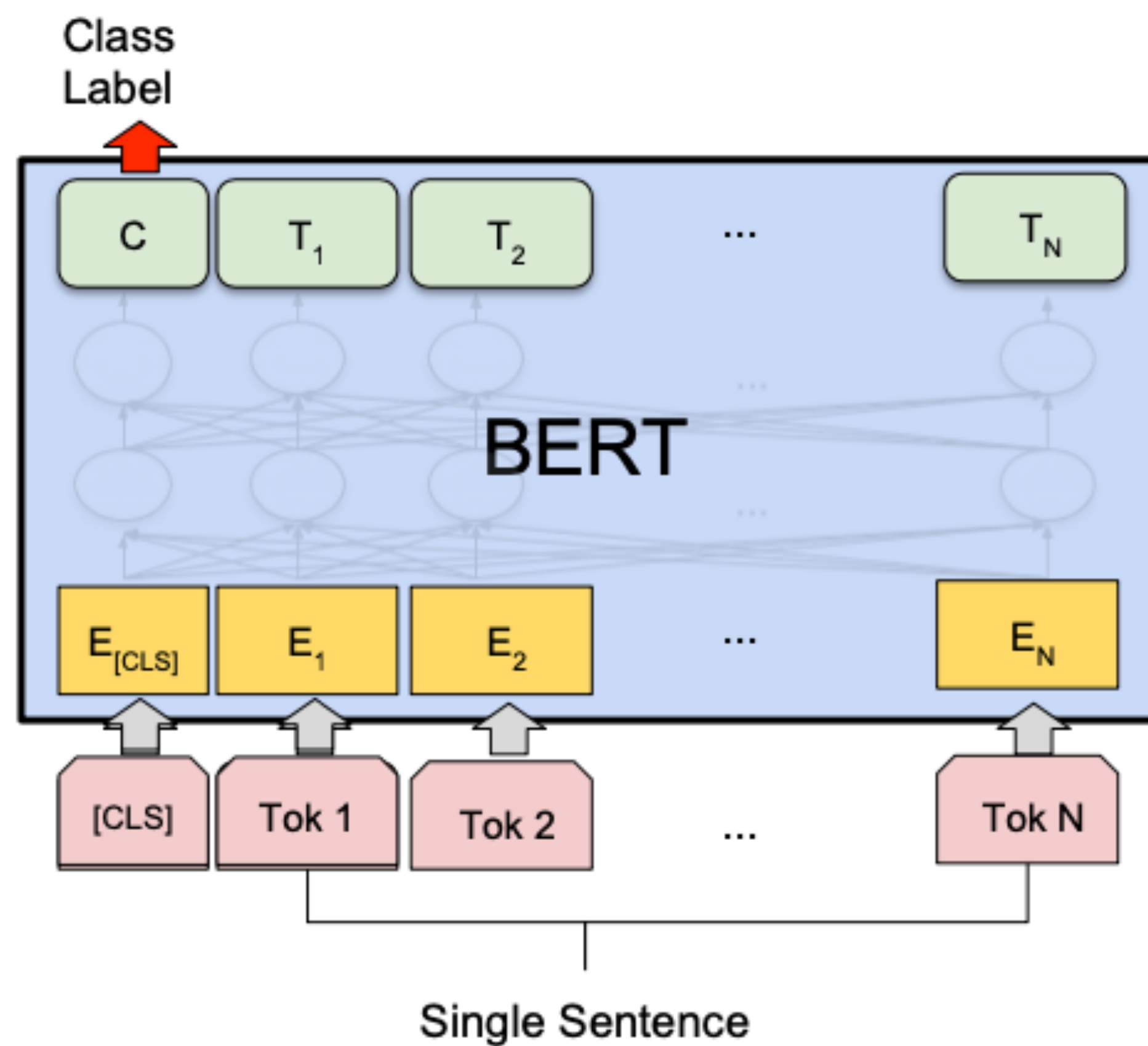


(c) Question Answering Tasks:
SQuAD v1.1



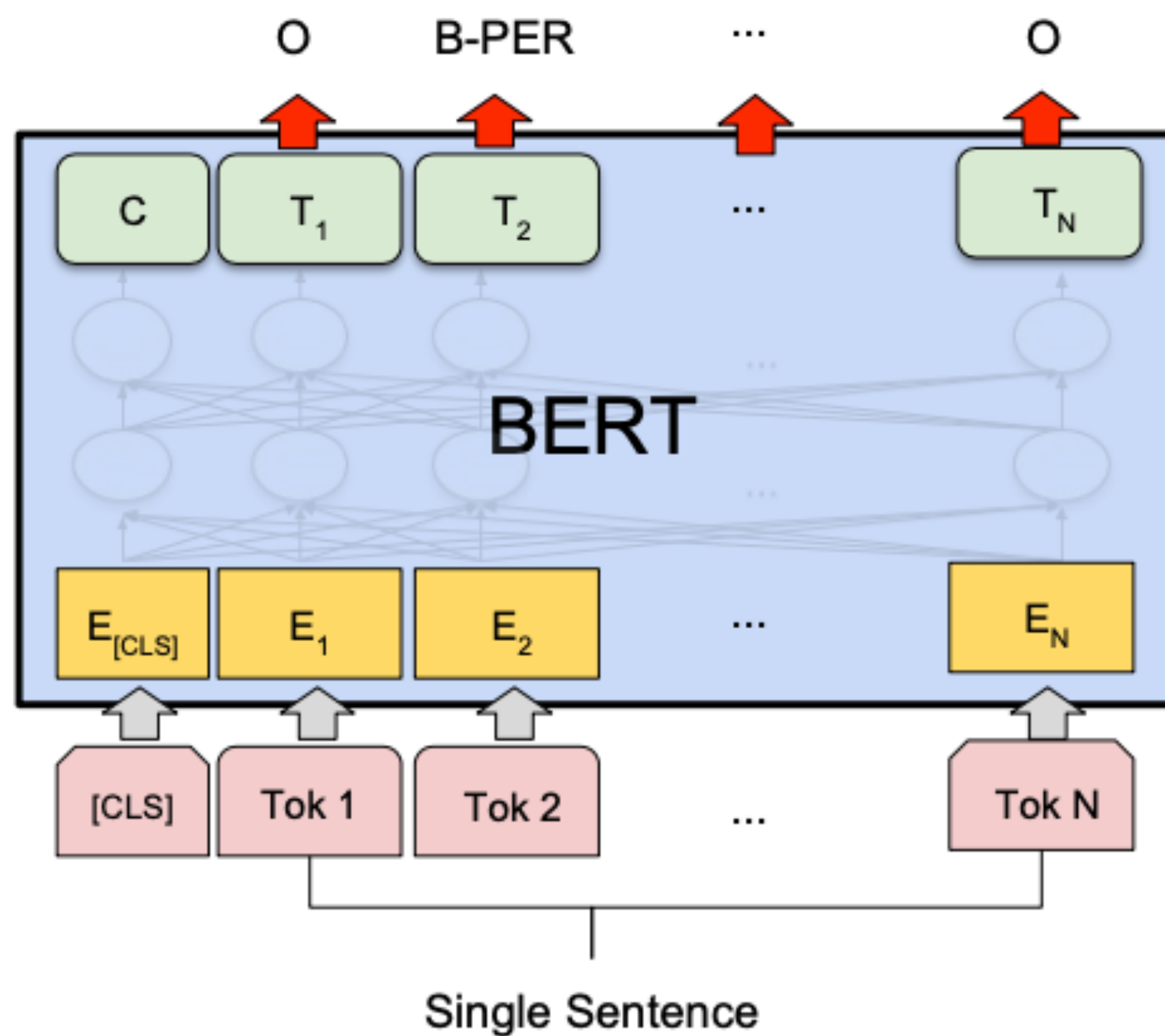
(d) Single Sentence Tagging Tasks:
CoNLL-2003 NER

Applications



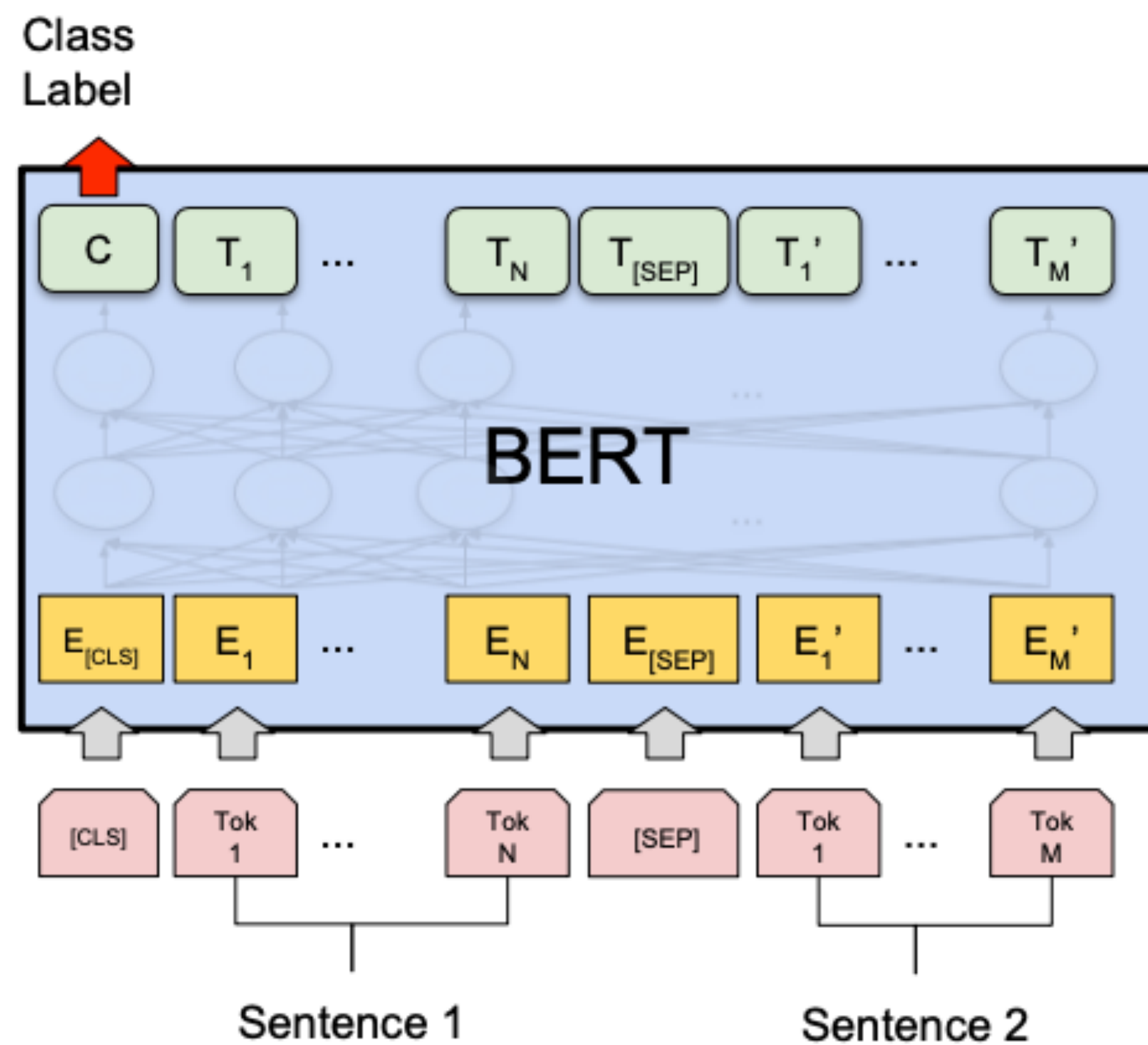
(b) Single Sentence Classification Tasks:
SST-2, CoLA

Applications



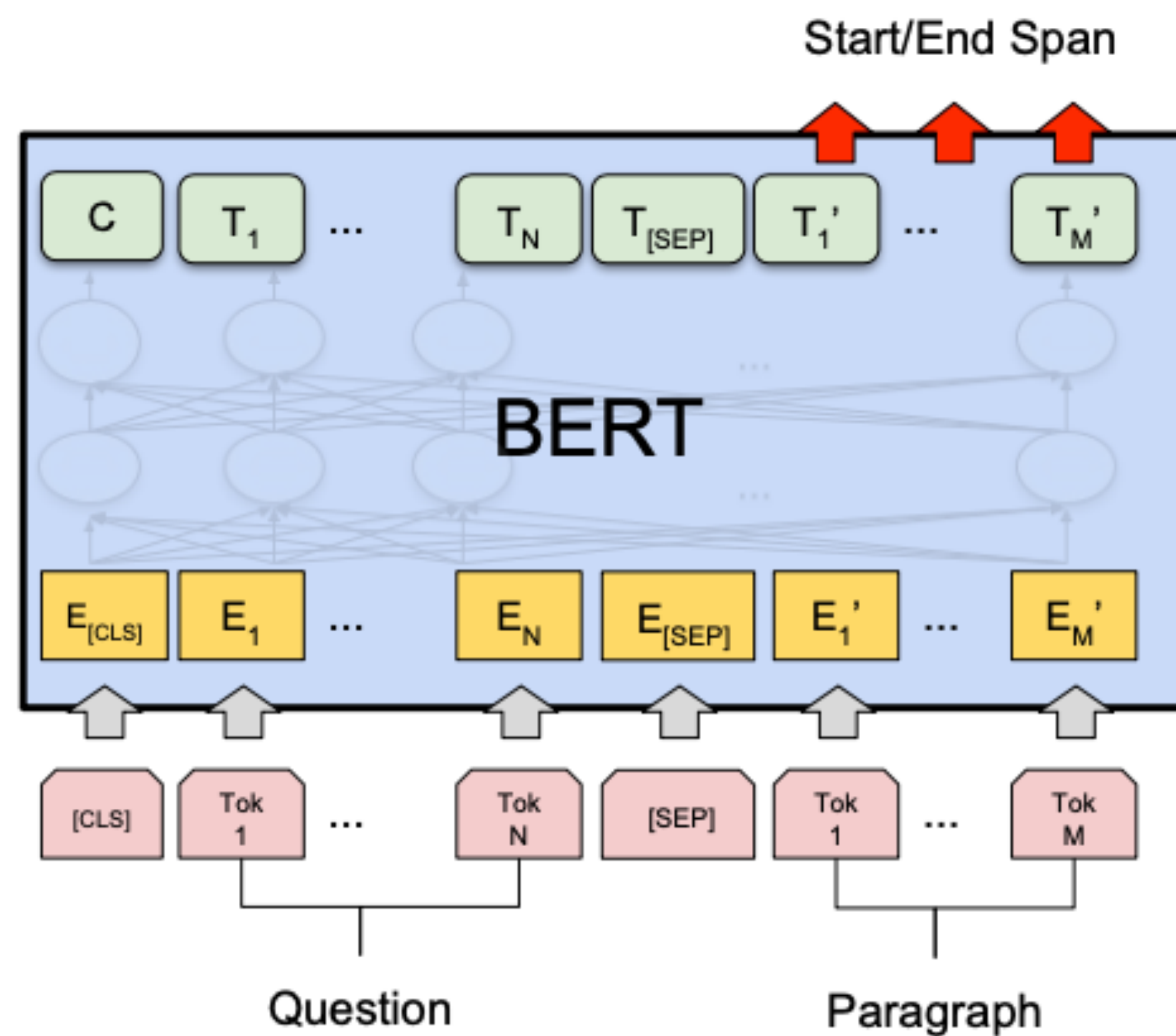
(d) Single Sentence Tagging Tasks:
CoNLL-2003 NER

Applications



(a) Sentence Pair Classification Tasks:
MNLI, QQP, QNLI, STS-B, MRPC,
RTE, SWAG

Applications



(c) Question Answering Tasks:
SQuAD v1.1

BERT Details

Two models were released:

- BERT-base: 12 layers, 768-dim hidden states, 12 attention heads, 110 million params.
- BERT-large: 24 layers, 1024-dim hidden states, 16 attention heads, 340 million params.

Trained on:

- BooksCorpus (800 million words)
- English Wikipedia (2,500 million words)

Pretraining is expensive and impractical on a single GPU.

- BERT was pretrained with 64 TPU chips for a total of 4 days.
- (TPUs are special tensor operation acceleration hardware)

Finetuning is practical and common on a single GPU

- “Pretrain once, finetune many times.”

Experimental results

BiLSTM: 63.9

Entailment

- MNLI: multilingual NLI
- QNLI: NLI with SQuAD data
- RTE: Textual Entailment

Similarity

- QQP: Quora Question Pairs
- STS-B: Semantic Textual Similarity
- MRPC: MS Research Paraphrase Corpus

Other

- SST-2: sentiment analysis
- CoLA: Linguistic acceptability
- SQuAD: question answering

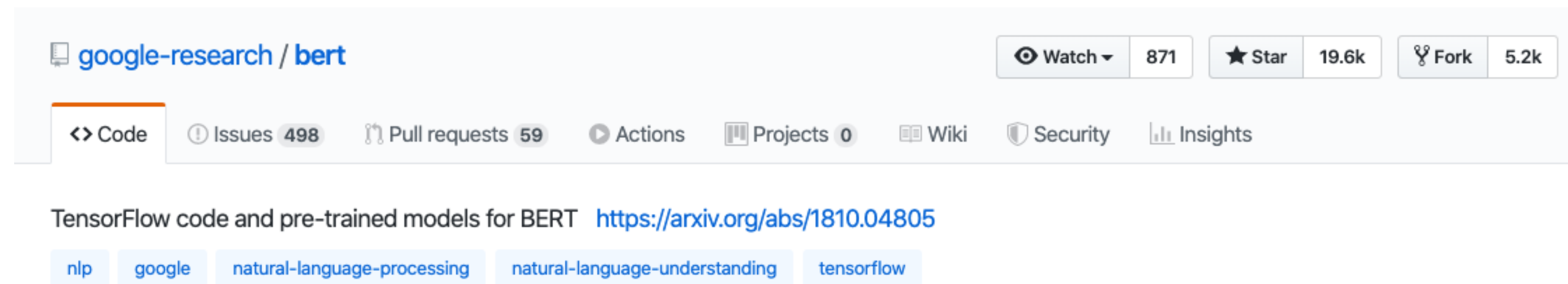
System	MNLI-(m/mm) 392k	QQP 363k	QNLI 108k	SST-2 67k	CoLA 8.5k	STS-B 5.7k	MRPC 3.5k	RTE 2.5k	Average
Pre-OpenAI SOTA	80.6/80.1	66.1	82.3	93.2	35.0	81.0	86.0	61.7	74.0
BiLSTM+ELMo+Attn	76.4/76.1	64.8	79.9	90.4	36.0	73.3	84.9	56.8	71.0
OpenAI GPT	82.1/81.4	70.3	88.1	91.3	45.4	80.0	82.3	56.0	75.2
BERT _{BASE}	84.6/83.4	71.2	90.1	93.5	52.1	85.8	88.9	66.4	79.6
BERT _{LARGE}	86.7/85.9	72.1	91.1	94.9	60.5	86.5	89.3	70.1	81.9

Model	data	bsz	steps	SQuAD (v1.1/2.0)	MNLI-m	SST-2
RoBERTa						
with BOOKS + WIKI	16GB	8K	100K	93.6/87.3	89.0	95.3
+ additional data (§3.2)	160GB	8K	100K	94.0/87.7	89.3	95.6
+ pretrain longer	160GB	8K	300K	94.4/88.7	90.0	96.1
+ pretrain even longer	160GB	8K	500K	94.6/89.4	90.2	96.4
BERT _{LARGE}						
with BOOKS + WIKI	13GB	256	1M	90.9/81.8	86.6	93.7
XLNet _{LARGE}						
with BOOKS + WIKI	13GB	256	1M	94.0/87.8	88.4	94.4
+ additional data	126GB	2K	500K	94.5/88.8	89.8	95.6

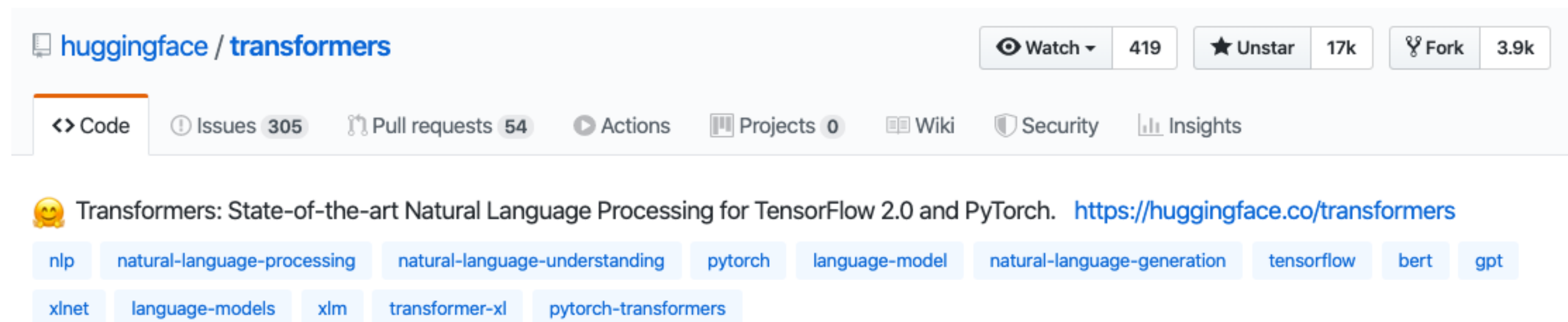
(Wang et al, 2018): GLUE: A Multi-Task Benchmark and Analysis Platform for Natural Language Understanding

Use BERT in practice

TensorFlow: <https://github.com/google-research/bert>



PyTorch: <https://github.com/huggingface/transformers>



Contextualized word embeddings in context

- TagLM (Peters et, 2017)
- CoVe (McCann et al. 2017)
- ULMfit (Howard and Ruder, 2018)
- **ELMo (Peters et al, 2018)**
- OpenAI GPT (Radford et al, 2018)
- **BERT (Devlin et al, 2018)**
- OpenAI GPT-2 (Radford et al, 2019)
- XLNet (Yang et al, 2019)
- SpanBERT (Joshi et al, 2019)
- RoBERTa (Liu et al, 2019)
- ALBERT (Lan et al, 2019)
- DistilBERT (Sanh et al, 2019)
- ELECTRA (Clark et al, 2020)
- ...



<https://github.com/huggingface/transformers>

See <https://huggingface.co/transformers/> for more information and models

Semi-supervised Sequence Learning

context2Vec

Pre-trained seq2seq



ELMo

ULMFiT

Multi-lingual

MultiFiT

Transformer

Bidirectional LM

GPT

Larger model
More data

GPT-2

Defense



Grover

Cross-lingual

Multi-task

+ Generation

XLM

Udify

MT-DNN

Knowledge distillation

MT-DNN_{KD}

MASS

UniLM

Span prediction
Remove NSP

Longer time
Remove NSP
More data

SpanBERT

RoBERTa

Permutation LM
Transformer-XL
More data

XLNet

Noise

+Knowledge Graph



ERNIE
(Tsinghua)

Neural entity linker

KnowBert

Reduced size

Cross-modal

BART

VideoBERT

CBT

ViLBERT

VisualBERT

B2T2

Unicoder-VL

LXMERT

VL-BERT

UNITER

Whole Word Masking

ALBERT

DistilBERT



ERNIE (Baidu)
BERT-wwm

RoBERTa

- Train with more data and for more epochs
 - Vocabulary size of 50K subword units vs 30K for BERT
 - Larger batch size and more training data
- No need for NSP

Model	data	bsz	steps	SQuAD (v1.1/2.0)	MNLI-m	SST-2
RoBERTa						
with BOOKS + WIKI	16GB	8K	100K	93.6/87.3	89.0	95.3
+ additional data (§3.2)	160GB	8K	100K	94.0/87.7	89.3	95.6
+ pretrain longer	160GB	8K	300K	94.4/88.7	90.0	96.1
+ pretrain even longer	160GB	8K	500K	94.6/89.4	90.2	96.4
BERT _{LARGE}						
with BOOKS + WIKI	13GB	256	1M	90.9/81.8	86.6	93.7

pretrain with **1024 V100 GPUs** for ~1 day

RoBERTa: A Robustly Optimized BERT Pretraining Approach
Liu et al, UW and Facebook, arXiv 2019

RoBERTa

- Train with more data and for more epochs
 - Vocabulary size of 50K subword units vs 30K for BERT
 - Larger batch size and more training data
- No need for NSP

Dynamic masking (masking changes)

Masking	SQuAD 2.0	MNLI-m	SST-2
reference	76.3	84.3	92.8
<i>Our reimplementation:</i>			
static	78.3	84.3	92.5
dynamic	78.7	84.0	92.9

Better results with careful reimplementation.
Mean over 5 random seeds.

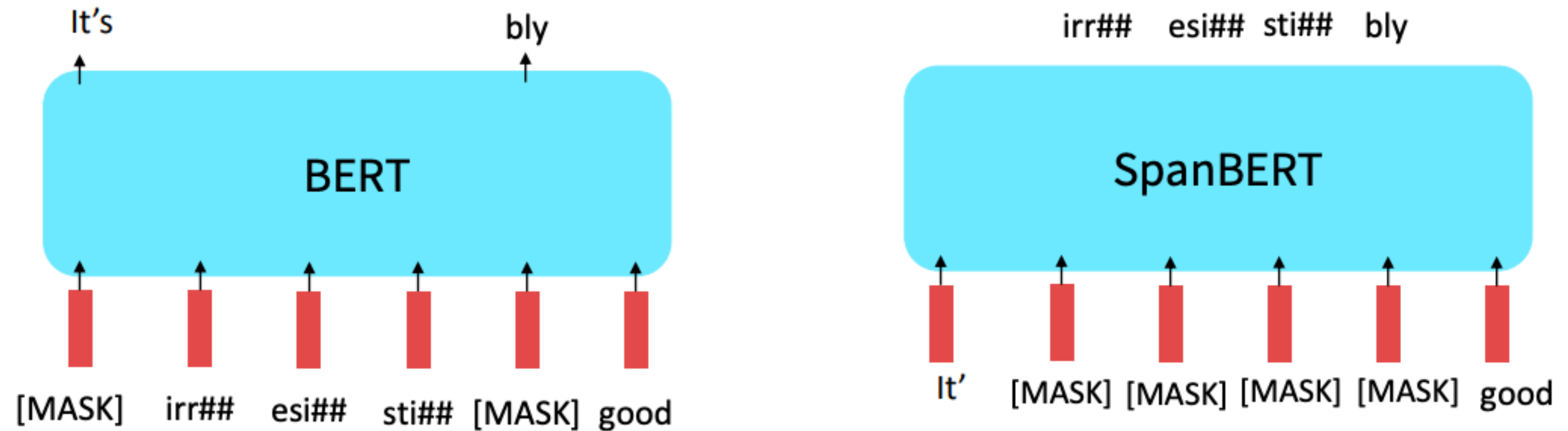
Model	SQuAD 1.1/2.0	MNLI-m	SST-2	RACE
<i>Our reimplementation (with NSP loss):</i>				
SEGMENT-PAIR	90.4/78.7	84.0	92.9	64.2
SENTENCE-PAIR	88.7/76.2	82.9	92.1	63.0
<i>Our reimplementation (without NSP loss):</i>				
FULL-SENTENCES	90.4/79.1	84.7	92.5	64.8
DOC-SENTENCES	90.6/79.7	84.7	92.7	65.6
BERT _{BASE}	88.5/76.3	84.3	92.8	64.3
XLNet _{BASE} (K = 7)	–/81.3	85.8	92.7	66.1
XLNet _{BASE} (K = 6)	–/81.0	85.6	93.4	66.7

RoBERTa: A Robustly Optimized BERT Pretraining Approach

Liu et al, UW and Facebook, arXiv 2019

SpanBERT

- Mask out spans!

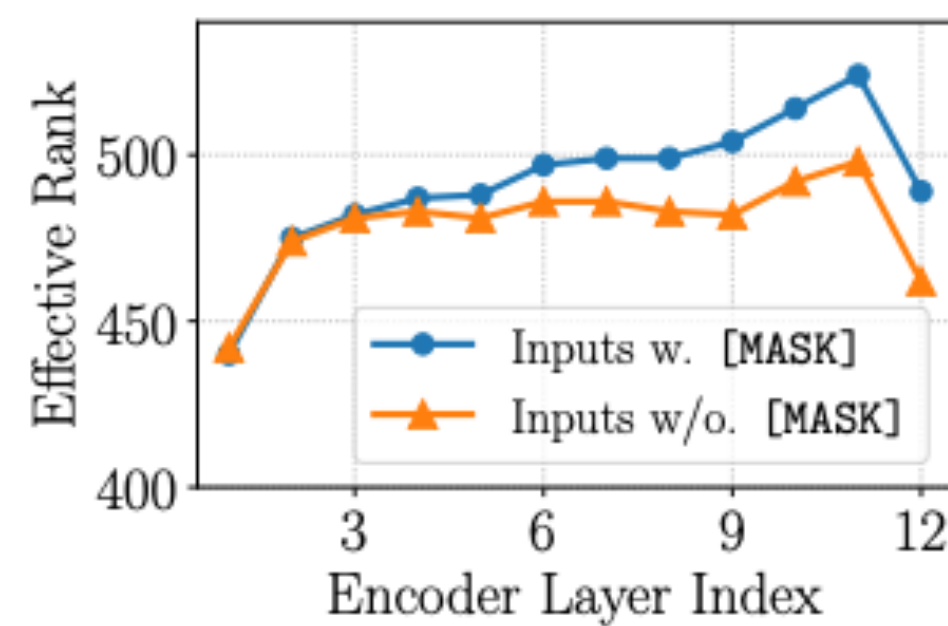


	NewsQA	TriviaQA	SearchQA	HotpotQA	Natural Questions	Avg.
Google BERT	68.8	77.5	81.7	78.3	79.9	77.3
Our BERT	71.0	79.0	81.8	80.5	80.5	78.6
Our BERT-1seq	71.9	80.4	84.0	80.3	81.8	79.7
SpanBERT	73.6	83.6	84.8	83.0	82.5	81.5

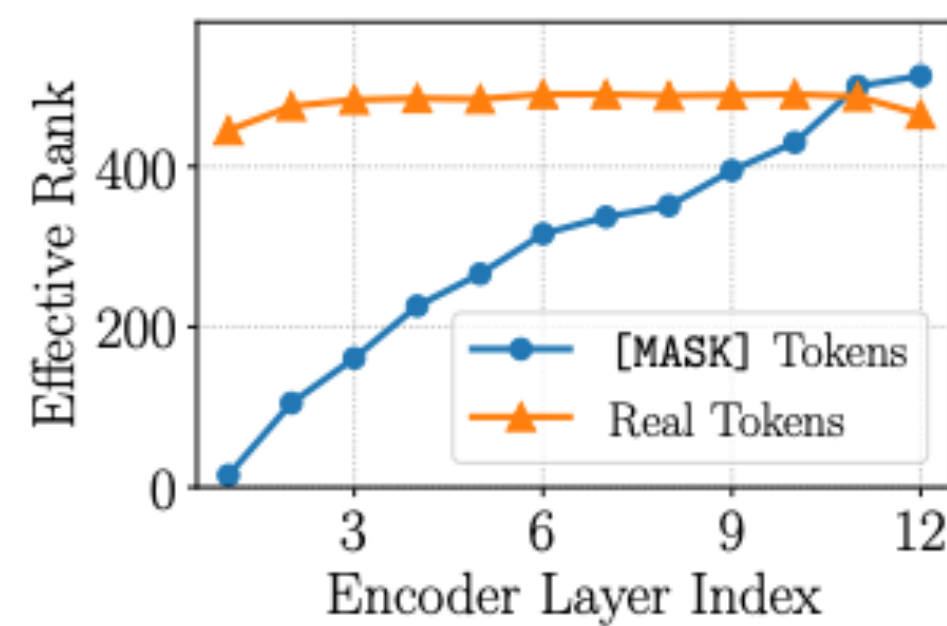
Table 2: Performance (F1) on the five MRQA extractive question answering tasks.

MAE-LM (Masked Autoencoder LM)

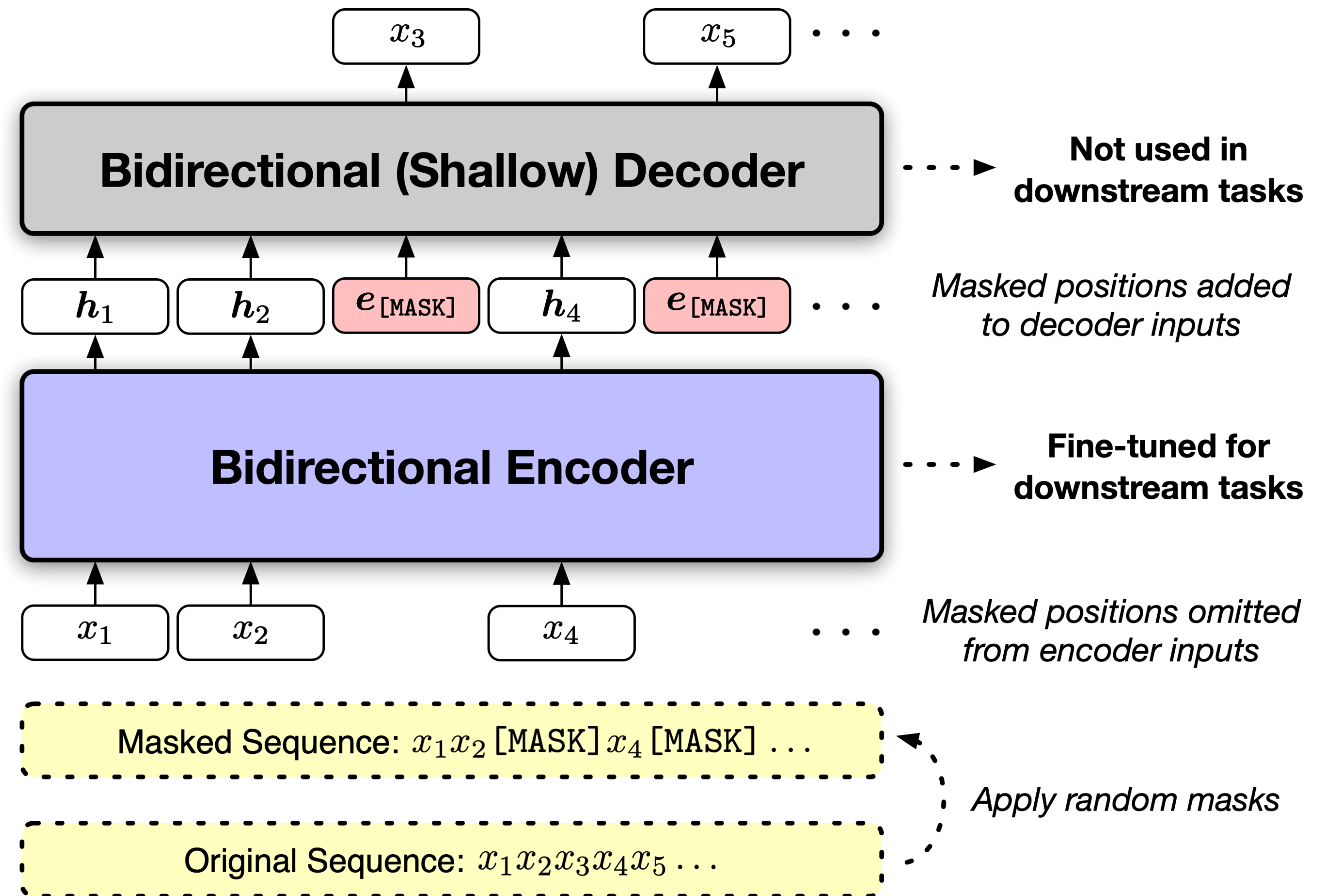
- [MASK] tokens are not observed in downstream tasks
- Model capacity wasted for [MASK] tokens
- Only feed non-masked tokens into encoder, have separate decoder (discarded) that predicts masked tokens



(a)



(b)



Representation Deficiency in Masked Language Modeling [Meng et al. 2024]

MAE-LM (Masked Autoencoder LM)

- [MASK] tokens are not observed in downstream tasks
- Model capacity wasted for [MASK] tokens
- Only feed non-masked tokens into encoder, have separate decoder (discarded) that predicts masked tokens

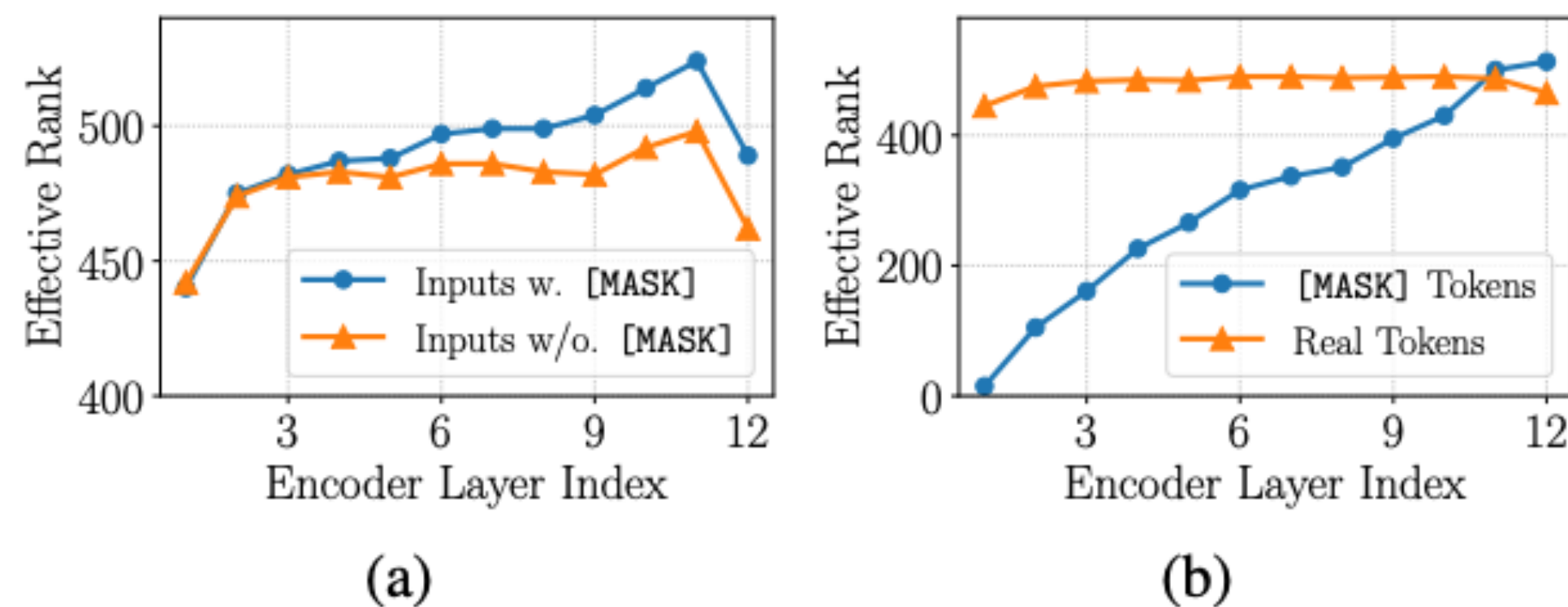


Table 2: Ablations evaluated with GLUE average scores. The setting of MAE-LM_{base} is: enc. w/o. [MASK]; aligned position encoding w. relative position encoding; bi. self-attention; 4 layer, 768 dimension.

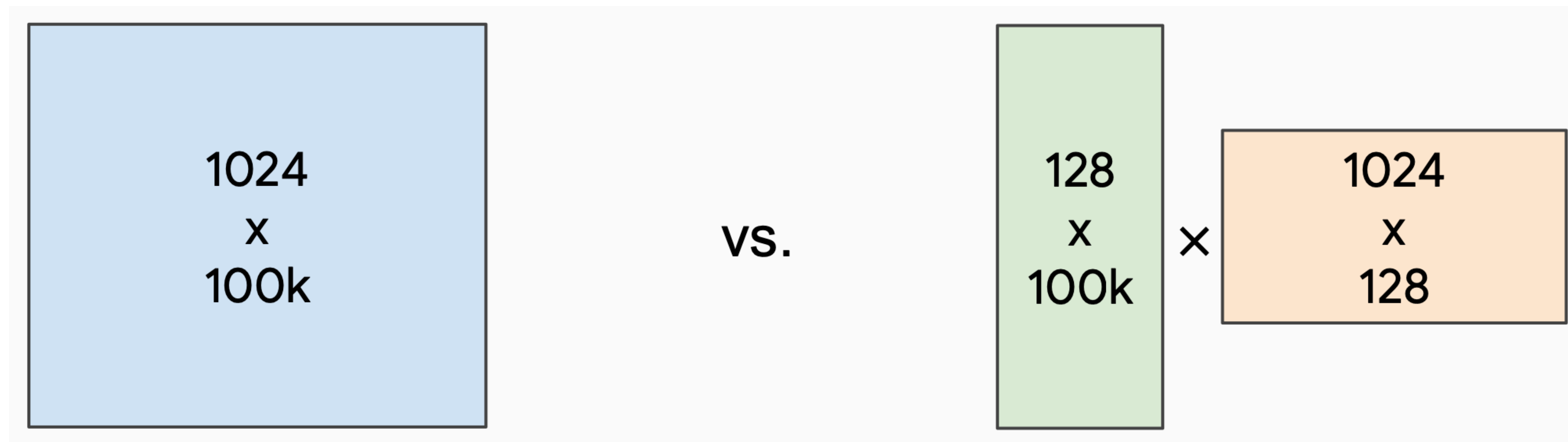
Group	Setting	GLUE
Original	MAE-LM _{base}	86.1
Naive	enc. w. [MASK] (<i>i.e.</i> , MLM)	85.2
	enc. w. [MASK] + dec.	85.1
Handling [MASK]	enc. w. [MASK], dec. resets [MASK]	85.9
	random replace w. real token	85.1
Position Encoding	misaligned position encoding	86.0
	no relative position encoding	86.1
Decoder Attention	bi. self-attention + cross-attention	85.4
	uni. self-attention + cross-attention	85.5
	cross-attention	86.0
Decoder Size	2 layer, 768 dimension	85.8
	6 layer, 768 dimension	84.8
	4 layer, 512 dimension	85.8
	4 layer, 1024 dimension	85.5

ALBERT

<https://arxiv.org/abs/1909.11942>

Lan+ 2019

- Factorized embedding parameterization
 - Use small embedding size (128) and project to Transformer hidden size (1024) using a parameter matrix



ALBERT

<https://arxiv.org/abs/1909.11942>

- Cross-layer parameter sharing
 - $h^{\ell+1}$ parameters are shared with h^{ℓ}

Models	MNLI	QNLI	QQP	RTE	SST	MRPC	CoLA	STS
<i>Single-task single models on dev</i>								
BERT-large	86.6	92.3	91.3	70.4	93.2	88.0	60.6	90.0
XLNet-large	89.8	93.9	91.8	83.8	95.6	89.2	63.6	91.8
RoBERTa-large	90.2	94.7	92.2	86.6	96.4	90.9	68.0	92.4
ALBERT (1M)	90.4	95.2	92.0	88.1	96.8	90.2	68.7	92.7
ALBERT (1.5M)	90.8	95.3	92.2	89.2	96.9	90.9	71.4	93.0

ALBERT

<https://arxiv.org/abs/1909.11942>

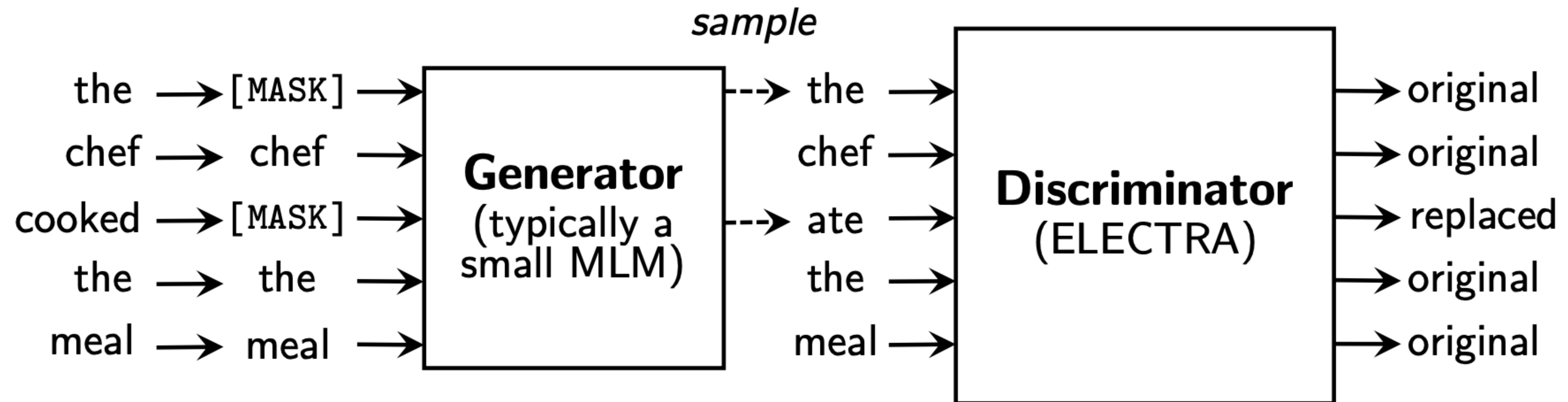
- Light on parameters; not necessarily faster than BERT

Model		Parameters	SQuAD1.1	SQuAD2.0	MNLI	SST-2	RACE	Avg	Speedup
BERT	base	108M	90.4/83.2	80.4/77.6	84.5	92.8	68.2	82.3	4.7x
	large	334M	92.2/85.5	85.0/82.2	86.6	93.0	73.9	85.2	1.0
ALBERT	base	12M	89.3/82.3	80.0/77.1	81.6	90.3	64.0	80.1	5.6x
	large	18M	90.6/83.9	82.3/79.4	83.5	91.7	68.5	82.4	1.7x
	xlarge	60M	92.5/86.1	86.1/83.1	86.4	92.4	74.8	85.5	0.6x
	xxlarge	235M	94.1/88.3	88.1/85.1	88.0	95.2	82.3	88.7	0.3x

Discriminative training

Loss is on all the training tokens vs just the masked ones, more compute efficient use of the training data

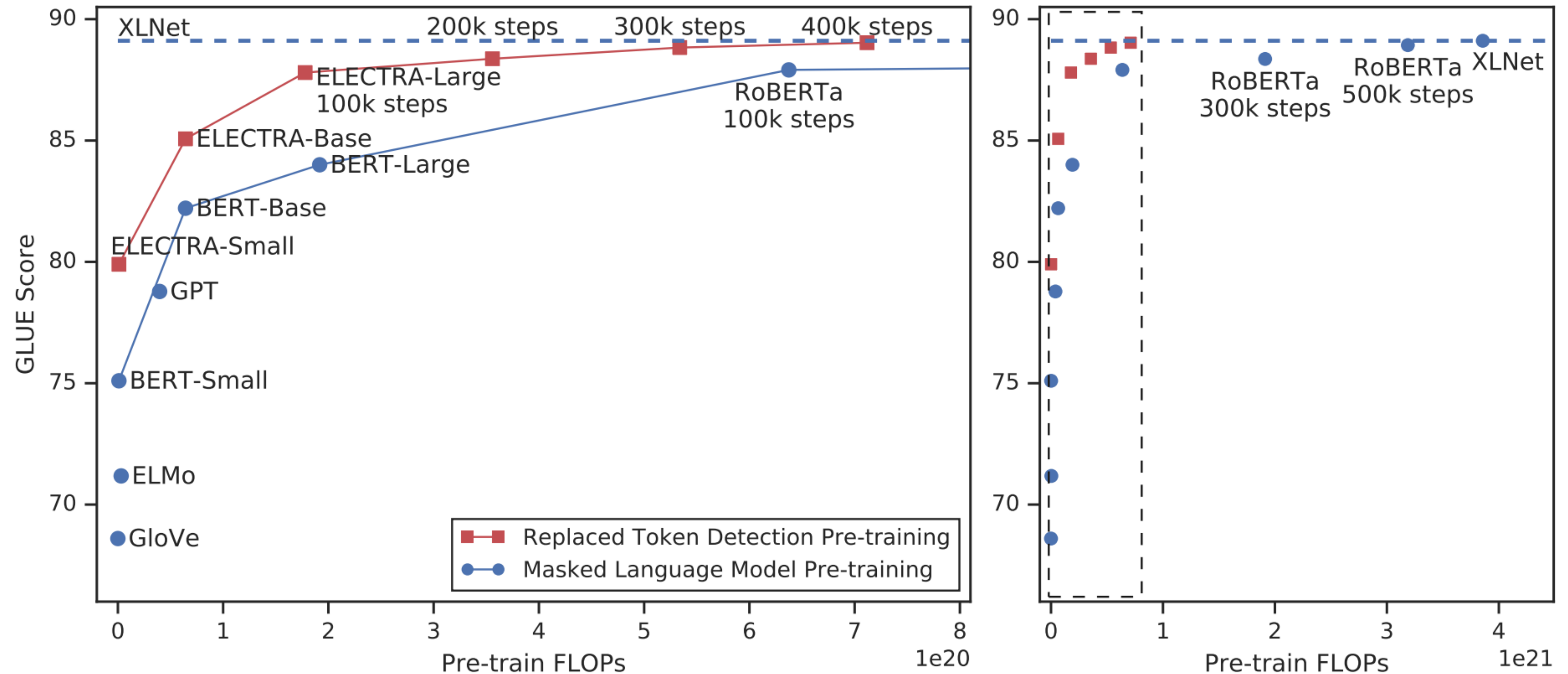
Train model to discriminate locally plausible text from real text



ELECTRA: Pre-training Text Encoders as Discriminators Rather Than Generators

Clark et al, ICLR 2020

Discriminative training



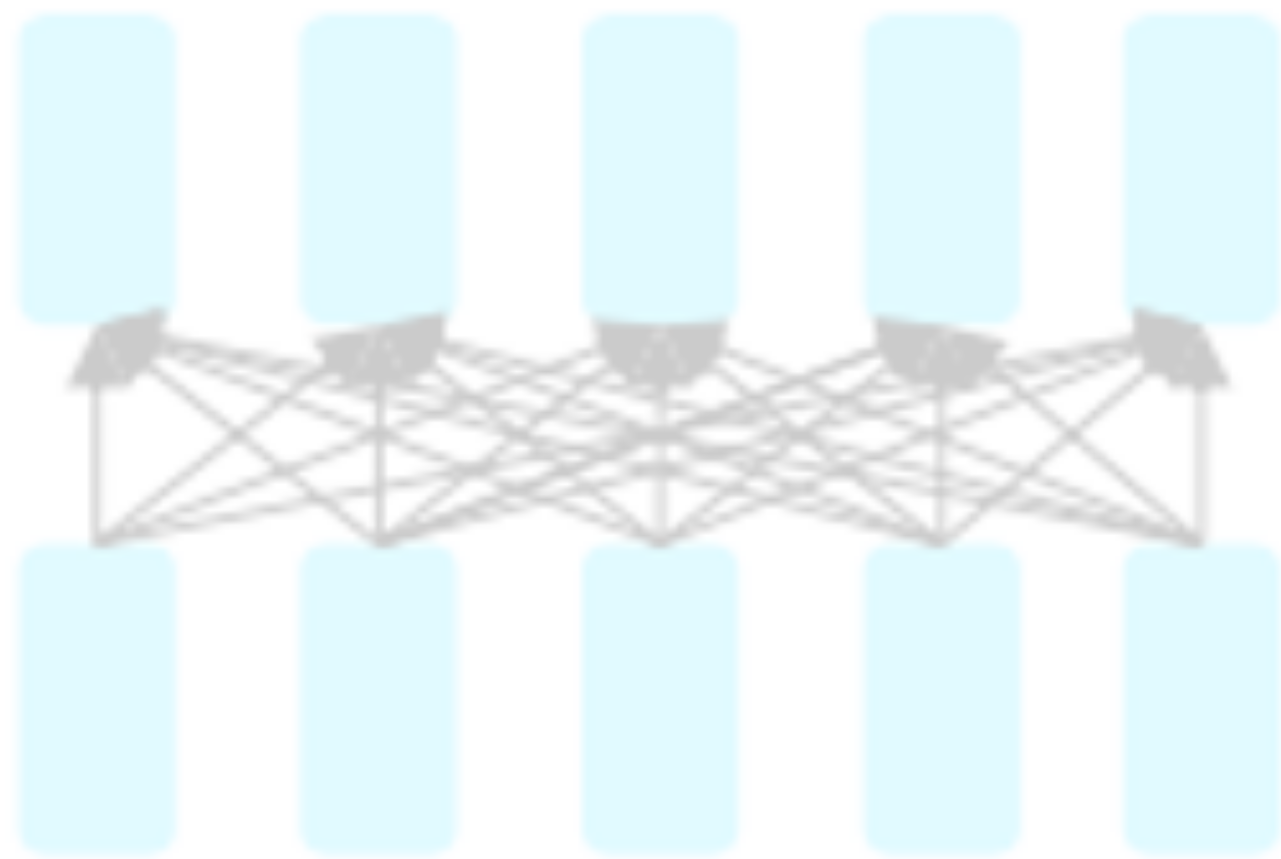
ELECTRA: Pre-training Text Encoders as Discriminators Rather Than Generators

Clark et al, ICLR 2020

Transformers for pretraining

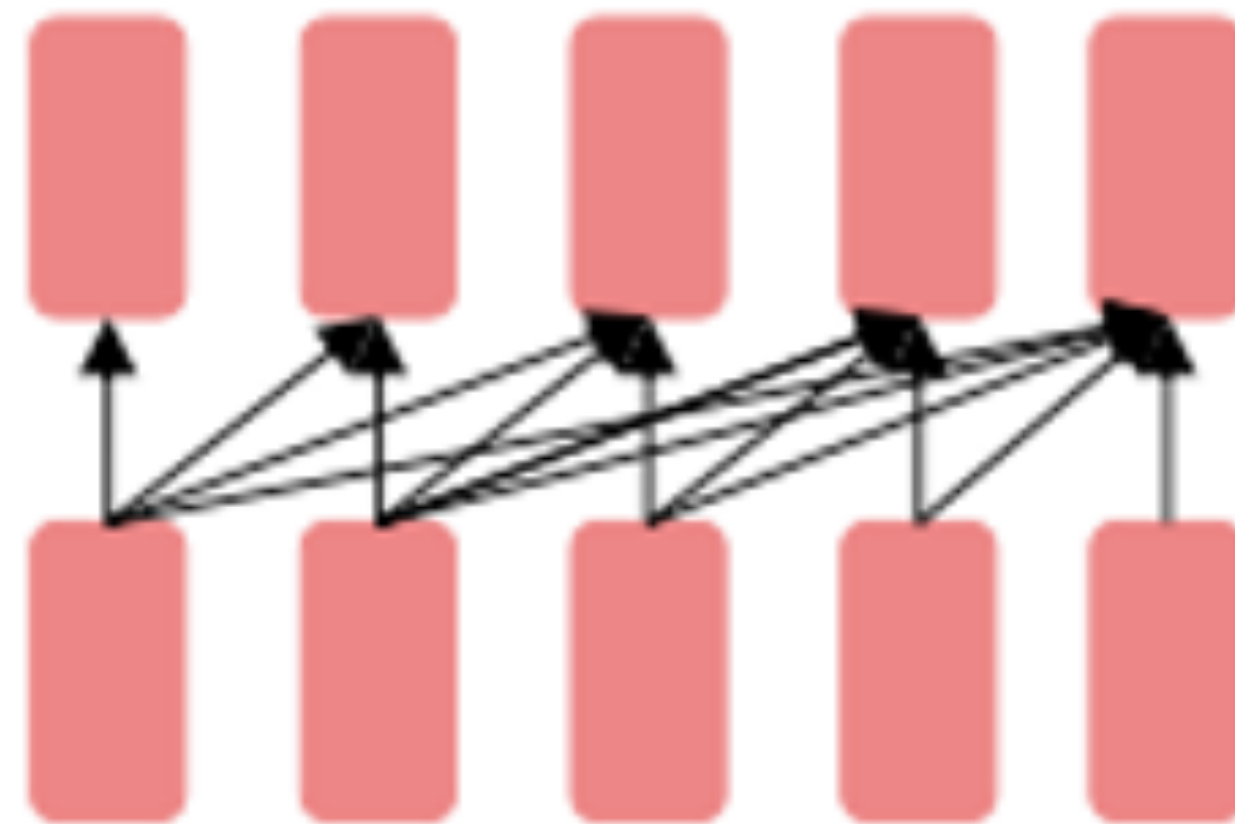
- Self-supervised Transformer based models shattered language understanding benchmarks in NLP in 2018.
- Trained on large text corpus with self-supervised objectives and then transferred.

Encoder only



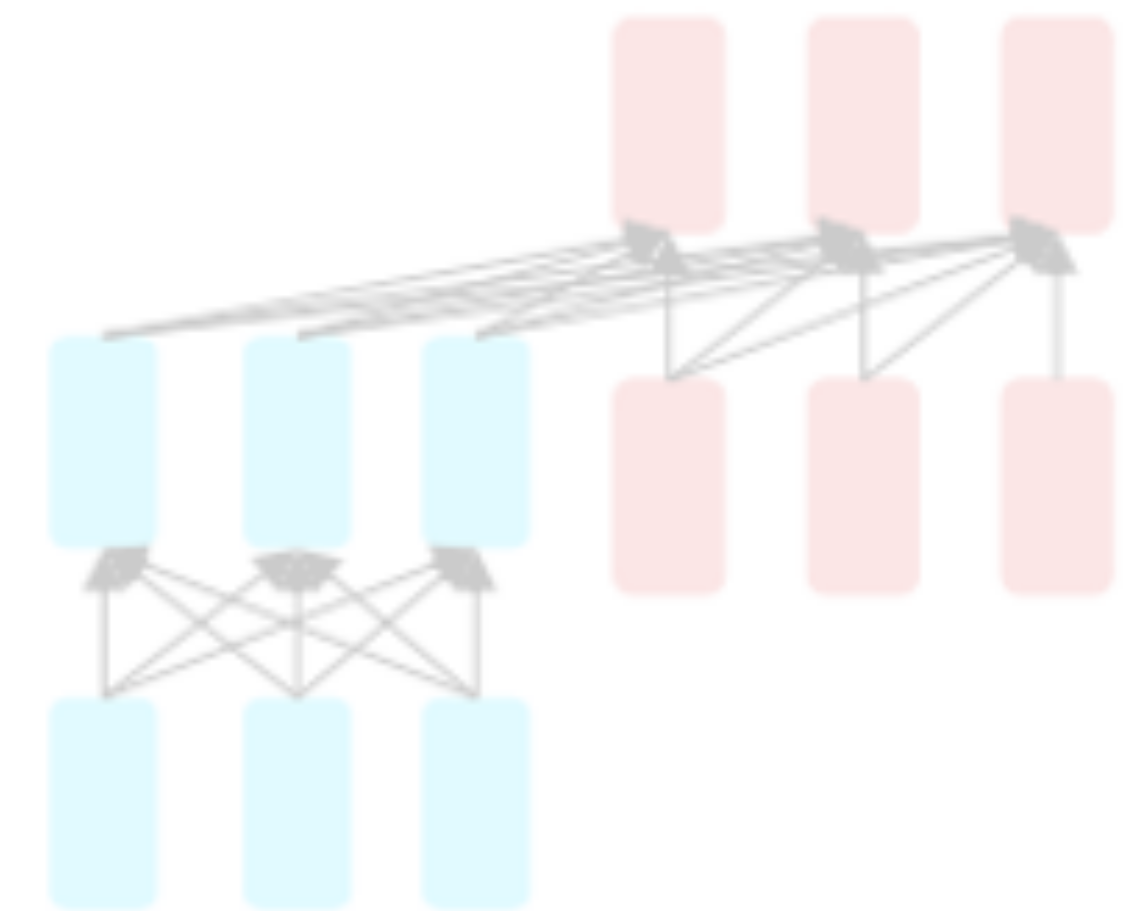
- Masked language models
- Bidirectional context
- BERT + variants (e.g. RoBERTa)
-

Decoder only



- Language models
- Can't condition on future words, good for generation
- GPT, LLaMa, PaLM

Encoder-Decoder



- Combine benefits of both
- Original Transformer, UniLM, BART, T5

Improving Language Understanding by Generative Pre-Training

GPT1

Alec Radford
OpenAI
alec@openai.com

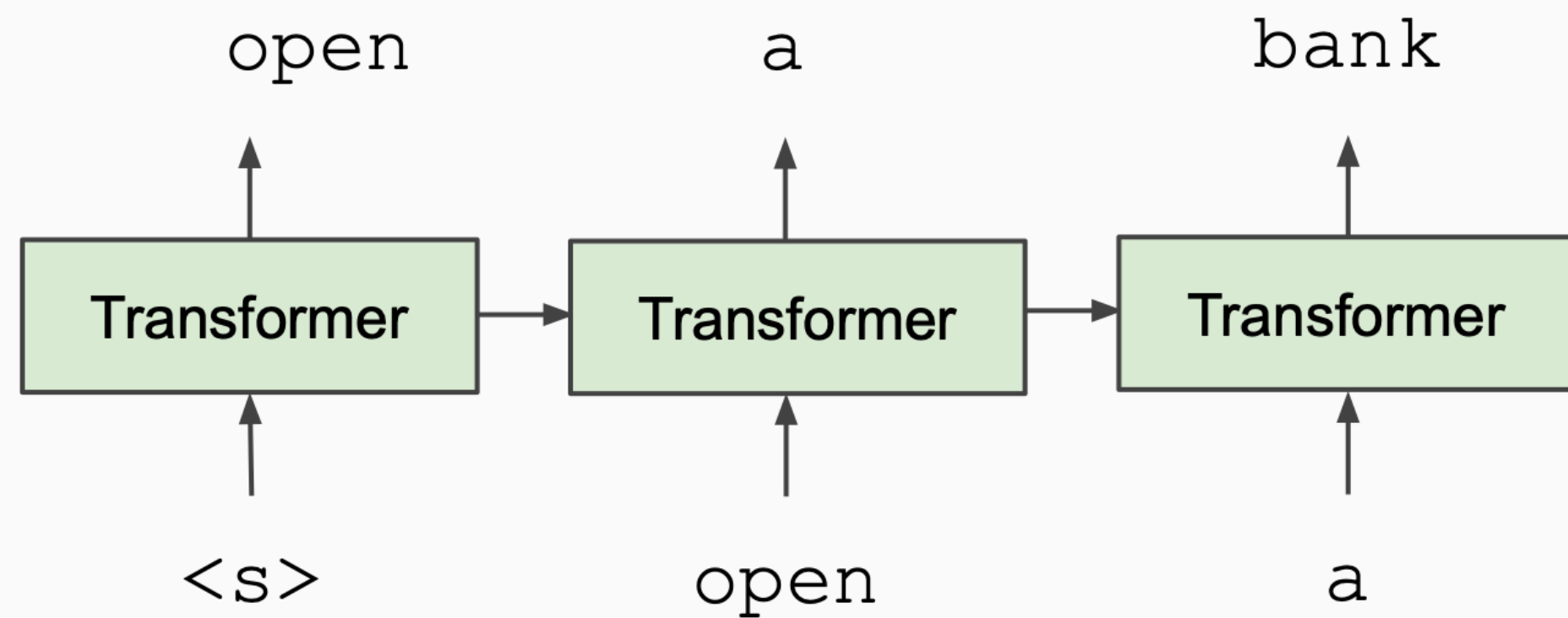
Karthik Narasimhan
OpenAI
karthikn@openai.com

Tim Salimans
OpenAI
tim@openai.com

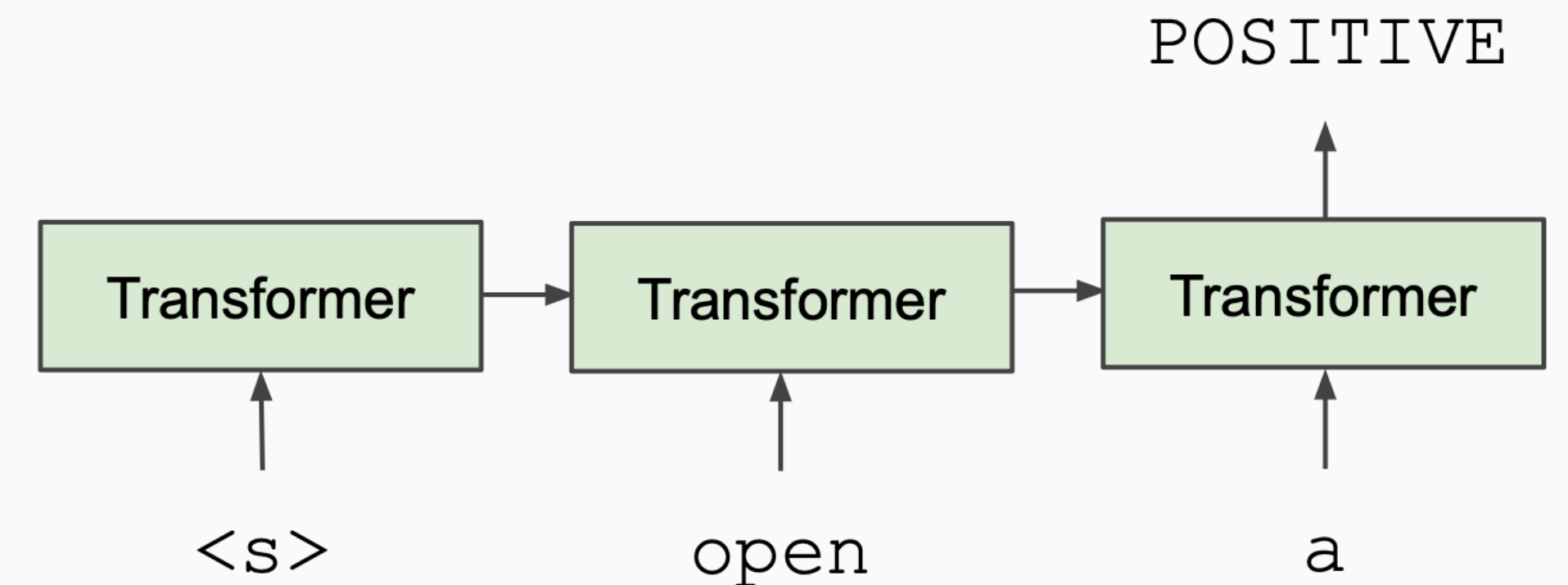
Ilya Sutskever
OpenAI
ilyasu@openai.com

GPT1

Train Deep (12-layer) Transformer LM



Fine-tune on Classification Task



GPT1

Pre-training an autoregressive language model

BooksCorpus: 7K
unpublished books
(1B words)

- Start with a large amount of unlabeled data $\mathcal{U} = \{u_1, \dots, u_n\}$
- Pre-training objective: Maximize the likelihood of predicting the next token

$$L_i(\mathcal{U}) = \sum_i \log P(u_i \mid u_{i-k}, \dots, u_{i-1}; \Theta)$$

$U = (u_{-k}, \dots, u_{-1})$ is the context vector of tokens

- This is equivalent to training a Transformer decoder

n is the number of Transformer layers

$$h_0 = U \boxed{W_e} + W_p$$

W_e is the token embedding matrix

$$h_\ell = \text{transformer_block}(h_{\ell-1}) \forall \ell \in [1, n]$$

W_p is the position embedding matrix

$$P(u) = \text{softmax}(h_n \boxed{W_e^T})$$

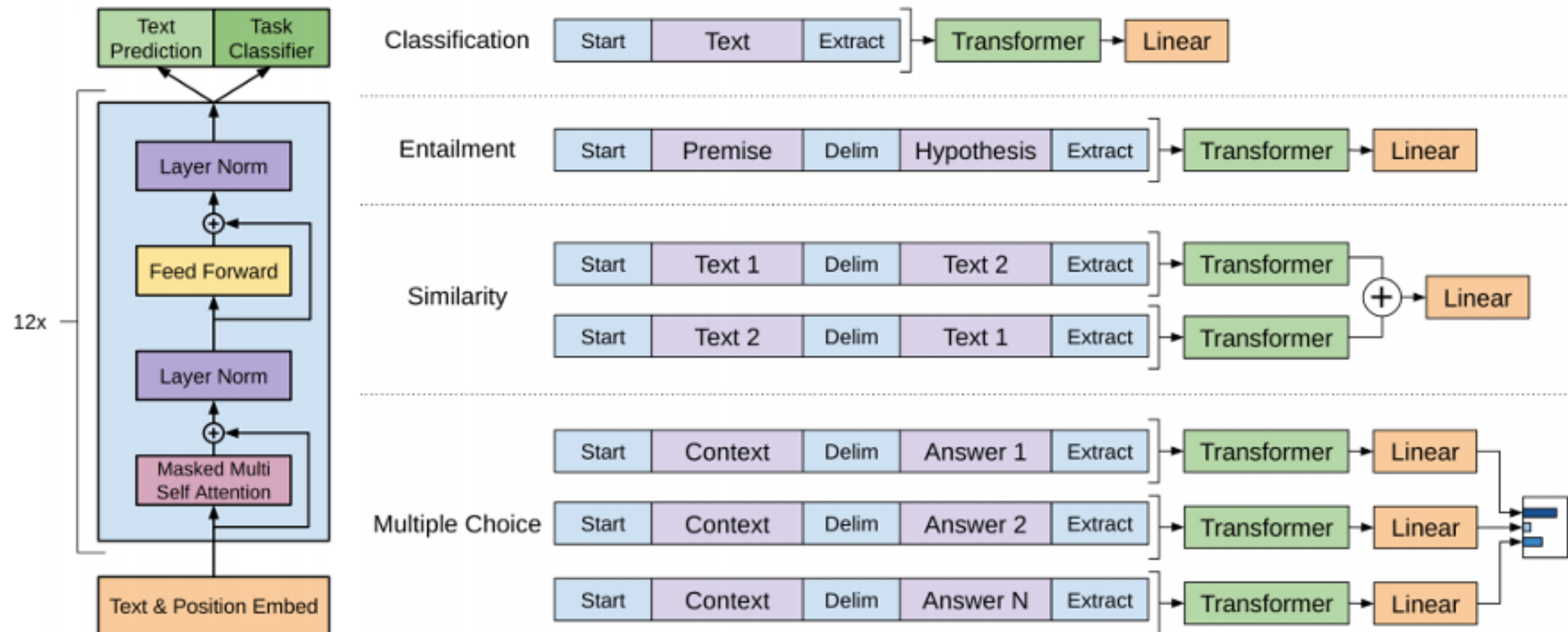
- Directionality is needed to generate a well-formed probability distribution

Dataset	Task	SOTA	GPT1
SNLI	Textual entailment	89.3	89.9
MNLI matched	Textual entailment	80.6	82.1
MNLI mismatched	Textual entailment	80.1	81.4
SciTail	Textual entailment	83.3	88.3
QNLI	Textual entailment	82.3	88.1
RTE	Textual entailment	61.7	56.0
STS-B	Semantic similarity	81.0	82.0
QQP	Semantic similarity	66.1	70.3
MRPC	Semantic similarity	86.0	82.3
RACE	Reading comprehension	53.3	59.0
ROCStories	Commonsense reasoning	77.6	86.5
COPA	Commonsense reasoning	71.2	78.6
SST-2	Sentiment analysis	93.2	91.3
CoLA	Linguistic acceptability	35.0	45.4
GLUE	Multi task benchmark	68.9	72.8

<https://openai.com/research/language-unsupervised>

GPT (Generative pretrained transformer)

- Unsupervised retraining: Standard language model loss
- Supervised fine-tuning: Use simple classifier (linear layer + softmax) trained to predict correct class (use combined loss)

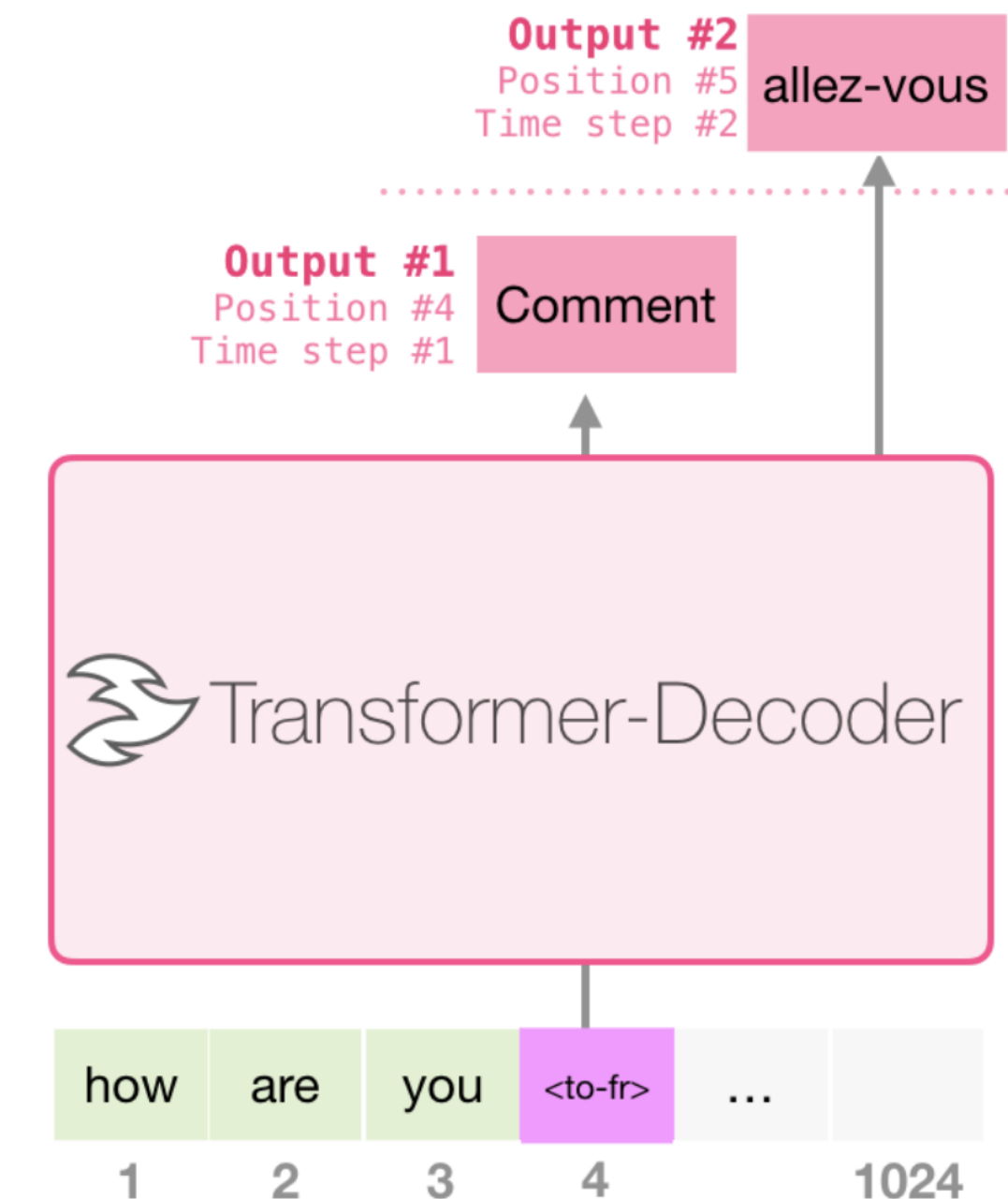


Improving language understanding by generative pre-training (Radford et al, 2018)

GPT-2

- Express all tasks as a language modelling task
- Training improvements
 - Improved initialization / additional layer normalization
 - Increased vocabulary / context /batch size
- Machine Translation

I	am	a	student	<to-fr>	je	suis	étudiant
let	them	eat	cake	<to-fr>	Qu'ils	mangent	de
good	morning	<to-fr>	Bonjour				



(figure credit: Jay Alammar
<http://jalammar.github.io/illustrated-gpt2/>)

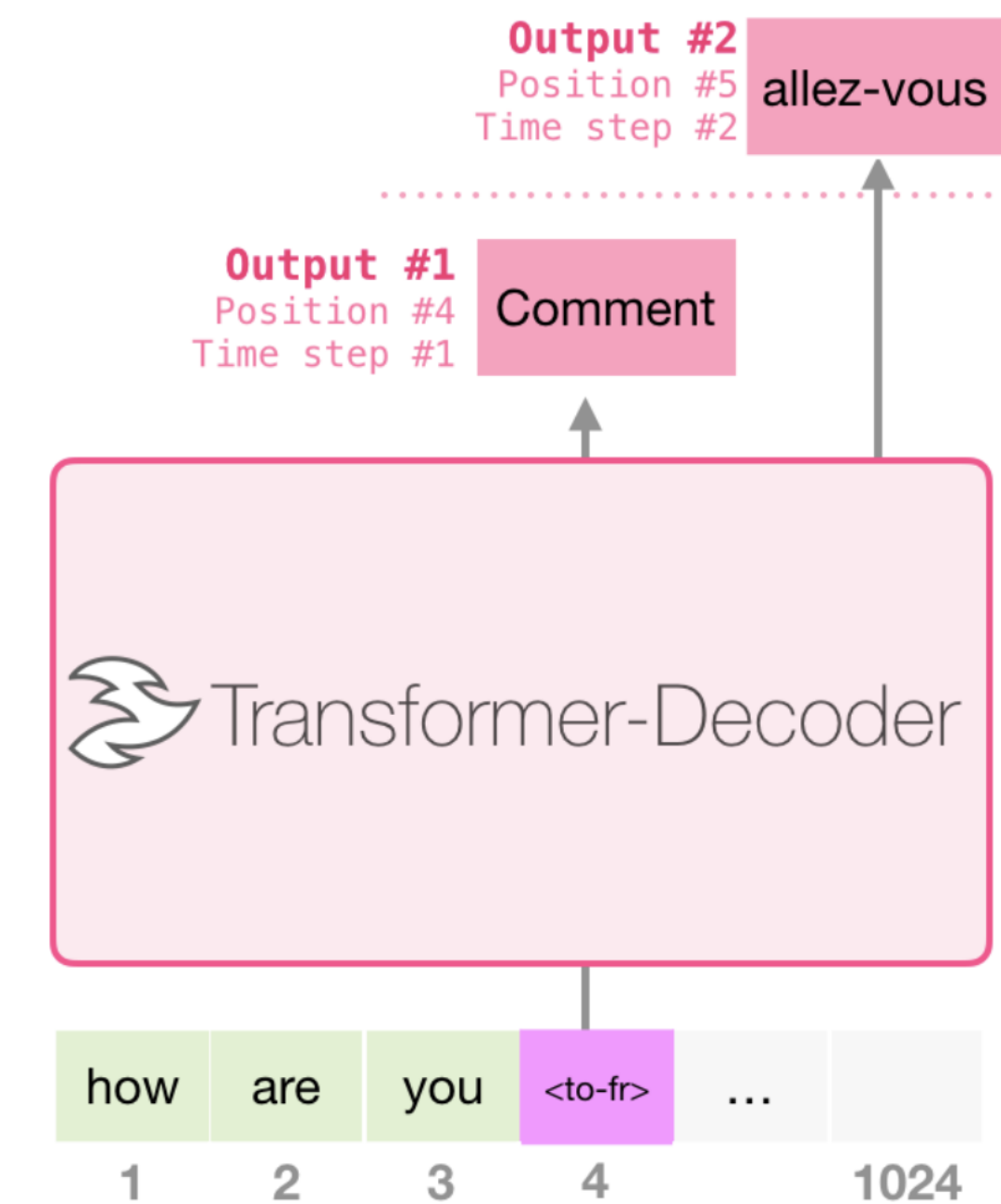
GPT-2

How can we use decoders for different tasks?

- Use special token to indicate task

Machine Translation

I	am	a	student	<to-fr>	je	suis	étudiant
let	them	eat	cake	<to-fr>	Qu'ils	mangent	de
good	morning	<to-fr>	Bonjour				



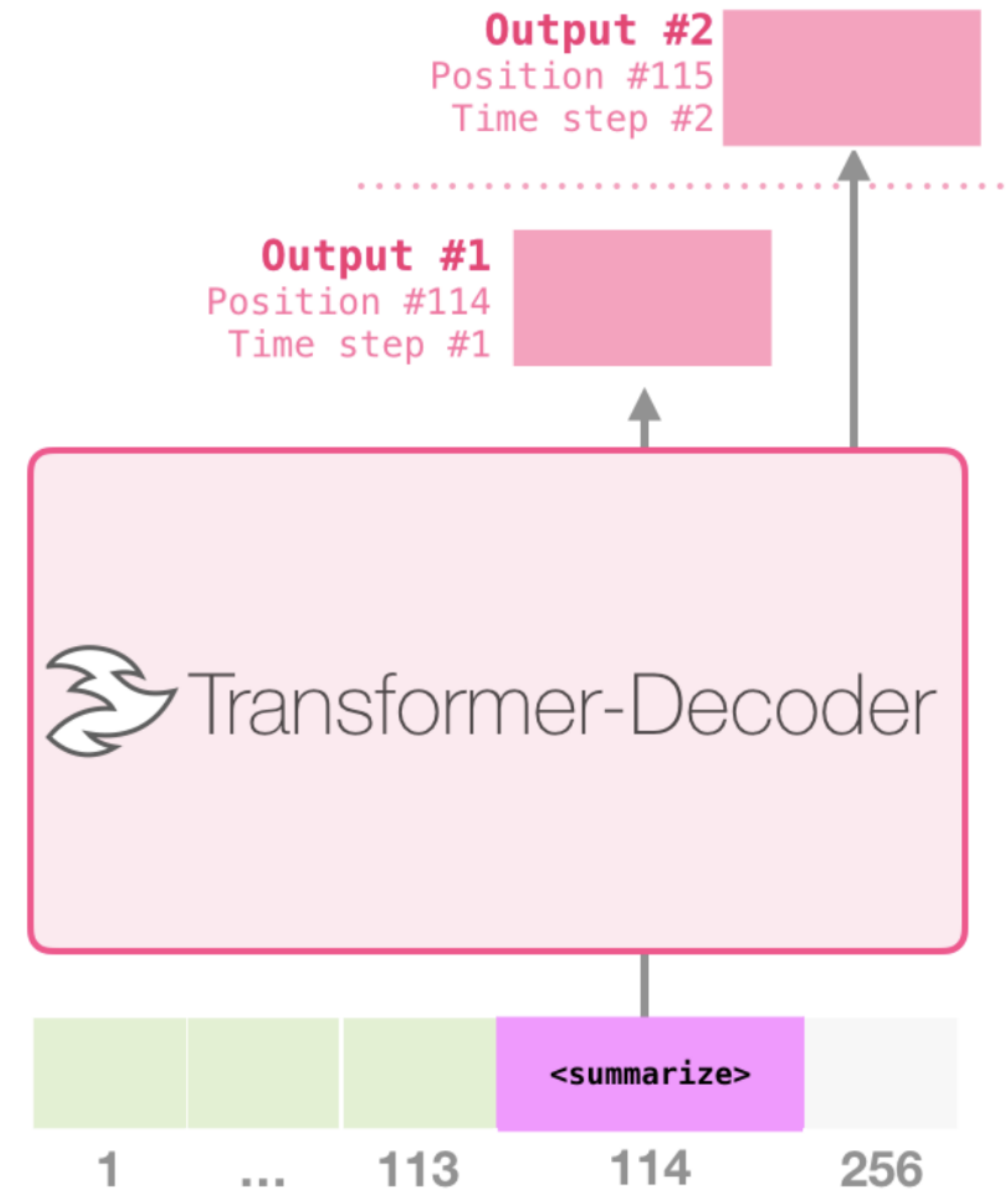
Summarization

Article #1 tokens		<summarize>	Article #1 Summary
Article #2 tokens	<summarize>	Article #2 Summary	padding
Article #3 tokens		<summarize>	Article #3 Summary

GPT-2

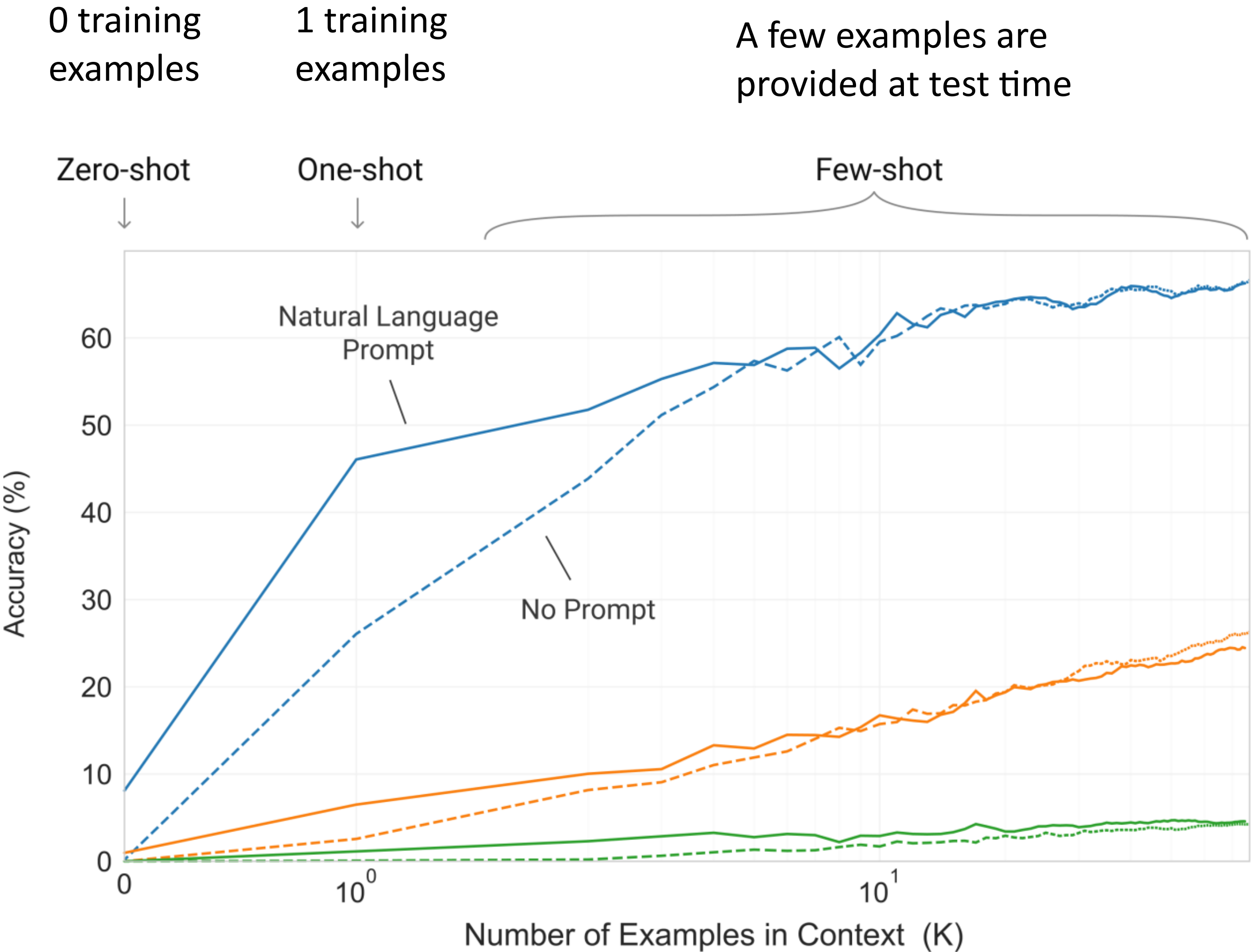
How can we use decoders for different tasks?

- Use special token to indicate task
- Summarization



(figure credit: Jay Alammar
<http://jalammar.github.io/illustrated-gpt2/>)

GPT-3: Few-shot learning



Zero-shot

The model predicts the answer given only a natural language description of the task. No gradient updates are performed.

```
1 Translate English to French: ← task description
2 cheese => ..... ← prompt
```

One-shot

In addition to the task description, the model sees a single example of the task. No gradient updates are performed.

```
1 Translate English to French: ← task description
2 sea otter => loutre de mer ← example
3 cheese => ..... ← prompt
```

Few-shot

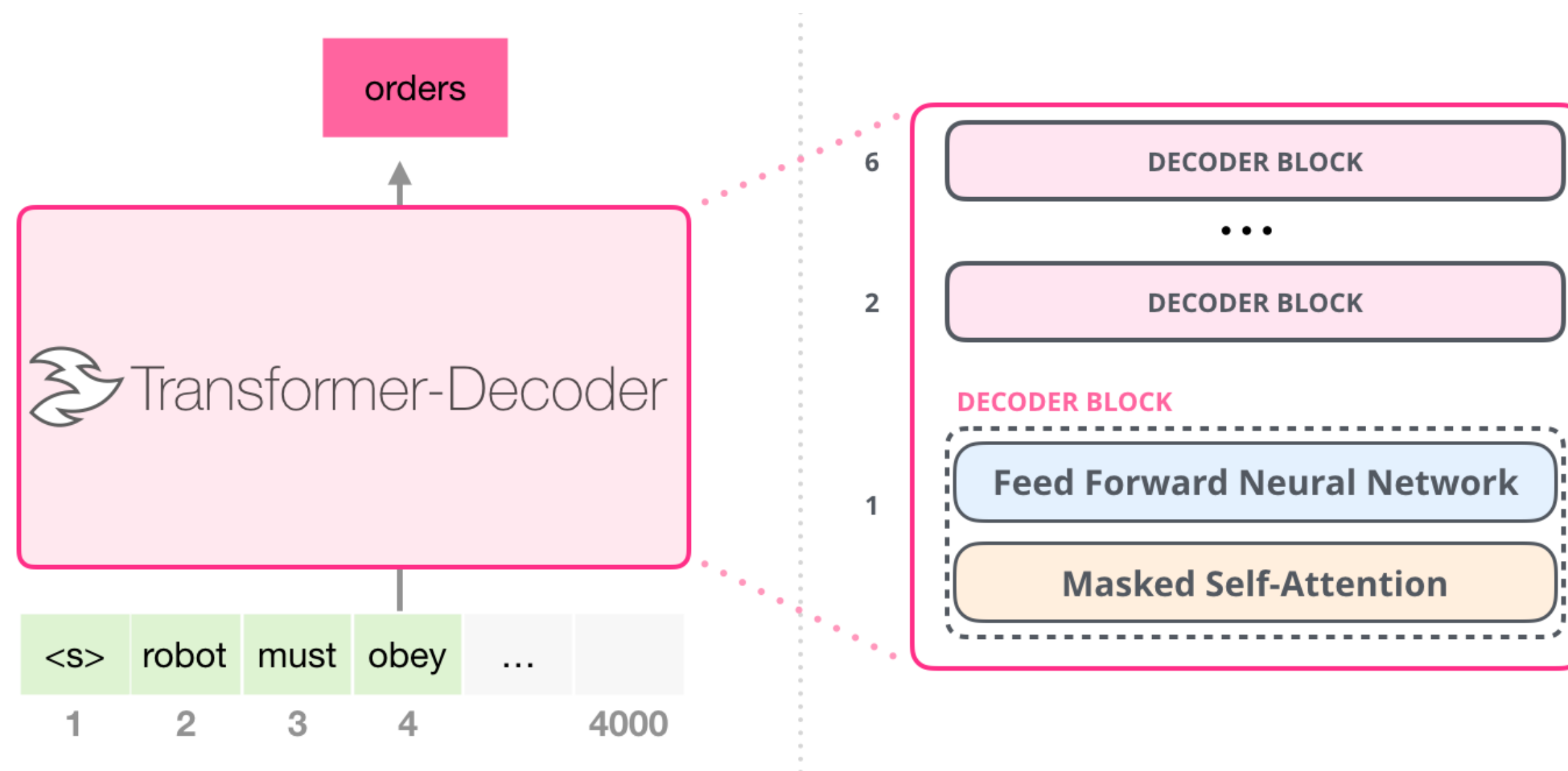
In addition to the task description, the model sees a few examples of the task. No gradient updates are performed.

```
1 Translate English to French: ← task description
2 sea otter => loutre de mer ← examples
3 peppermint => menthe poivrée ←
4 plush girafe => girafe peluche ←
5 cheese => ..... ← prompt
```

Autoregressive decoder-only models

GPT

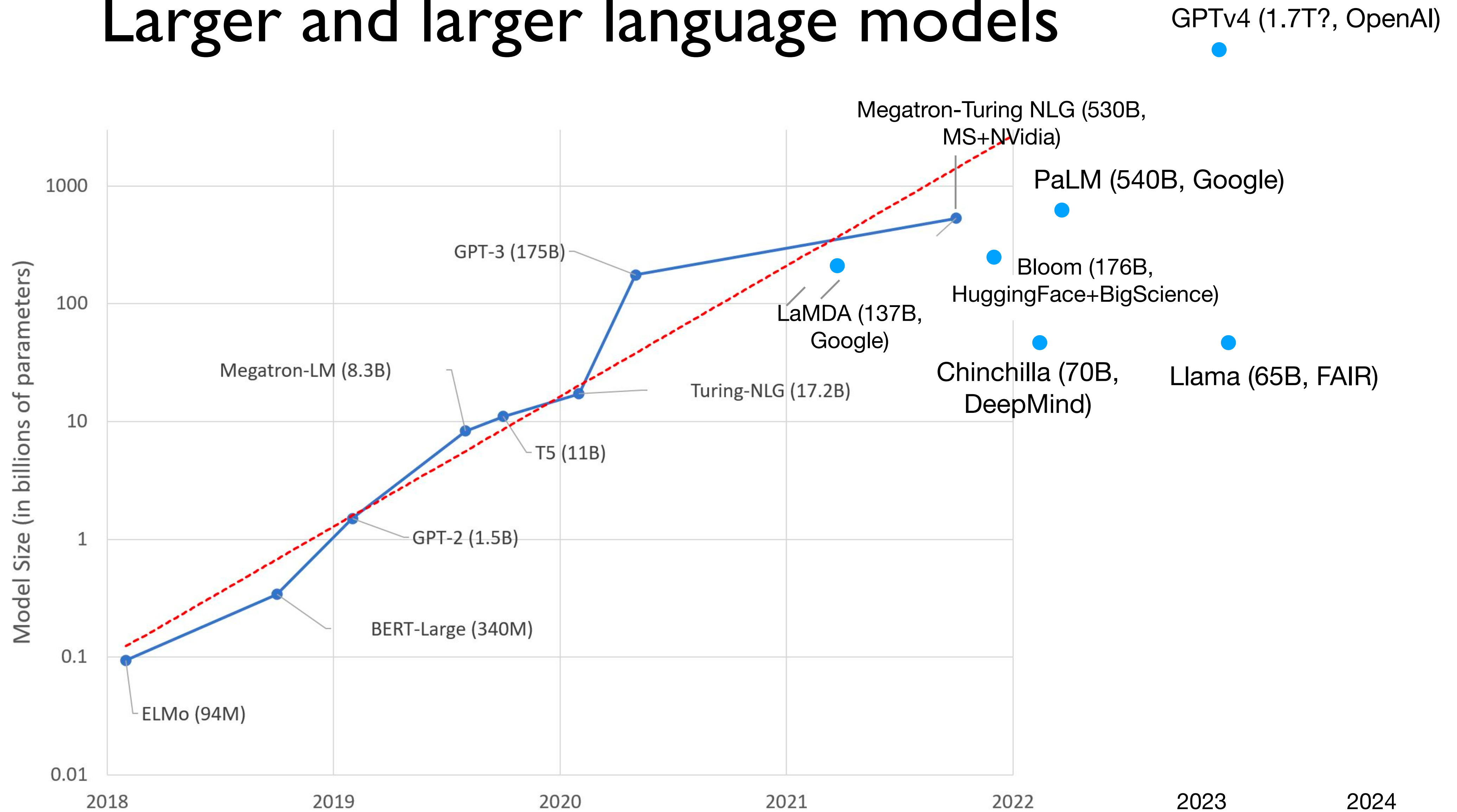
Rise of LLMs



- Multi-task training: modelling all tasks as autoregressive language modeling
- Scaling up to lots and lots lots and hundreds of billions of parameters
- Scaling up requires system engineering, tweaks to architecture for training stability
- Multi-lingual, multi-modal...

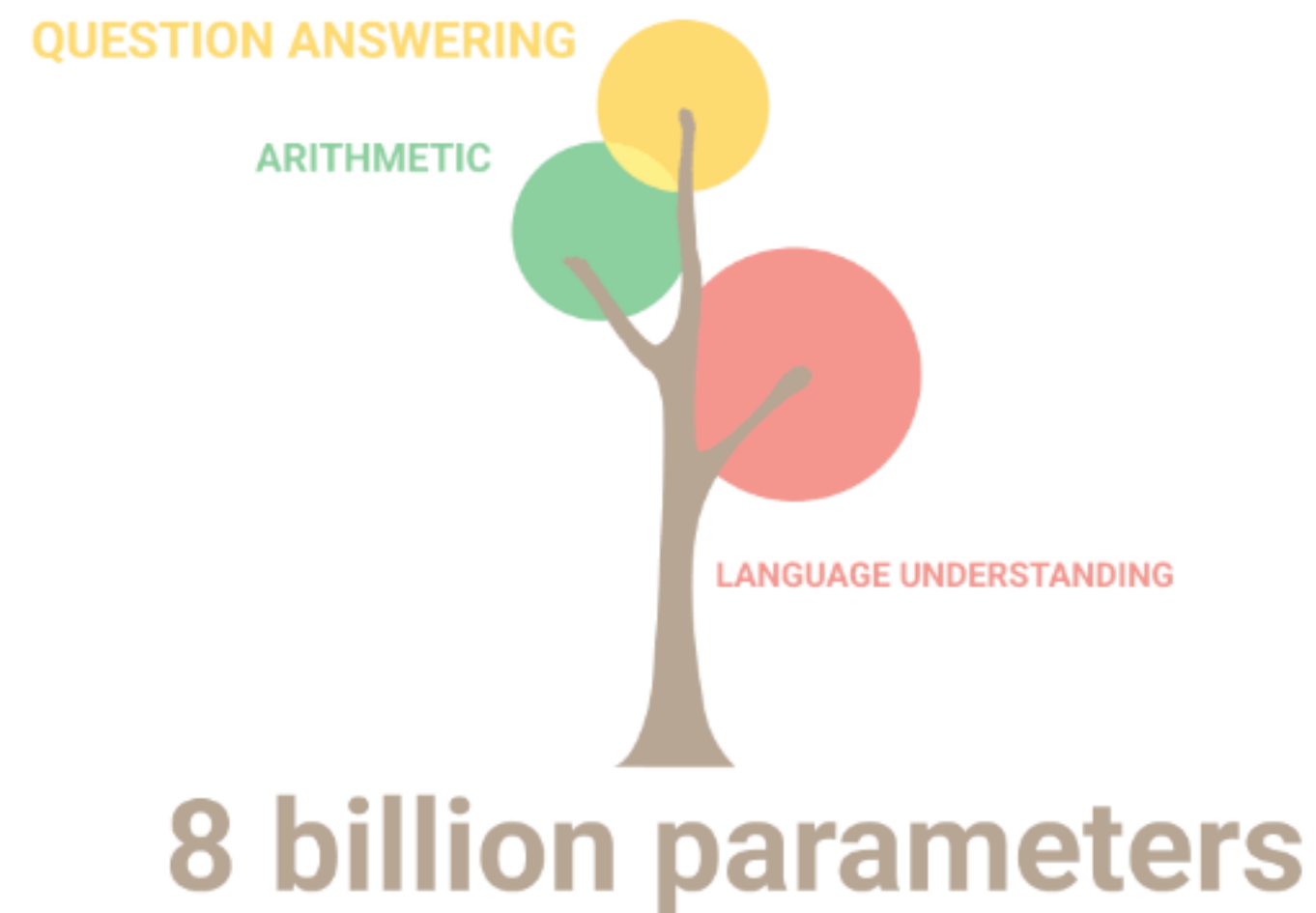
Objectives: next token prediction

Larger and larger language models



<https://huggingface.co/blog/large-language-models>

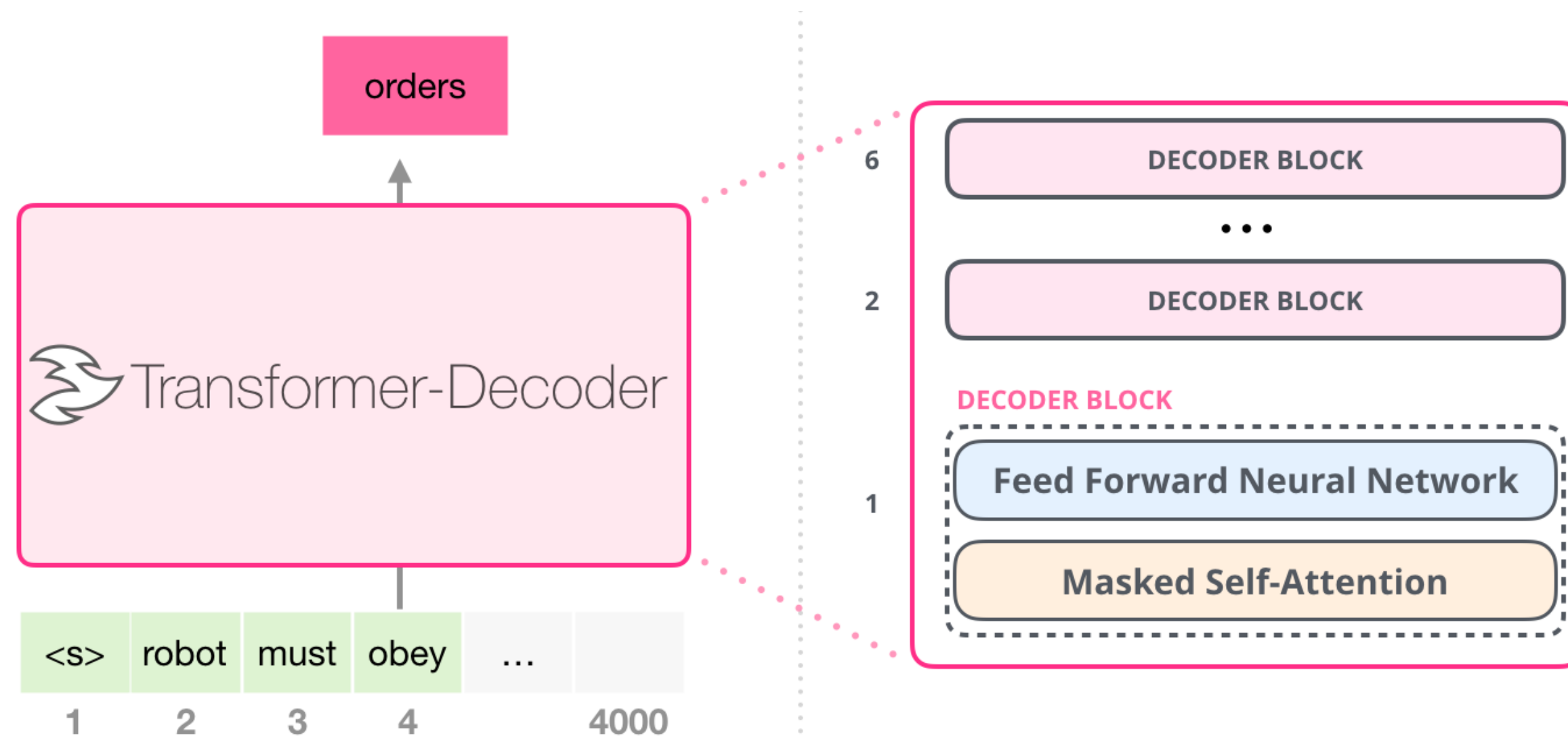
New capabilities emerge at scale



<https://ai.googleblog.com/2022/04/pathways-language-model-palm-scaling-to.html>

Autoregressive decoder-only models

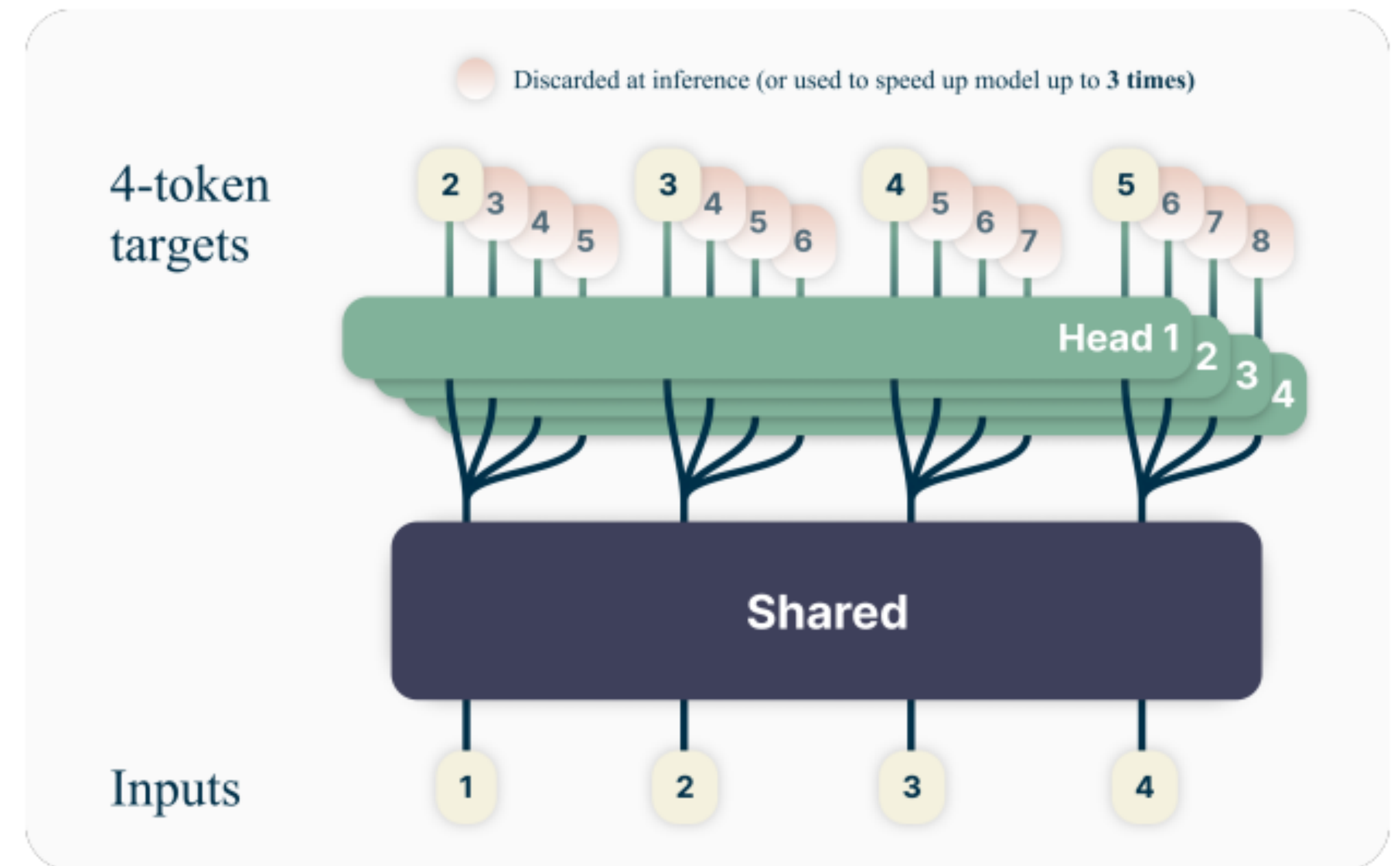
GPT



<https://jalammar.github.io/illustrated-gpt2/>

Objectives: next token prediction

Advances



Objectives: multi token prediction

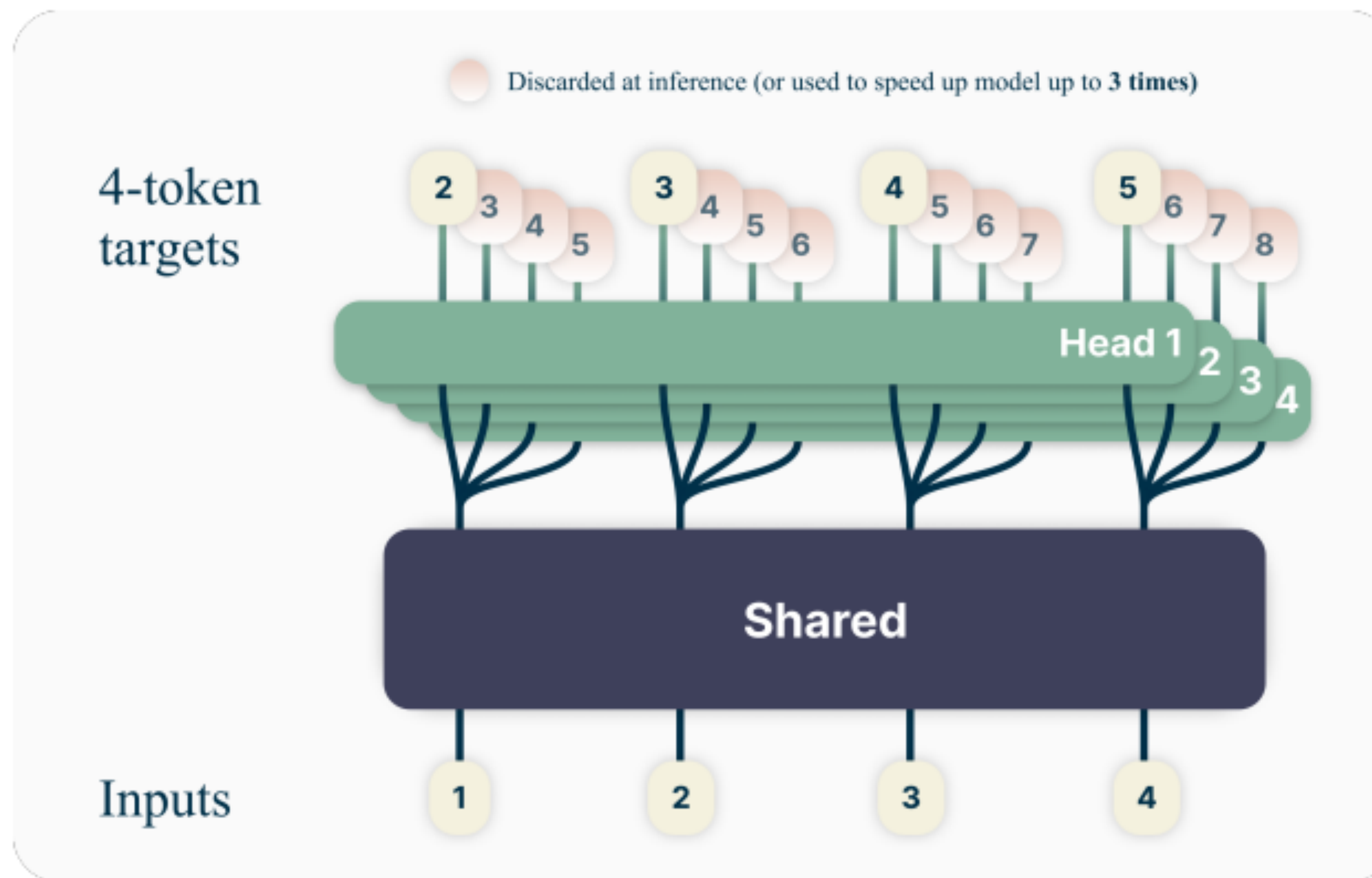
Multi-token prediction

- Predict multiple next tokens

$$\begin{aligned} L_n &= - \sum_t \log P_\theta(x_{t+n:t+1} \mid z_{t:1}) \cdot P_\theta(z_{t:1} \mid x_{t:1}) \\ &= - \sum_t \sum_{i=1}^n \log P_\theta(x_{t+i} \mid z_{t:1}) \cdot P_\theta(z_{t:1} \mid x_{t:1}). \end{aligned}$$

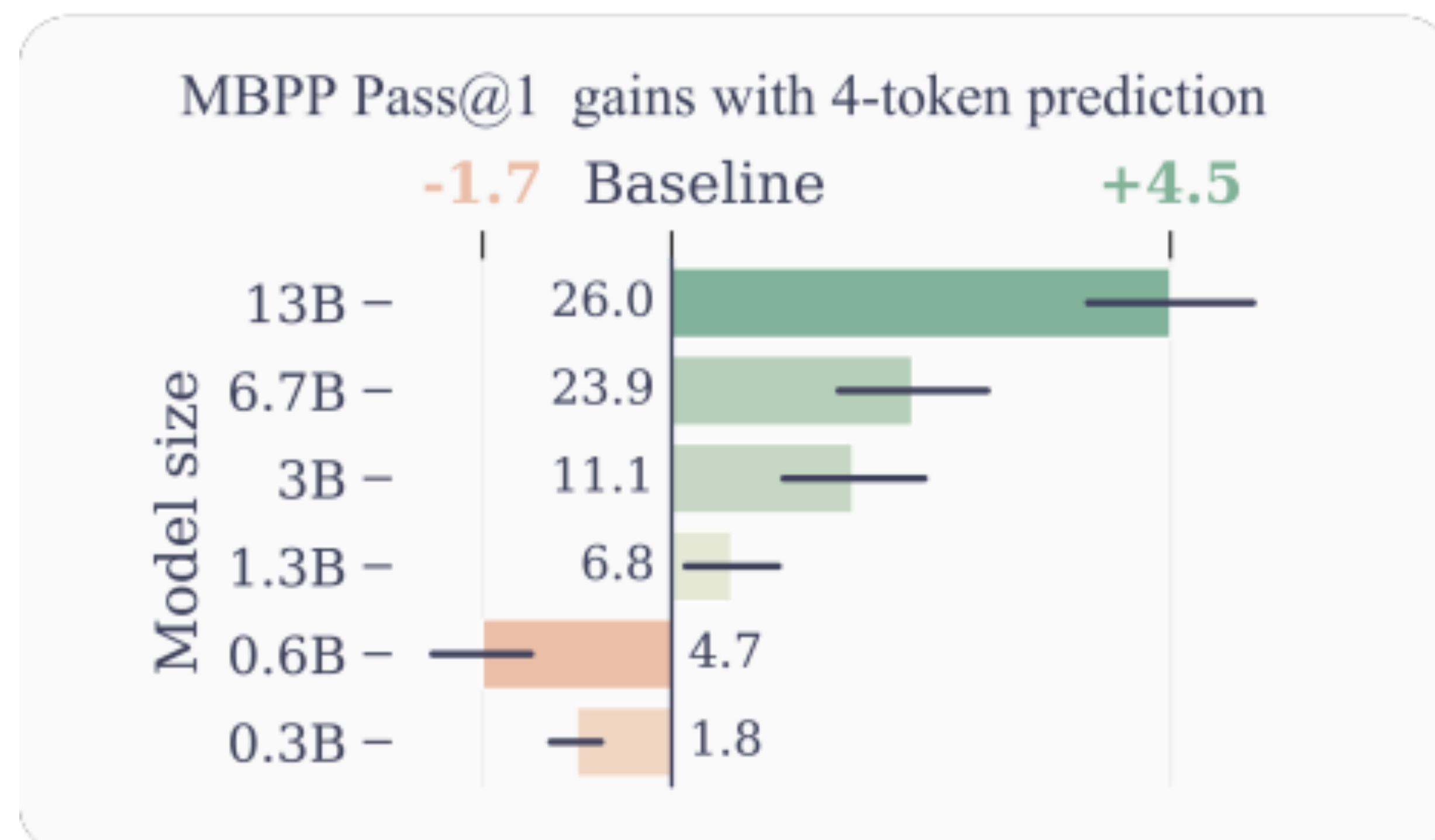
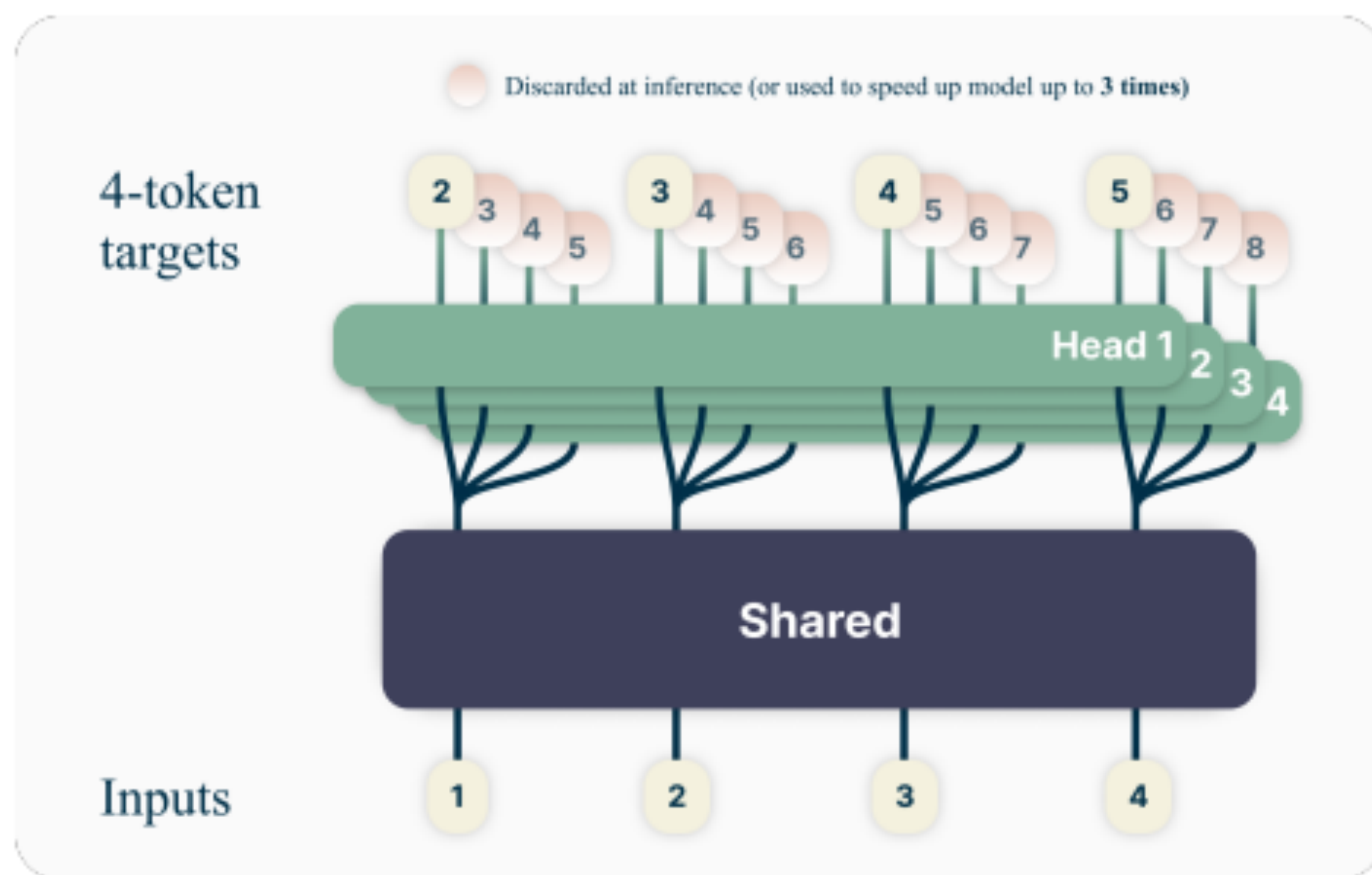
- Shared trunk / unembedding matrix

$$P_\theta(x_{t+i} \mid x_{t:1}) = \text{softmax}(f_u(f_{h_i}(f_s(x_{t:1}))))).$$



Multi-token prediction

- Predict next tokens with shared trunk
- Gains due to multi-token prediction



Multi-token prediction

- Predict in sequence to avoid large memory demands

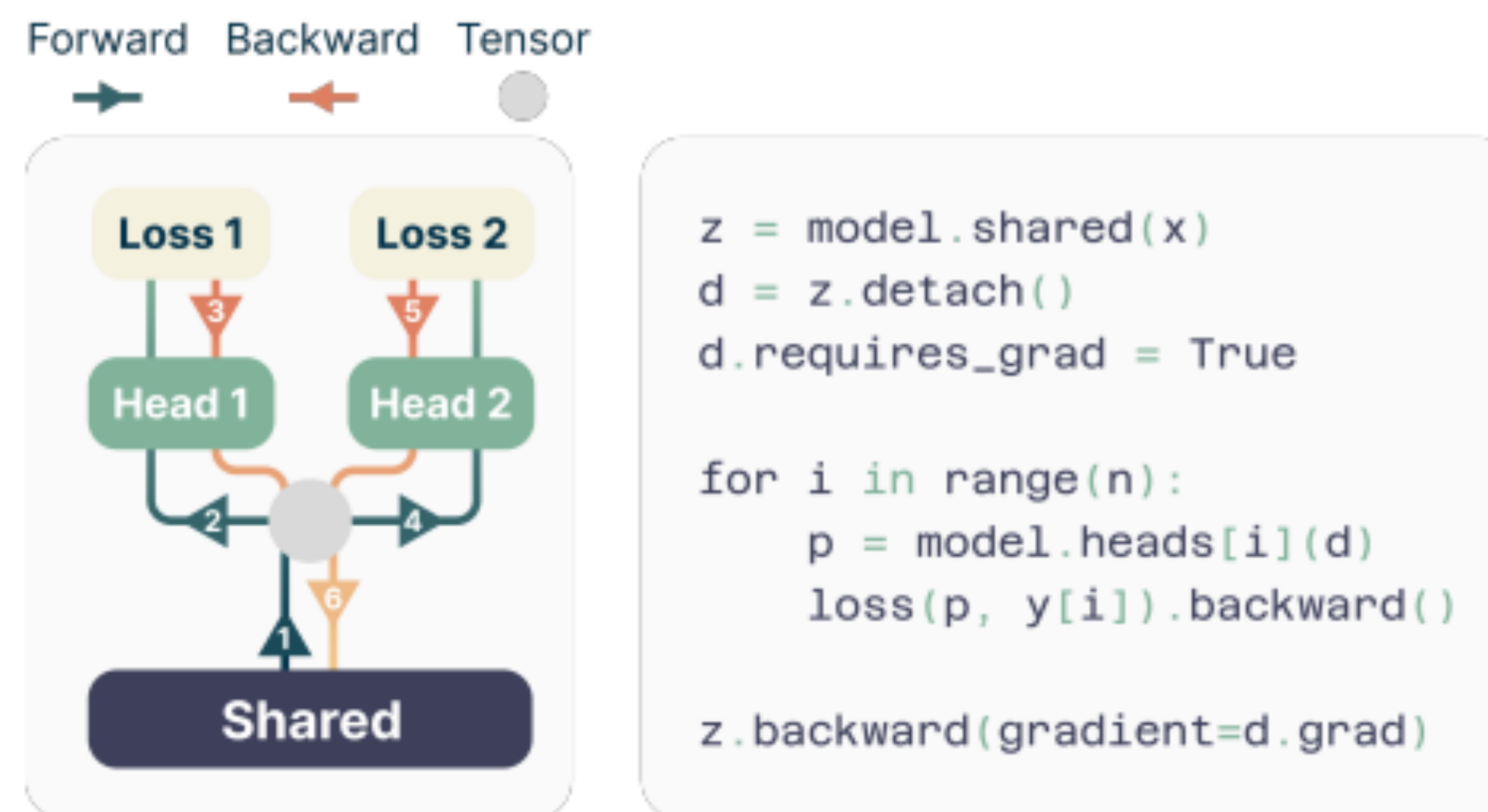


Figure 2: **Order of the forward/backward in an n -token prediction model with $n = 2$ heads.** By performing the forward/backward on the heads in sequential order, we avoid materializing all unembedding layer gradients in memory simultaneously and reduce peak GPU memory usage.

Multi-token prediction

- Predict multiple next tokens
- D sequential modules for predicting D tokens

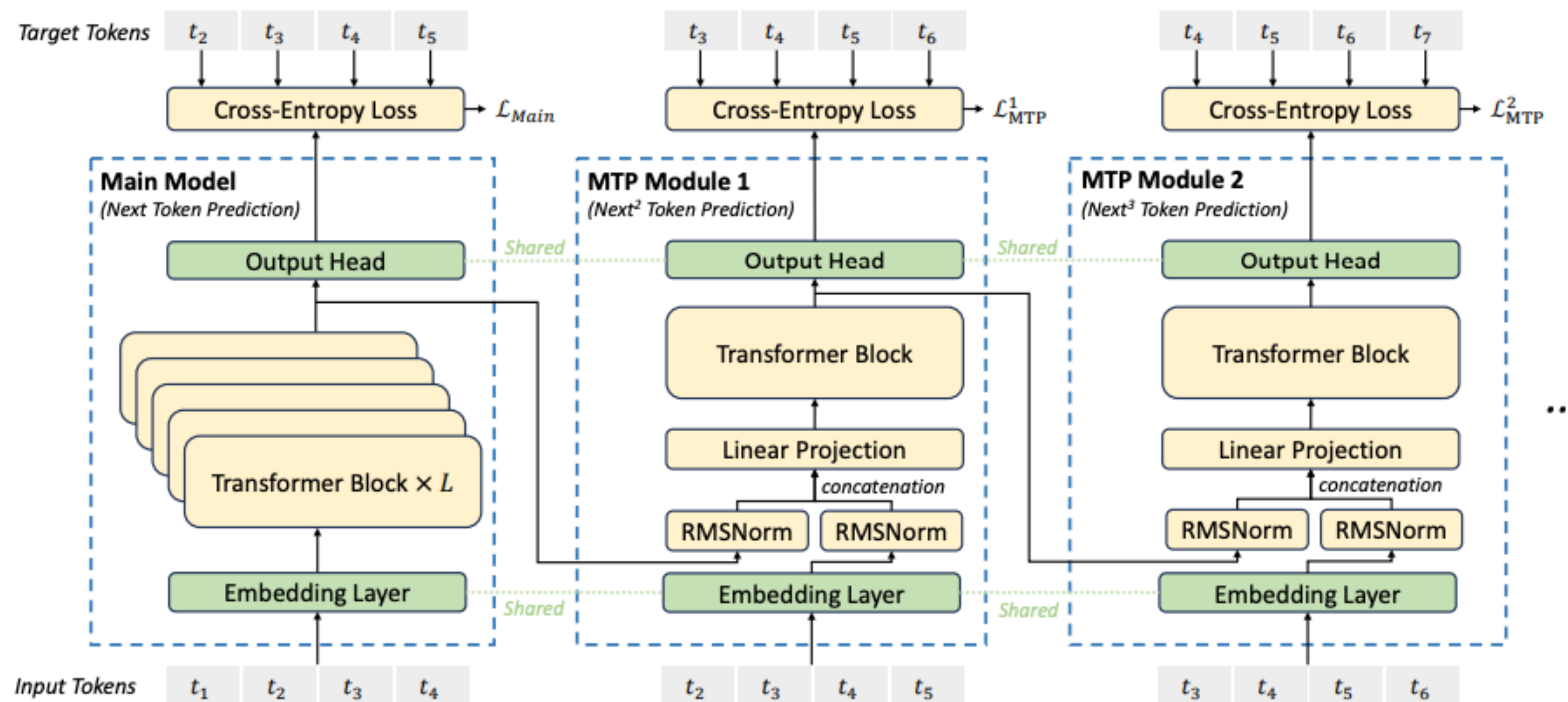


Figure 3 | Illustration of our Multi-Token Prediction (MTP) implementation. We keep the complete causal chain for the prediction of each token at each depth.

<https://arxiv.org/abs/2412.19437>

Continuous language modelling

- Autoencoder to encode K tokens into one vector
- Next vector prediction with efficient generative head

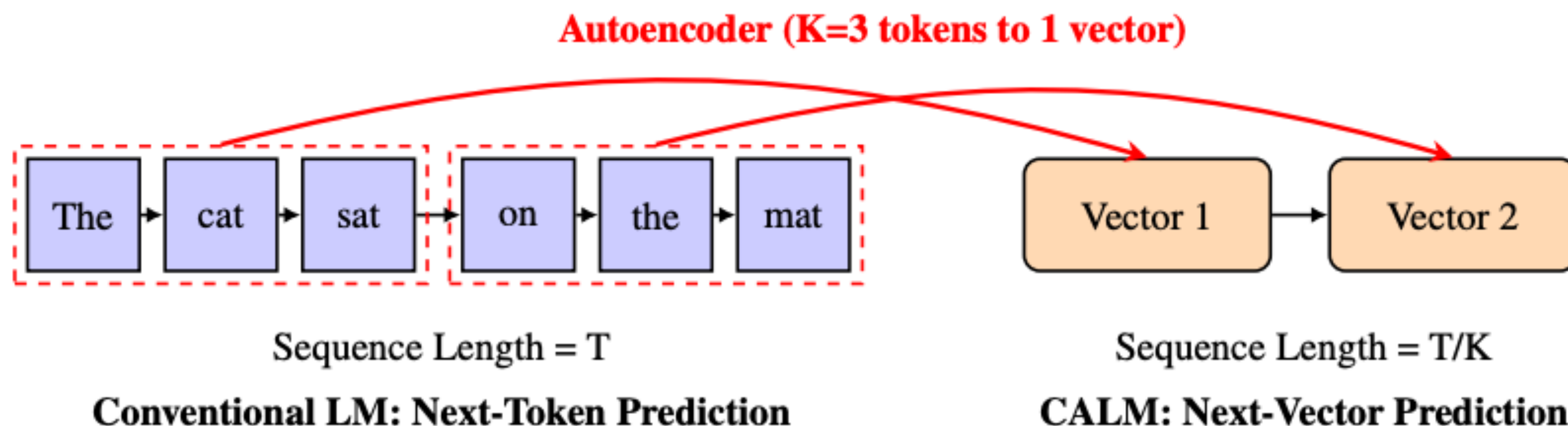
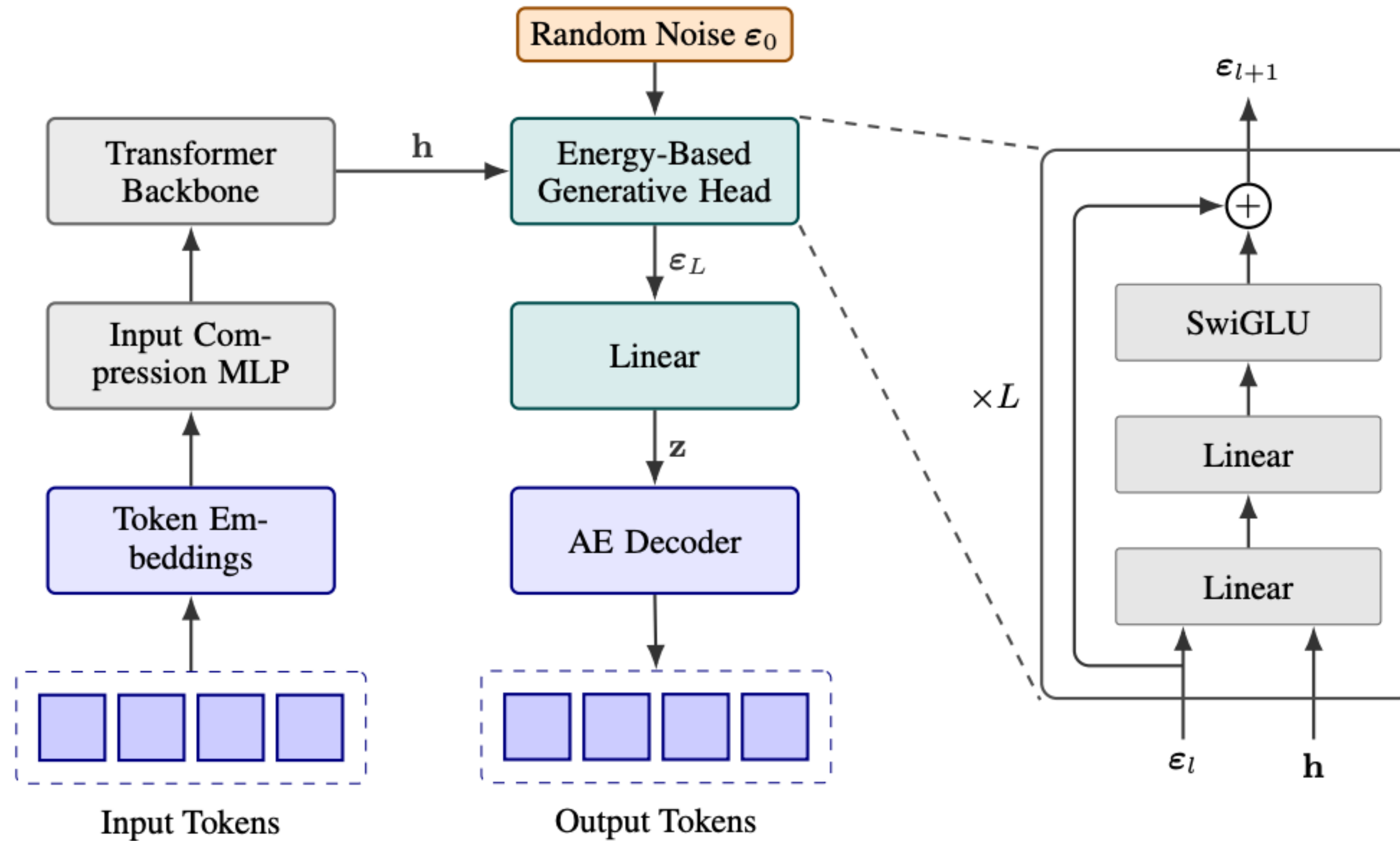


Figure 1: Comparison between conventional token-by-token generation and our proposed vector-by-vector framework (CALM). By compressing K tokens into a single vector, we reduce the sequence length K-fold, fundamentally improving computational efficiency.

Continuous language modelling

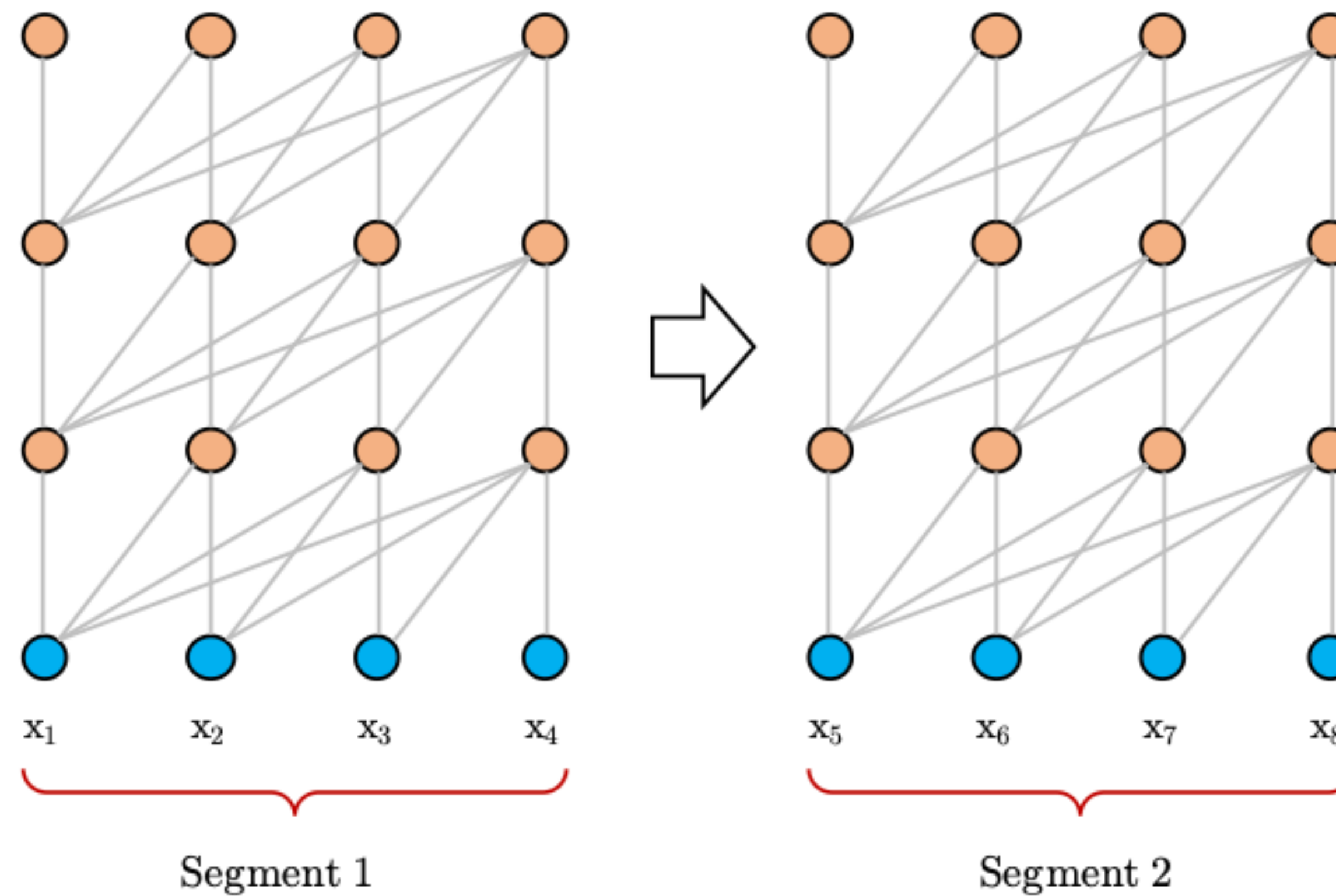


Transformer-XL

<https://arxiv.org/abs/1901.02860>

Dai+ 2019

- Vanilla Model



(a) Train phase.

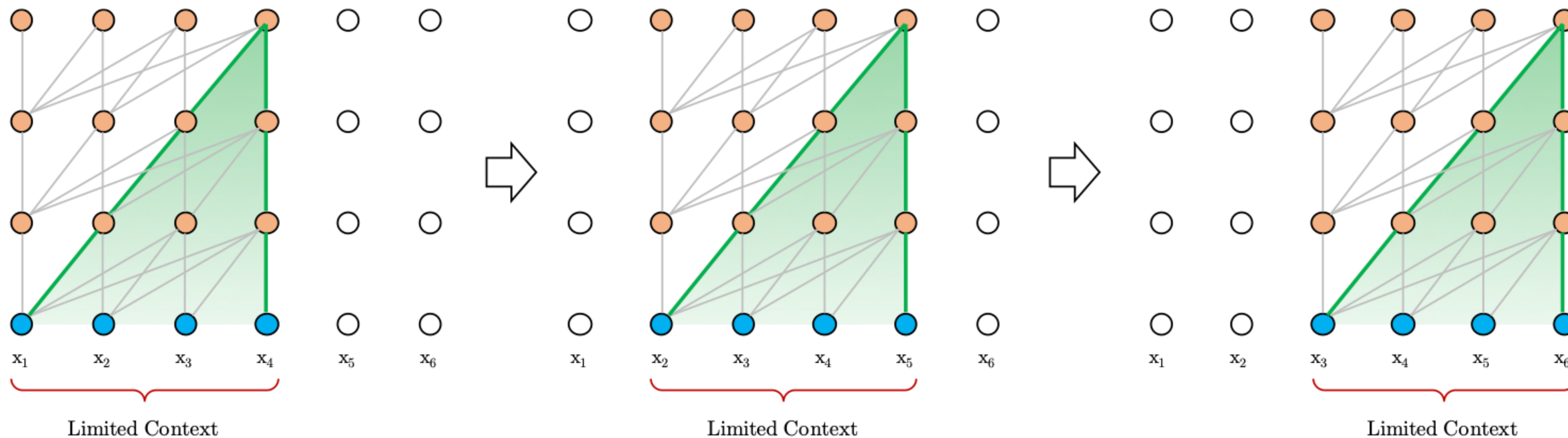
Transformer-XL

<https://arxiv.org/abs/1901.02860>

Dai+ 2019

Is there a better way to allow for long context?

- Vanilla Model



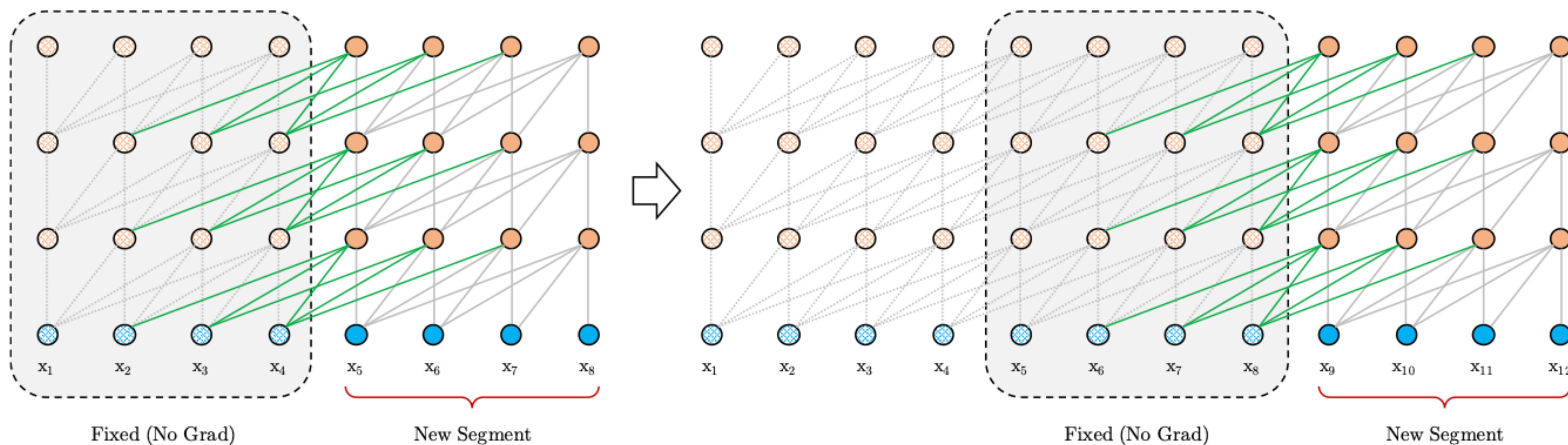
(b) Evaluation phase.

Transformer-XL

<https://arxiv.org/abs/1901.02860>

Dai+ 2019

- Autoregressive LM (different from GPT)
- segment level recurrence (reuse states) + relative positional embeddings



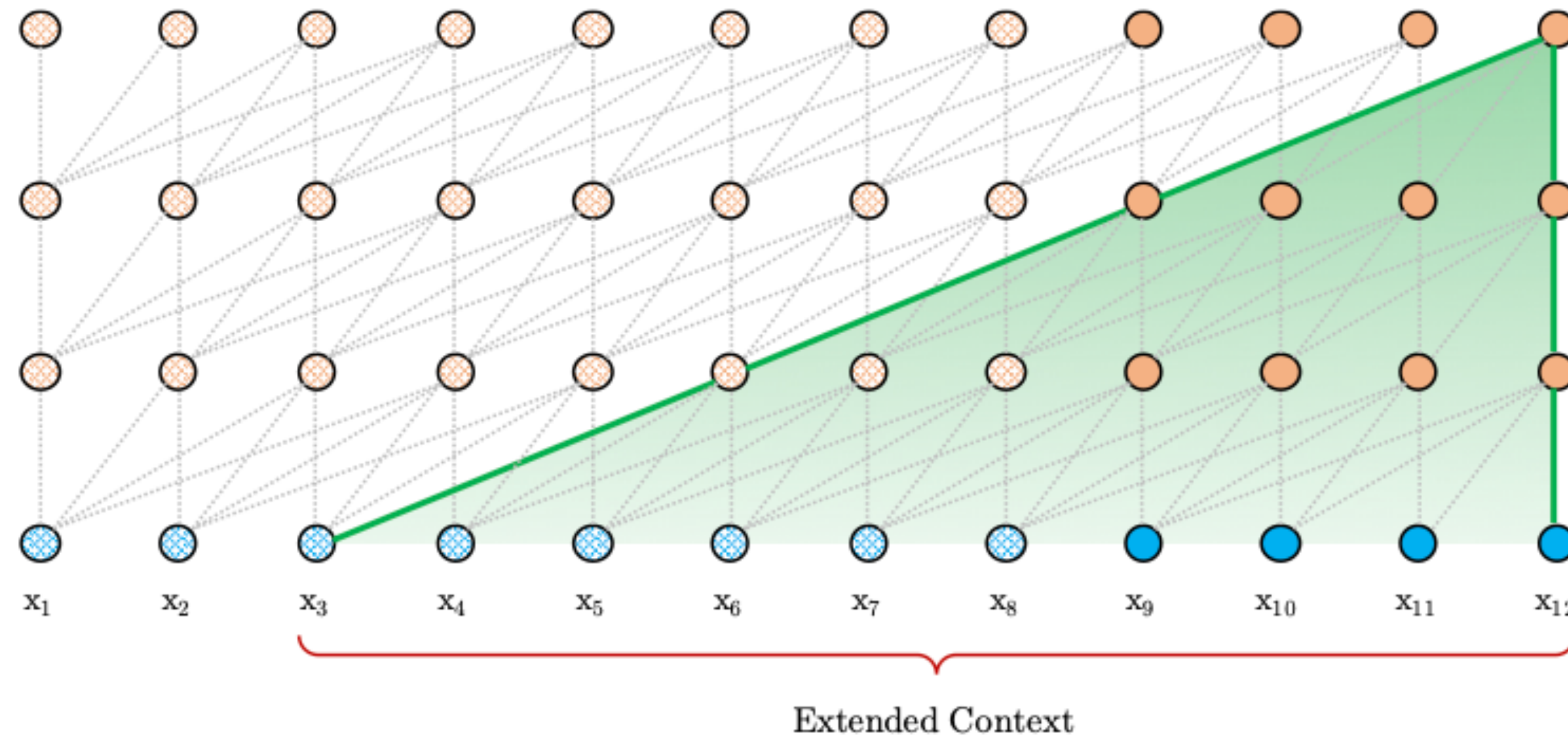
(a) Training phase.

Transformer-XL

<https://arxiv.org/abs/1901.02860>

Dai+ 2019

- Autoregressive LM (different from GPT)



(b) Evaluation phase.

XLNet

<https://arxiv.org/abs/1906.08237>

Yang+ 2019

- Autoregressive model for masked language modelling
 - Uses permutations (factorization order) to provide context
 - Allows for context from both sides through permutation
 - Avoid [MASK] token that does not appear in downstream tasks

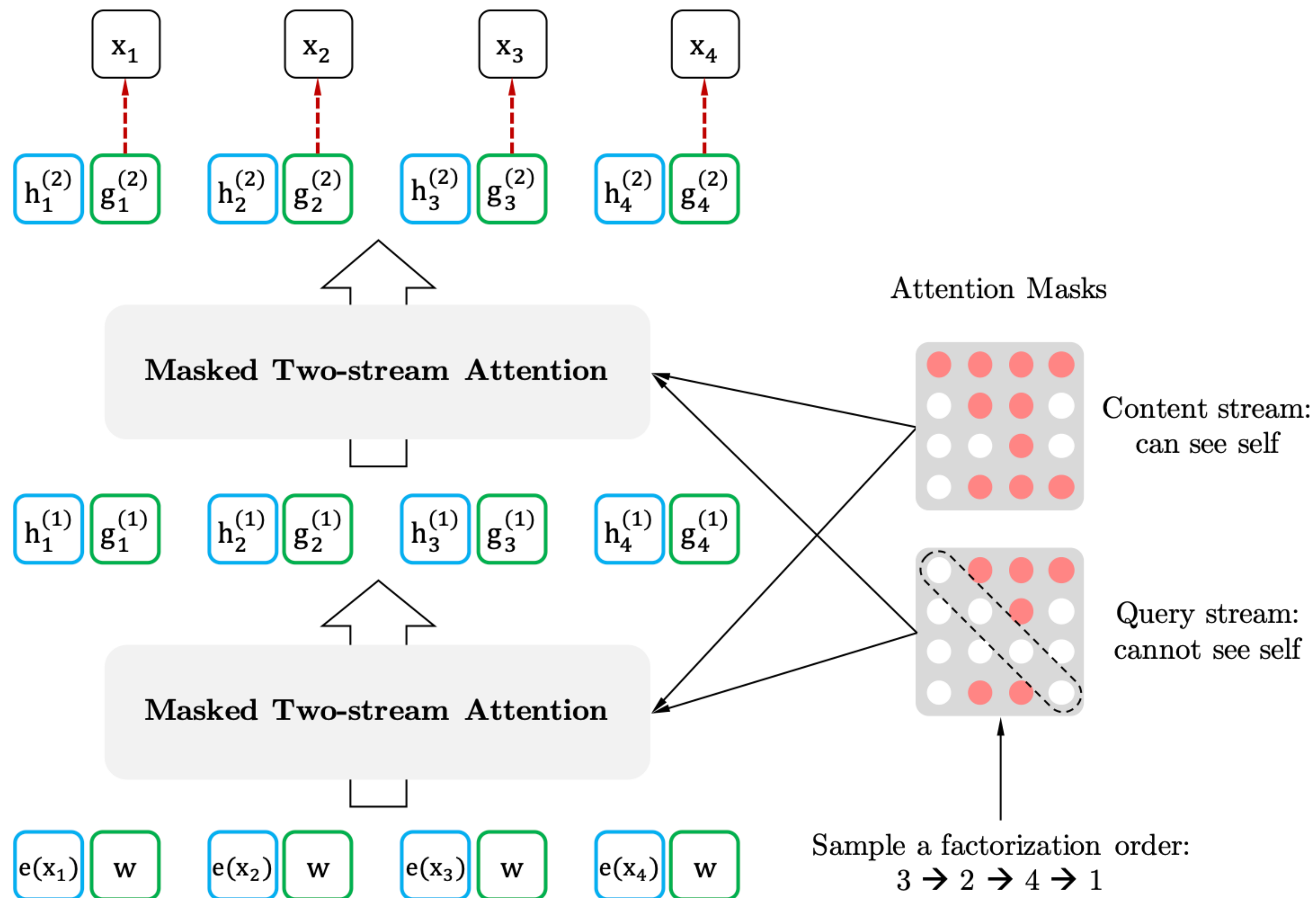
XLNet

<https://arxiv.org/abs/1906.08237>

Yang+ 2019

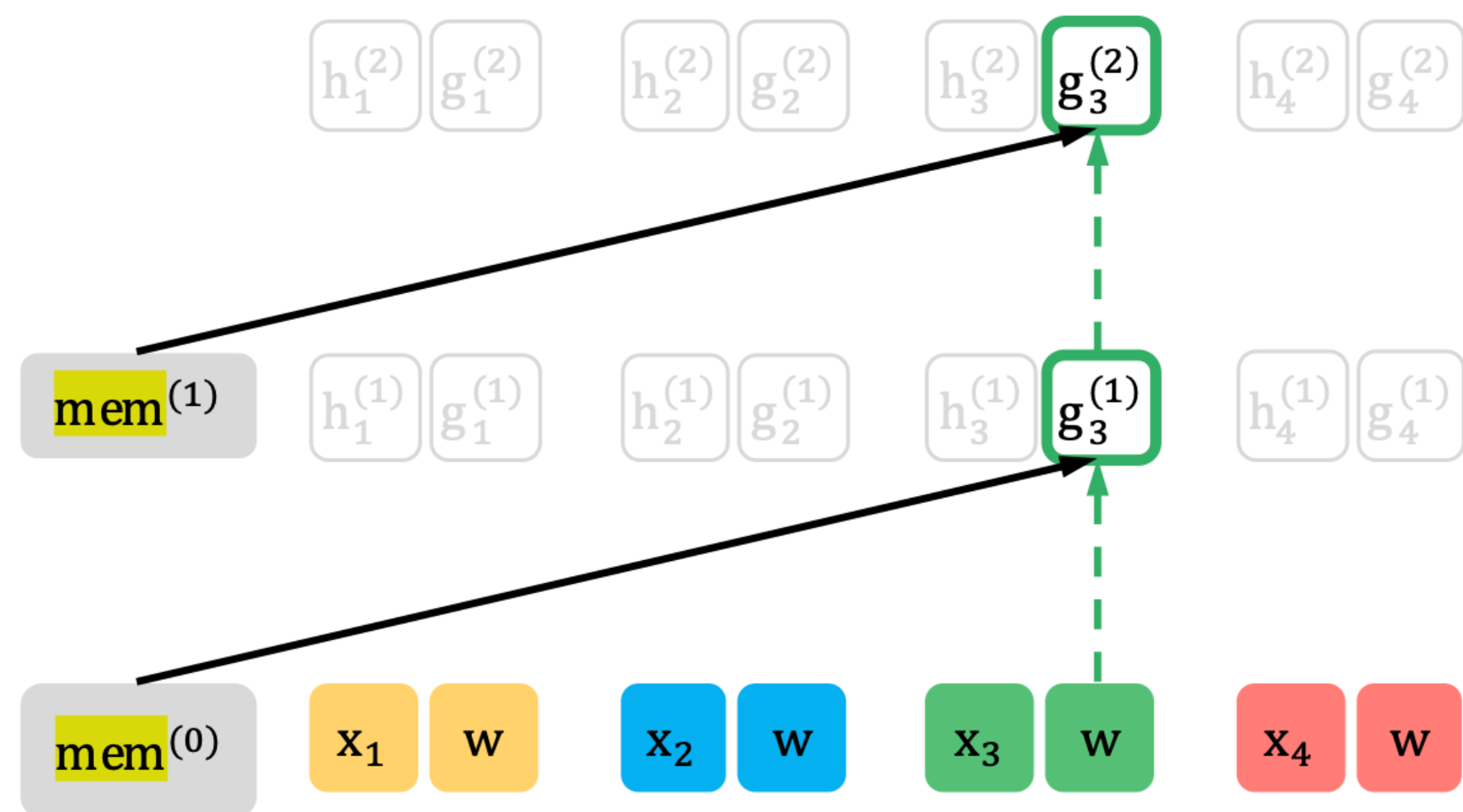
- Relative position embeddings (using auto-regressive TransformerXL)
 - Absolute attention: position $4 \rightarrow 5$; position $128 \rightarrow 129$
 - Relative attention: position $t \rightarrow (t - 1)$
- Mask prediction over all token positions using permutation on factorization order (sample a factorization order: $3 \rightarrow 2 \rightarrow 1 \rightarrow 4$)
 - Two stream self-attention: standard and query on [MASK] token
 - Permute only factorization order, not sequence order

XLNet

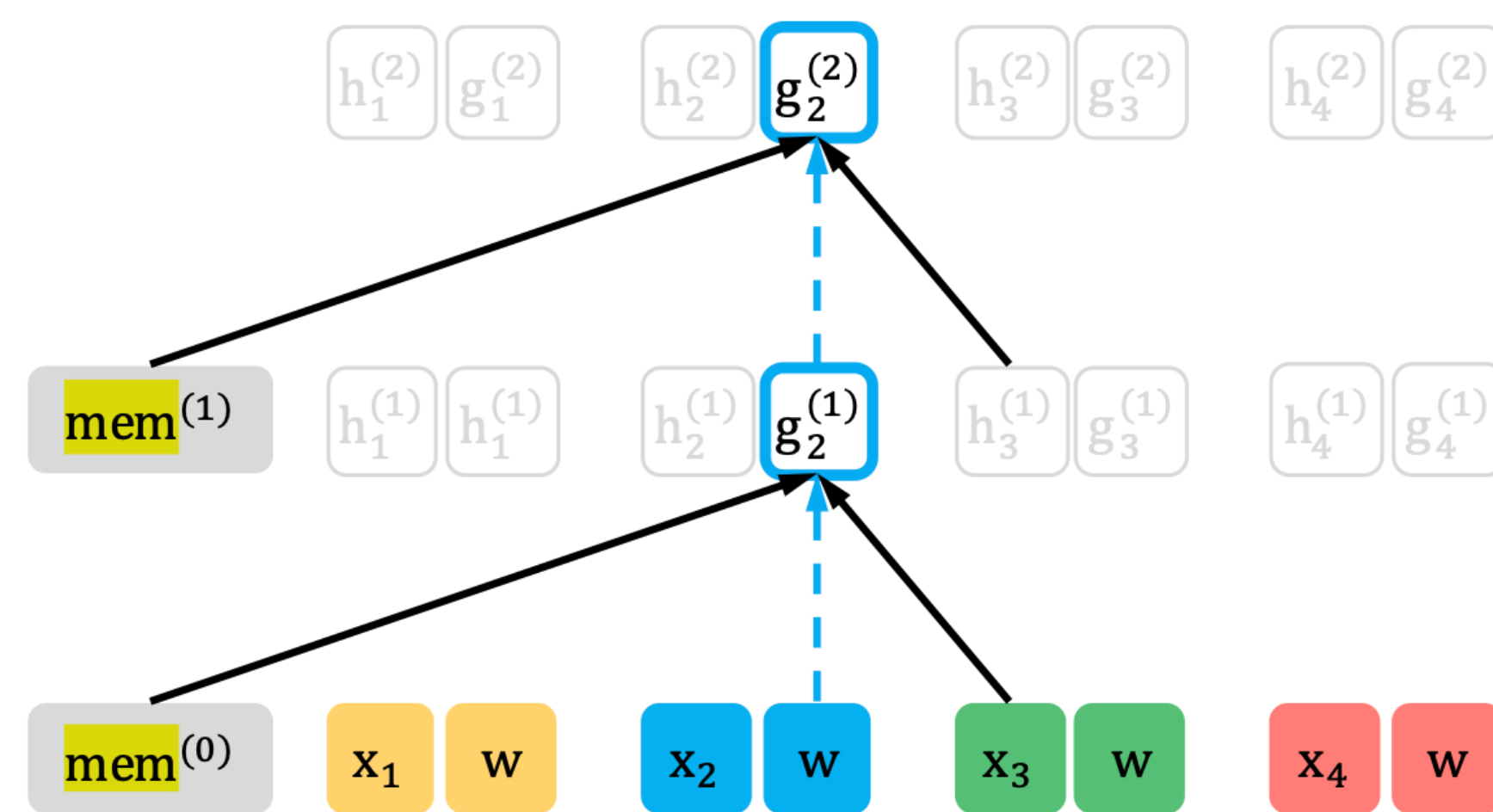


XLNet

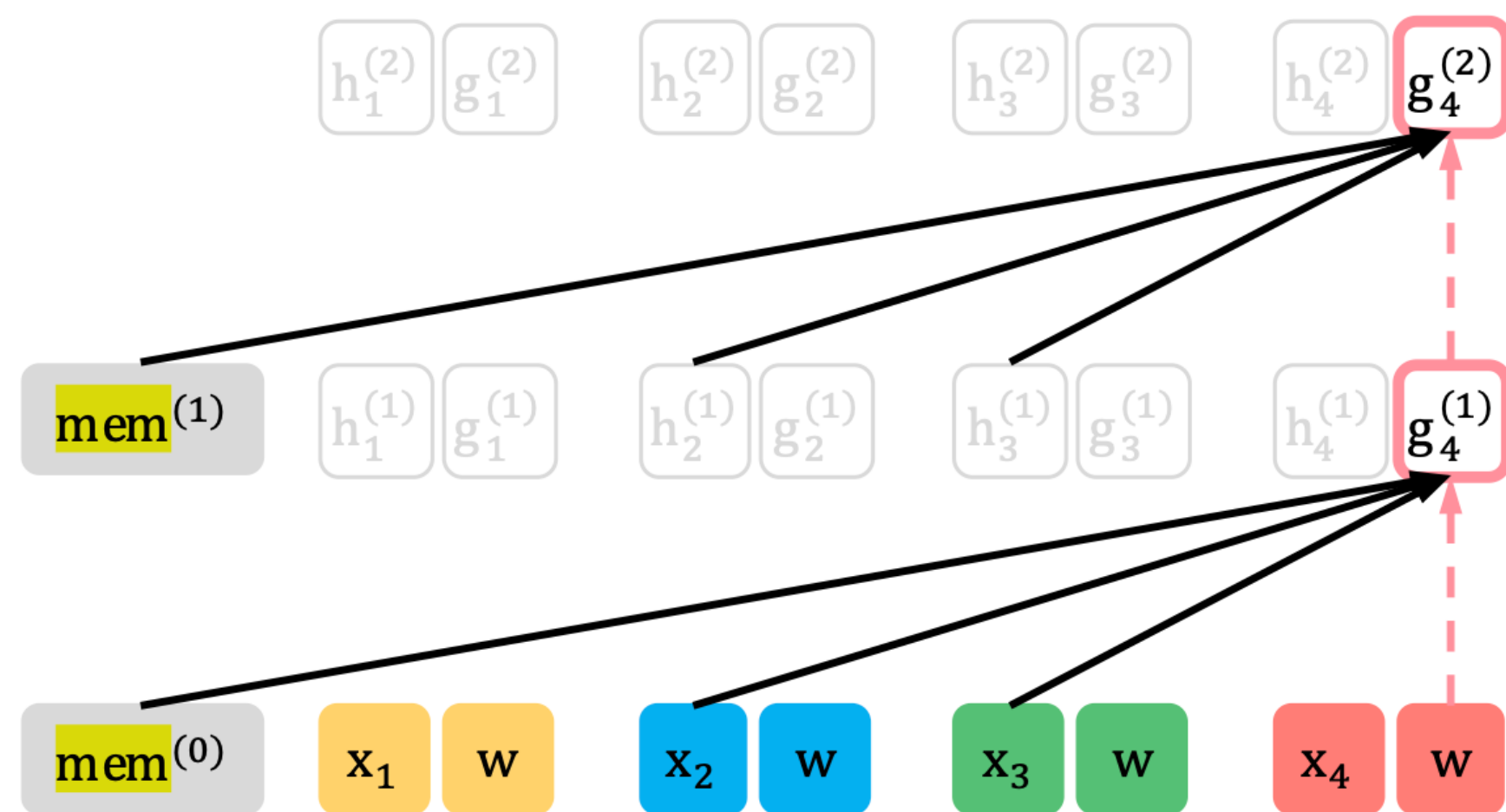
Split View of the Query Stream
(Factorization order: $3 \rightarrow 2 \rightarrow 4 \rightarrow 1$)



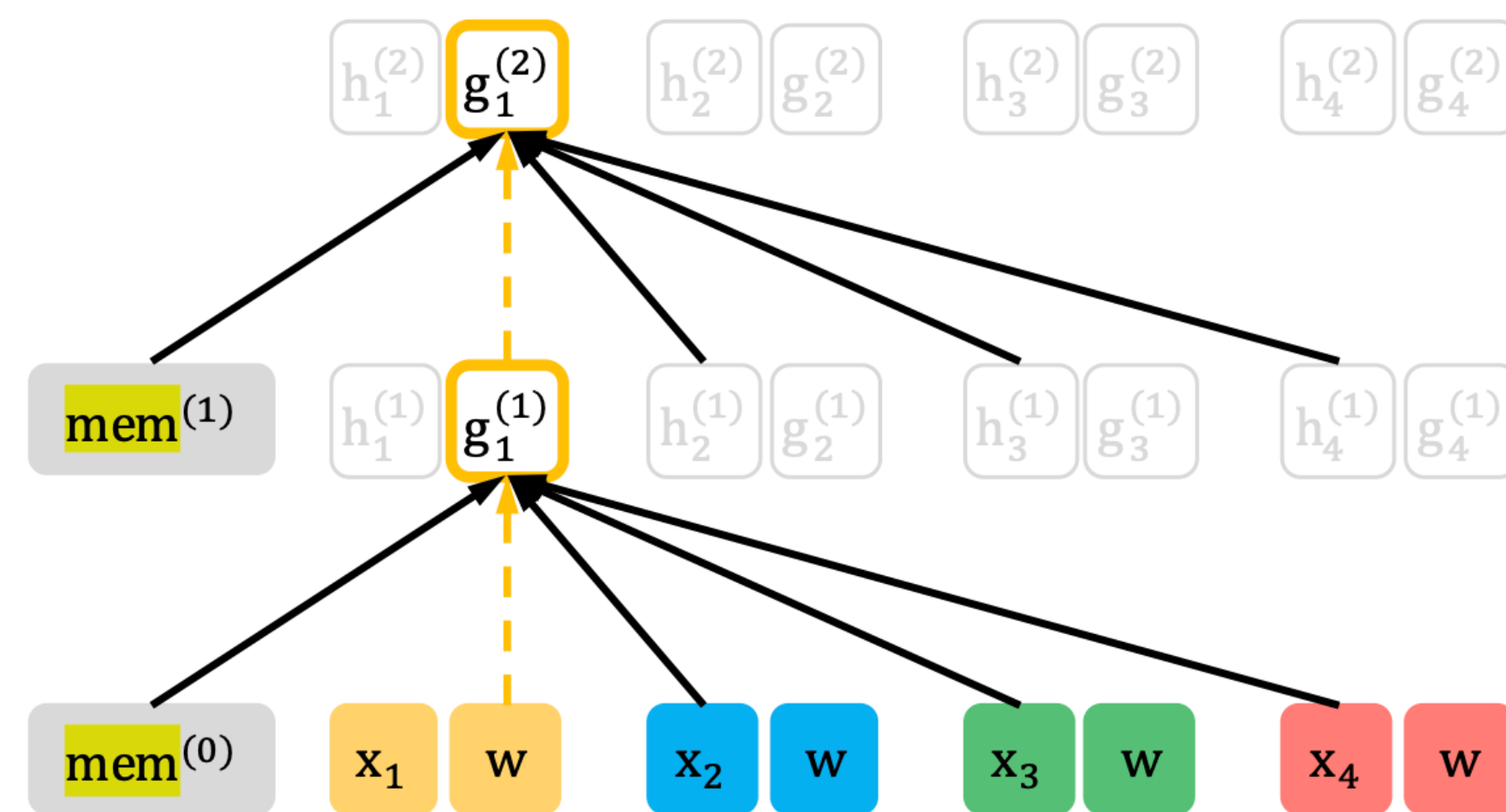
Position-3 View



Position-2 View



Position-4 View



Position-1 View

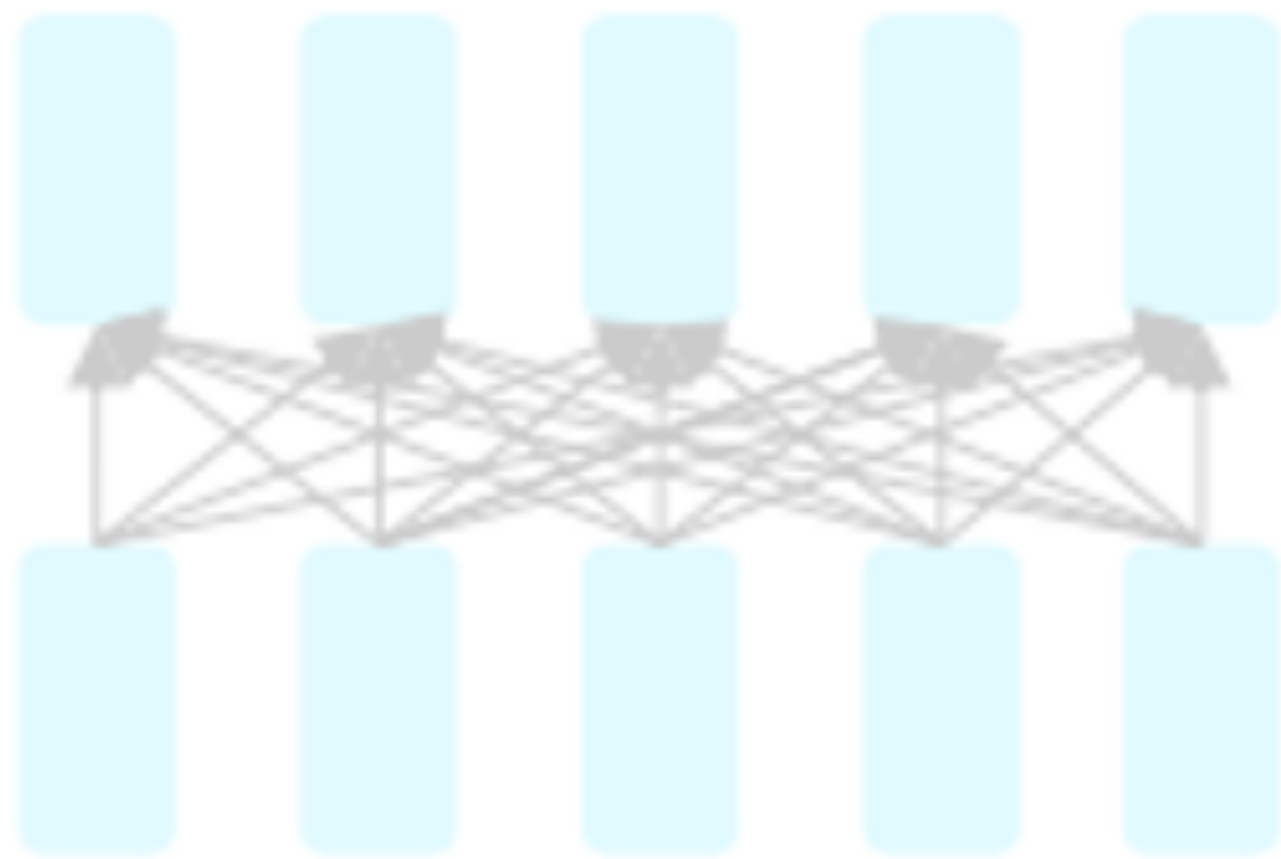
XLNet

Model	MNLI	QNLI	QQP	RTE	SST-2	MRPC	CoLA	STS-B
<i>Single-task single models on dev</i>								
BERT [2]	86.6/-	92.3	91.3	70.4	93.2	88.0	60.6	90.0
RoBERTa [21]	90.2/90.2	94.7	92.2	86.6	96.4	90.9	68.0	92.4
XLNet	90.8/90.8	94.9	92.3	85.9	97.0	90.8	69.0	92.5

Transformers for pretraining

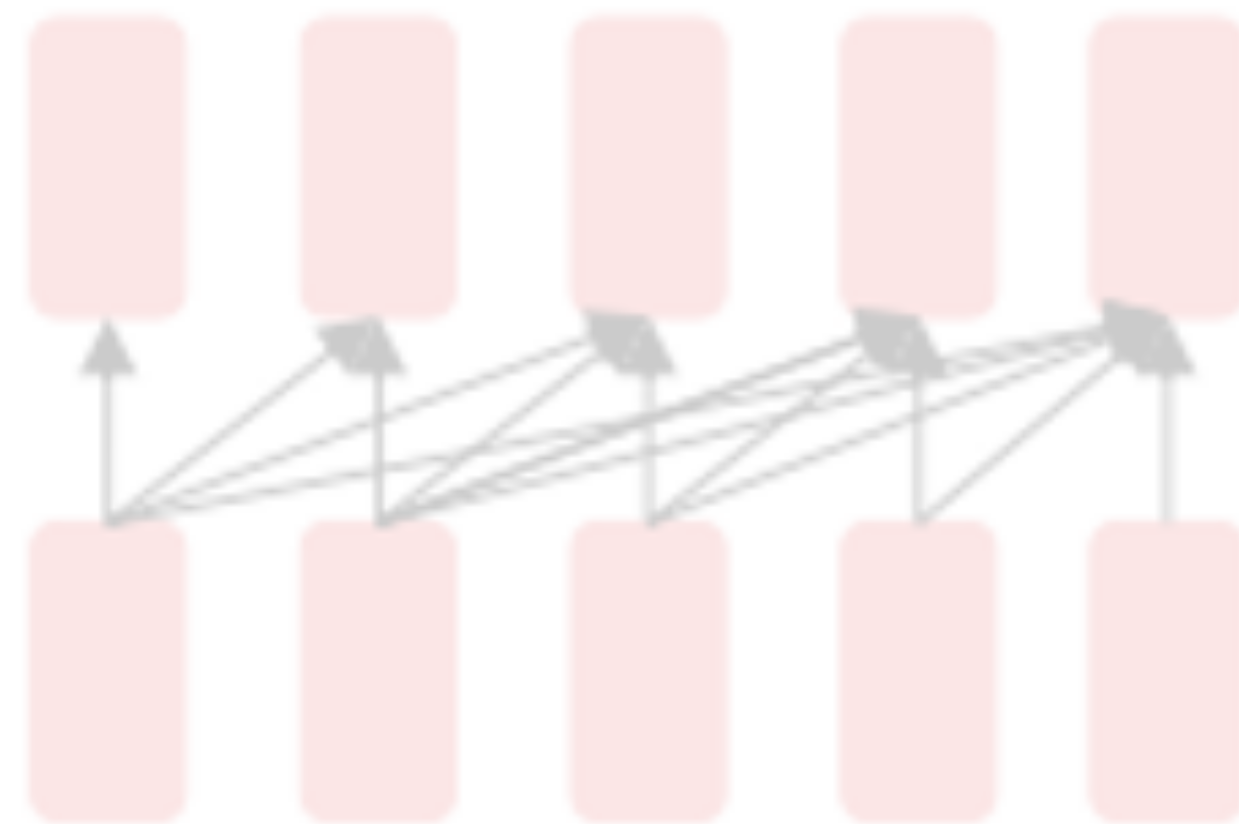
- Self-supervised Transformer based models shattered language understanding benchmarks in NLP in 2018.
- Trained on large text corpus with self-supervised objectives and then transferred.

Encoder only



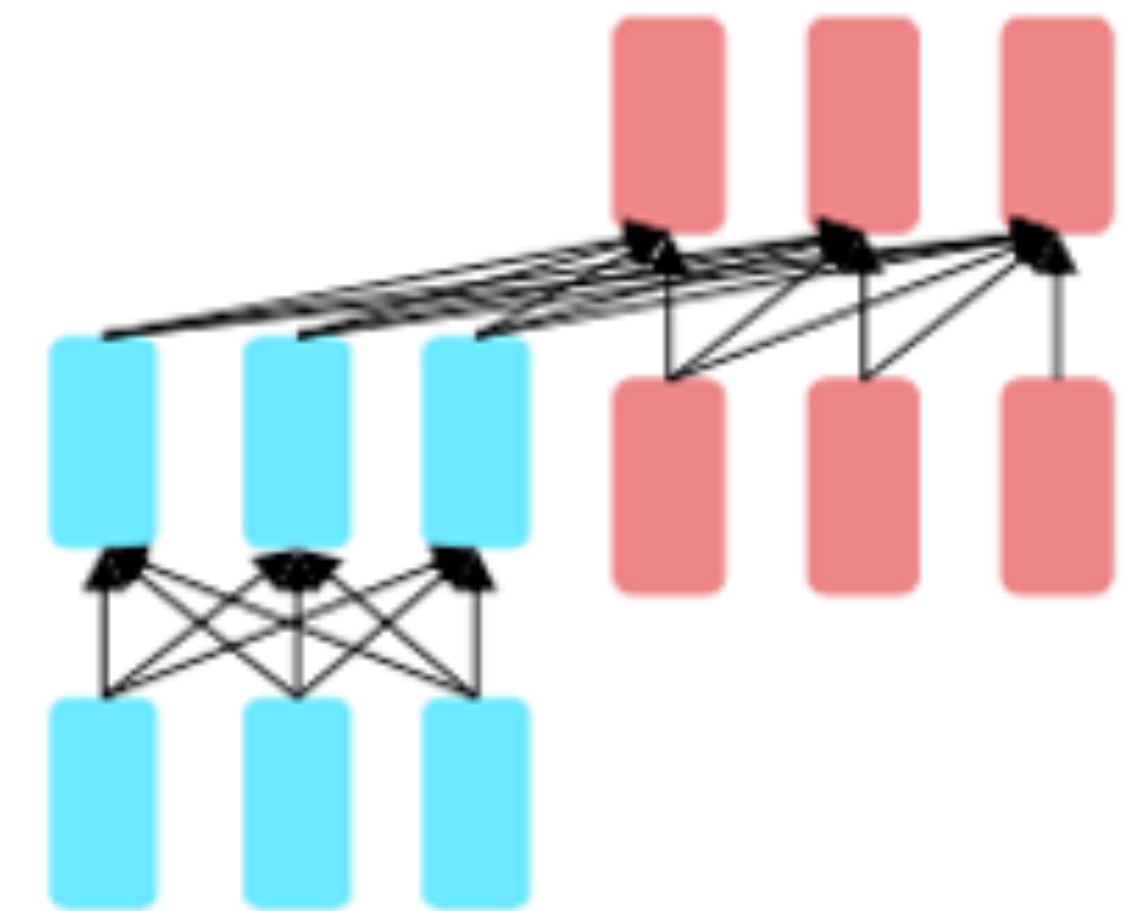
- Masked language models
- Bidirectional context
- BERT + variants (e.g. RoBERTa)
-

Decoder only



- Language models
- Can't condition on future words, good for generation
- GPT, LLaMa, PaLM

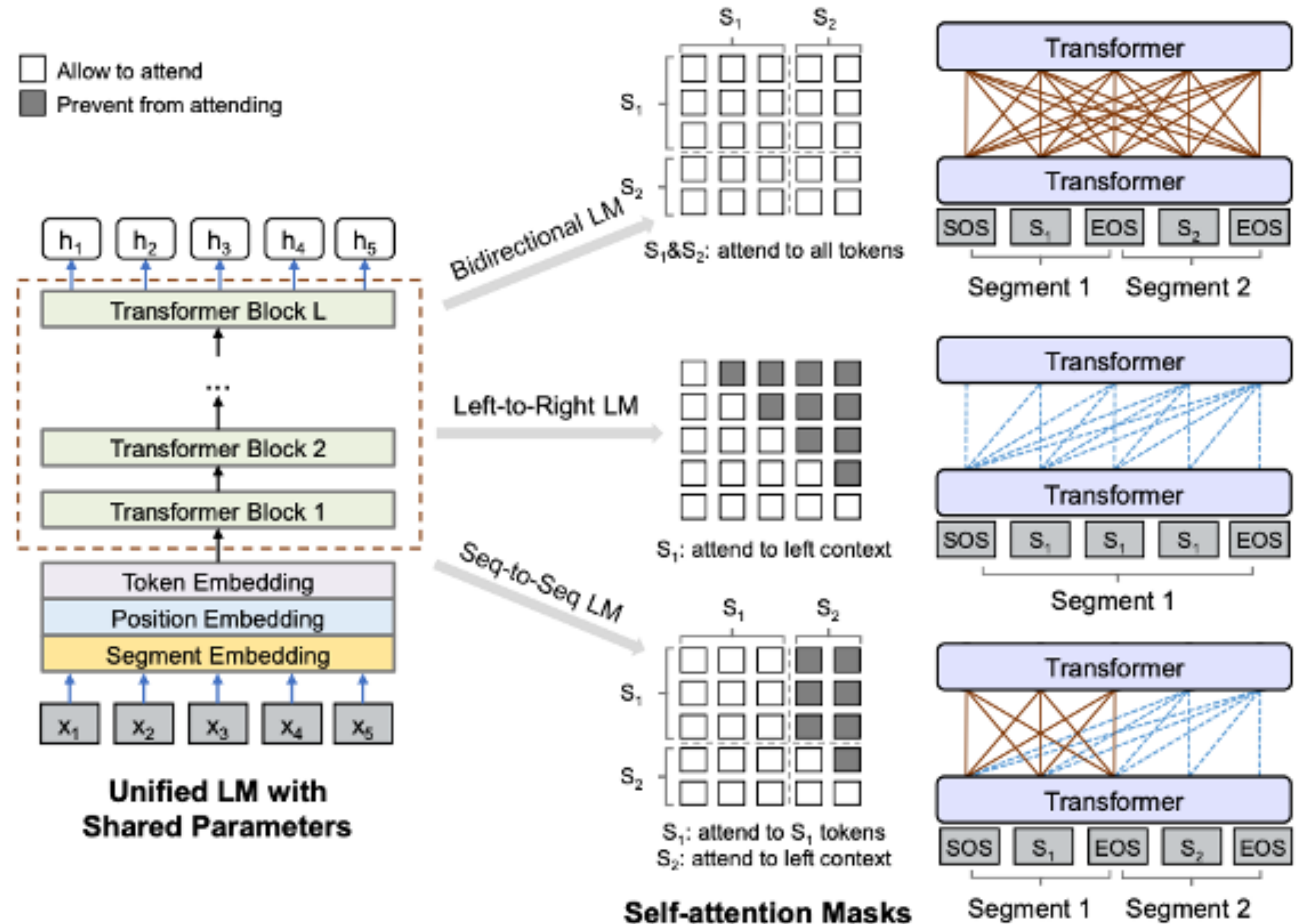
Encoder-Decoder



- Combine benefits of both
- Original Transformer, UniLM, BART, T5

Encoder-Decoder pretraining

- Combine advantages of both encoder and decoder
- Seq2Seq LM with masking
- Next sentence prediction

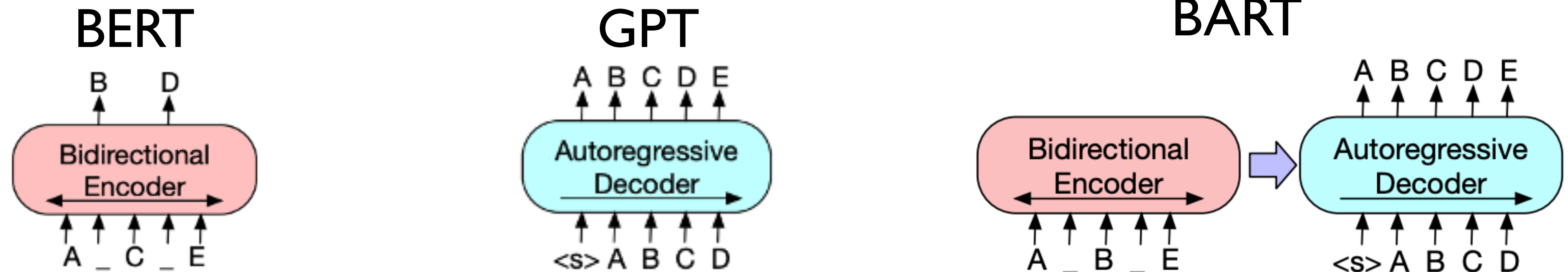


UniLM v1

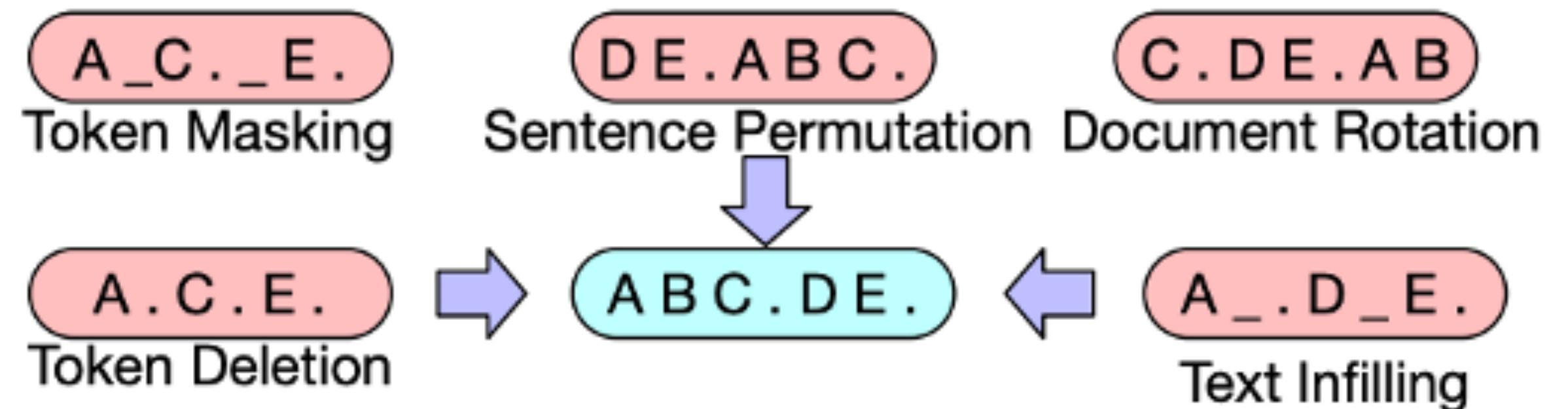
- Combine benefits of BERT (encoder) and GPT (decoder)

Model	CoLA MCC	SST-2 Acc	MRPC F1	STS-B S Corr	QQP F1	MNLI-m/mm Acc	QNLI Acc	RTE Acc	WNLI Acc	AX Acc	Score
GPT	45.4	91.3	82.3	80.0	70.3	82.1/81.4	87.4	56.0	53.4	29.8	72.8
BERT _{LARGE}	60.5	94.9	89.3	86.5	72.1	86.7/ 85.9	92.7	70.1	65.1	39.6	80.5
UNILM	61.1	94.5	90.0	87.7	71.7	87.0/85.9	92.7	70.9	65.1	38.4	80.8

BART: Denoising seq2seq training

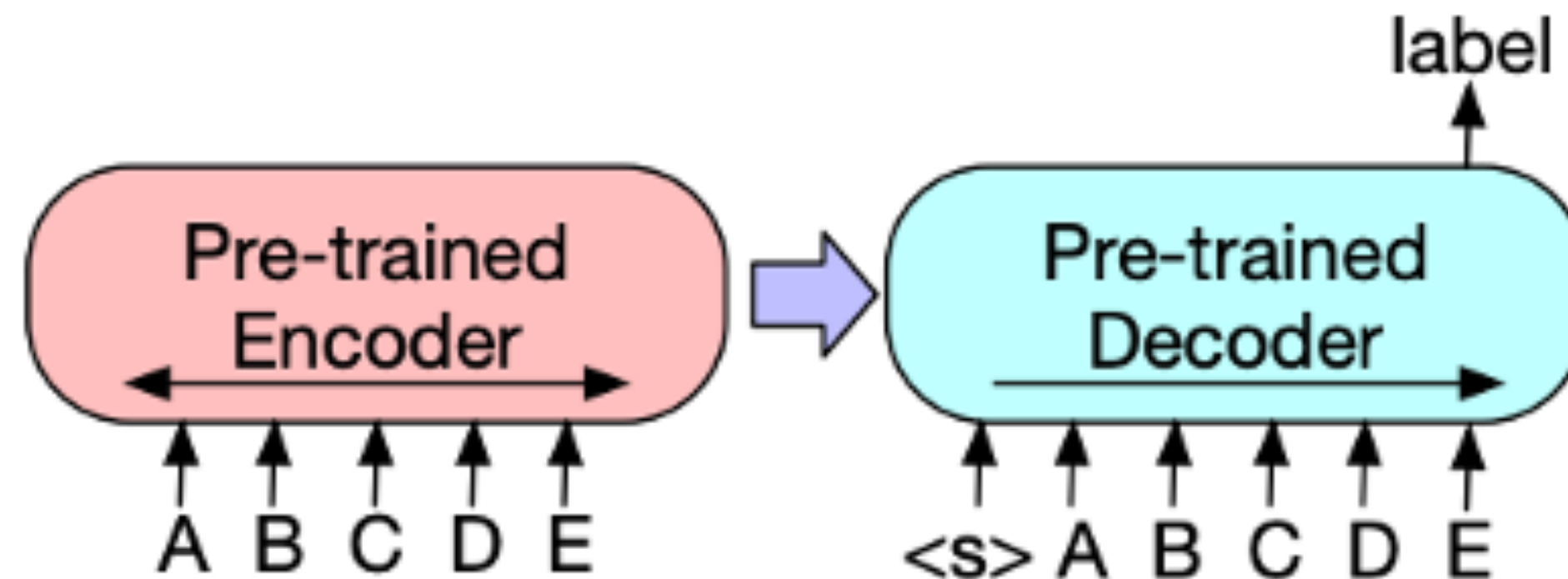


- Combine benefits of BERT (encoder) and GPT (decoder)
- More flexibility in noise generation

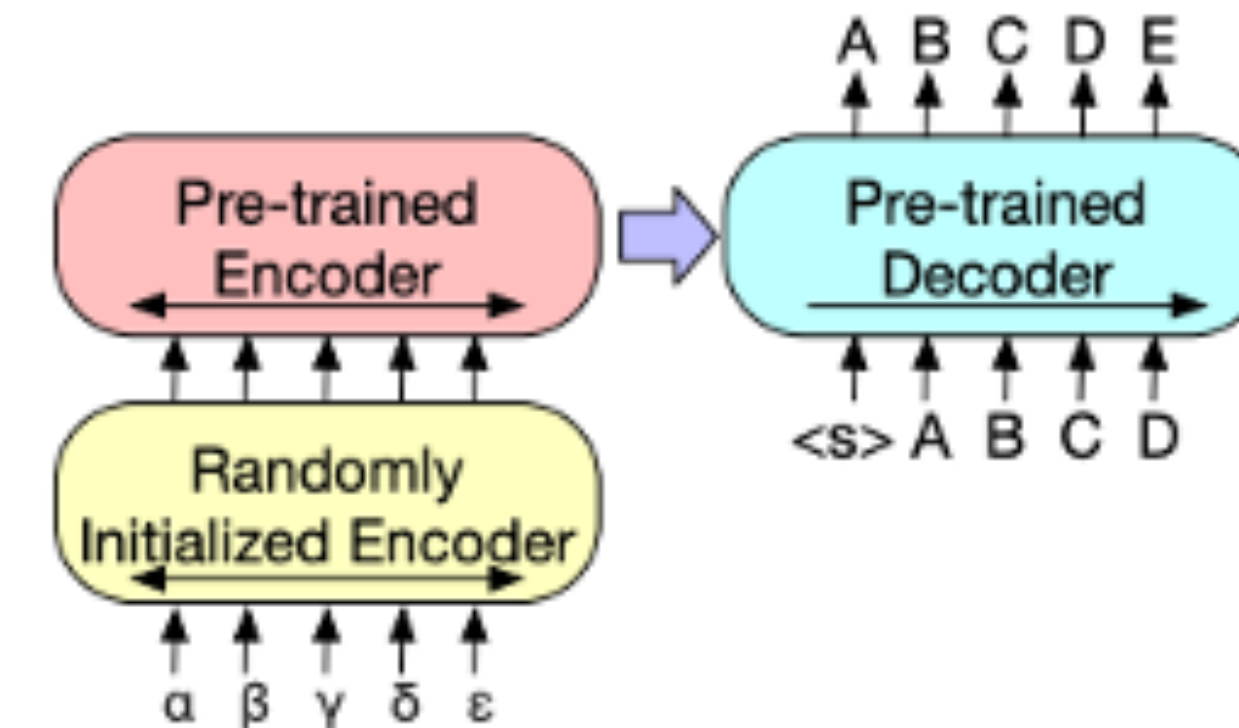


BART: Denoising seq2seq training

Classification



Machine Translation

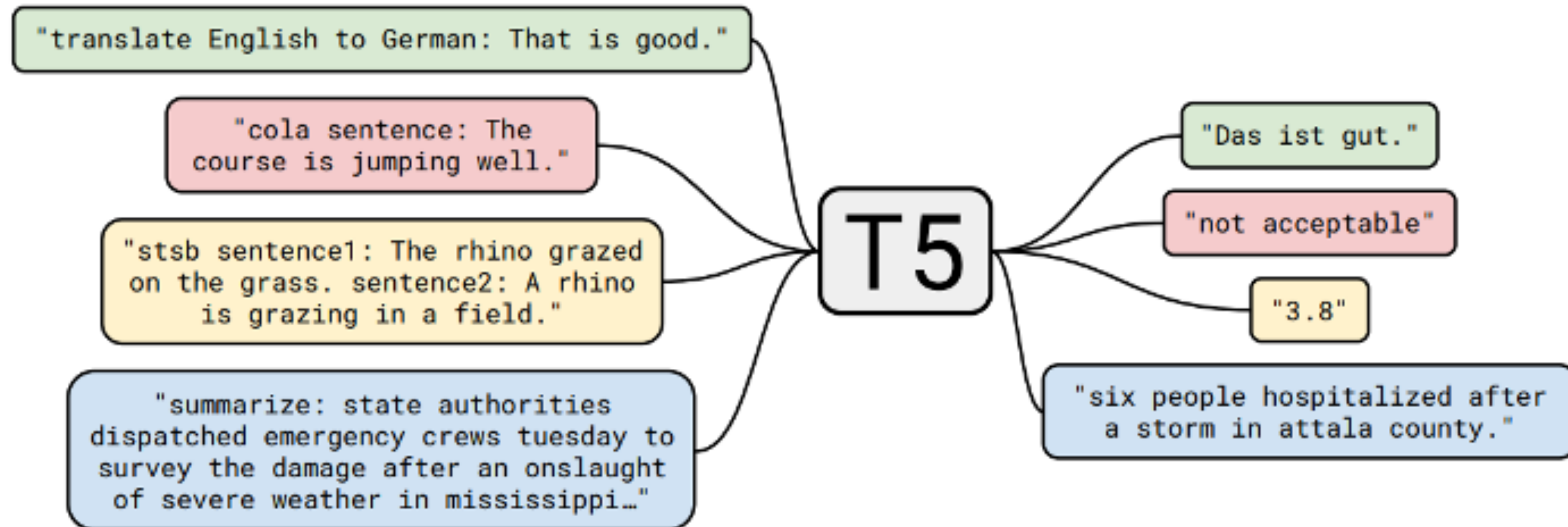


	SQuAD 1.1 EM/F1	SQuAD 2.0 EM/F1	MNLI m/mm	SST Acc	QQP Acc	QNLI Acc	STS-B Acc	RTE Acc	MRPC Acc	CoLA Mcc
BERT	84.1/90.9	79.0/81.8	86.6/-	93.2	91.3	92.3	90.0	70.4	88.0	60.6
UniLM	-/-	80.5/83.4	87.0/85.9	94.5	-	92.7	-	70.9	-	61.1
XLNet	89.0 /94.5	86.1/88.8	89.8/-	95.6	91.8	93.9	91.8	83.8	89.2	63.6
RoBERTa	88.9/ 94.6	86.5/89.4	90.2/90.2	96.4	92.2	94.7	92.4	86.6	90.9	68.0
BART	88.8/ 94.6	86.1/89.2	89.9/90.1	96.6	92.5	94.9	91.2	87.0	90.4	62.8

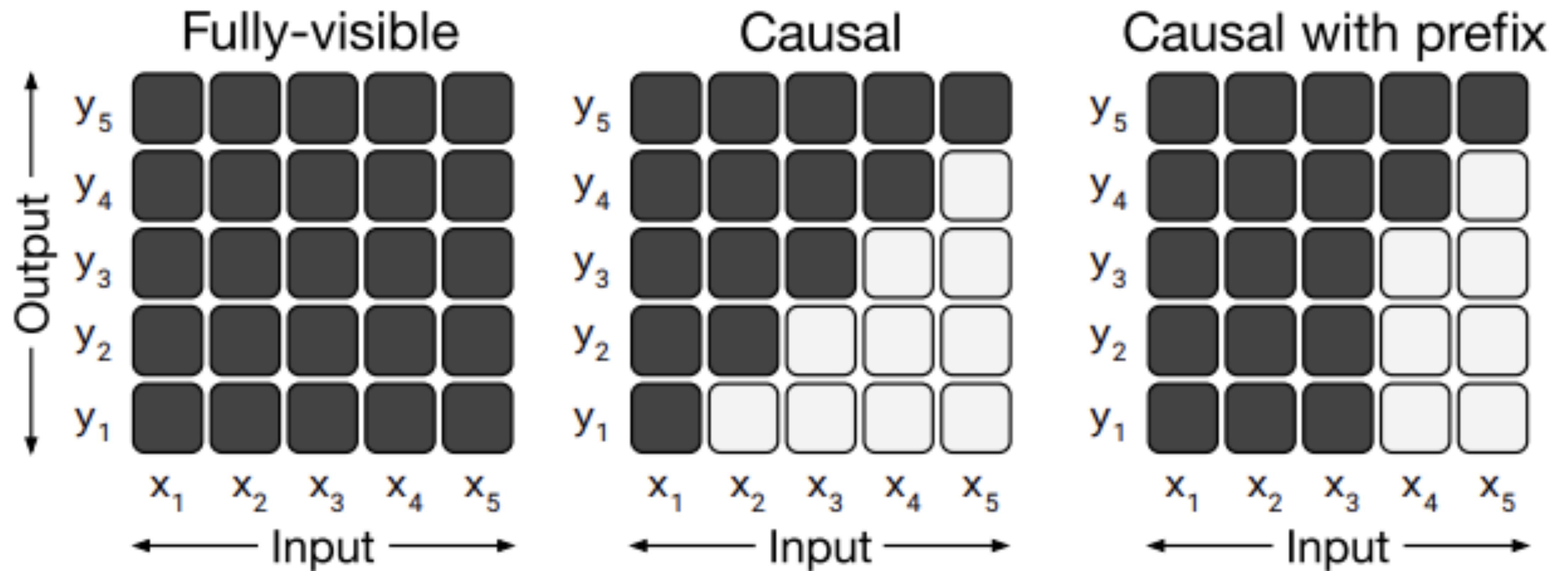
T5: Text to Text Transfer Transformer

<https://arxiv.org/abs/1910.10683>

- Treat all NLP problems as encoding text and generating text
- Trained on cleaned up version of Common Crawl

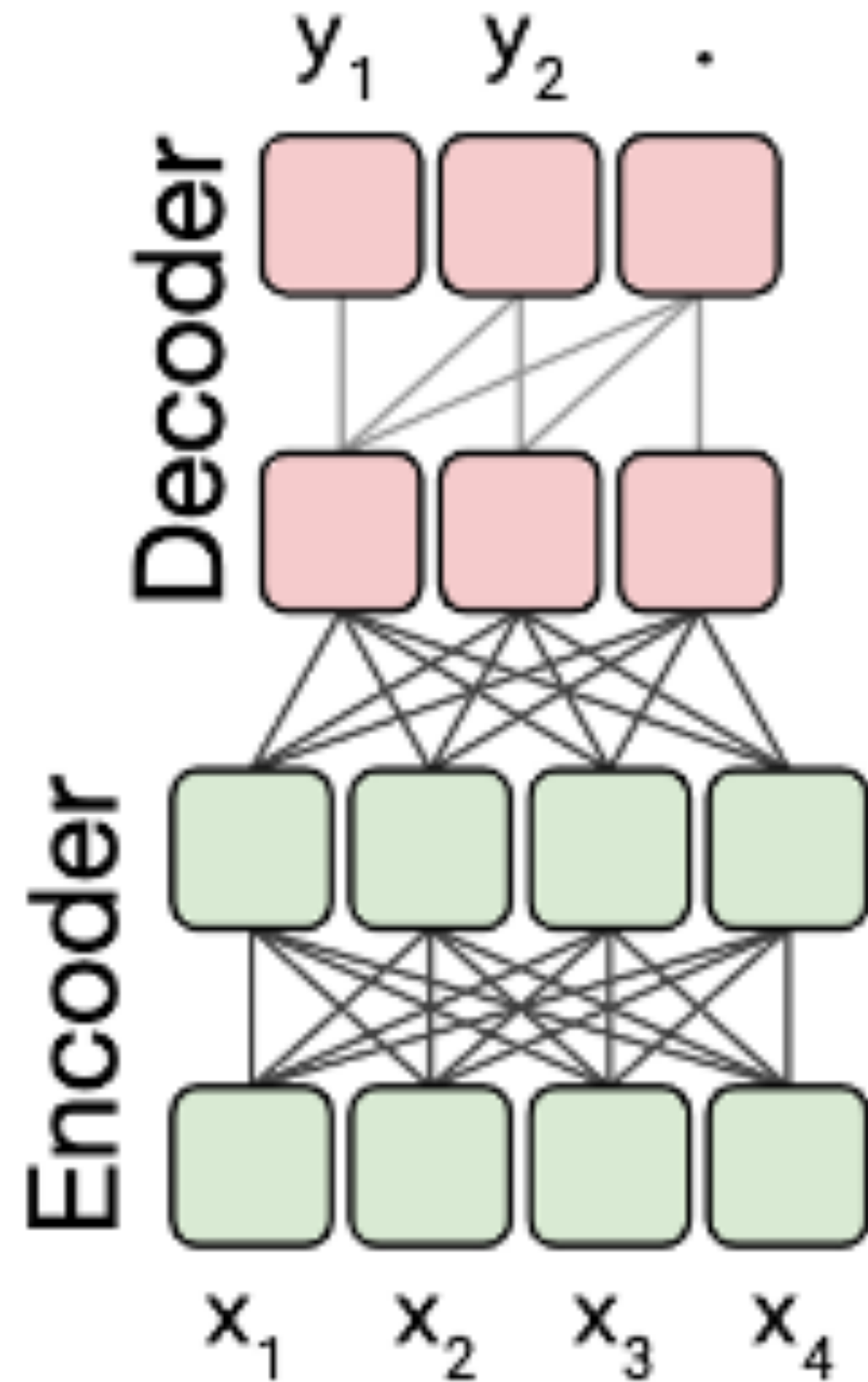


Exploring the Limits of Transfer Learning with a Unified Text-to-Text Transformer [Raffel et al, Google, JMLR 2020]



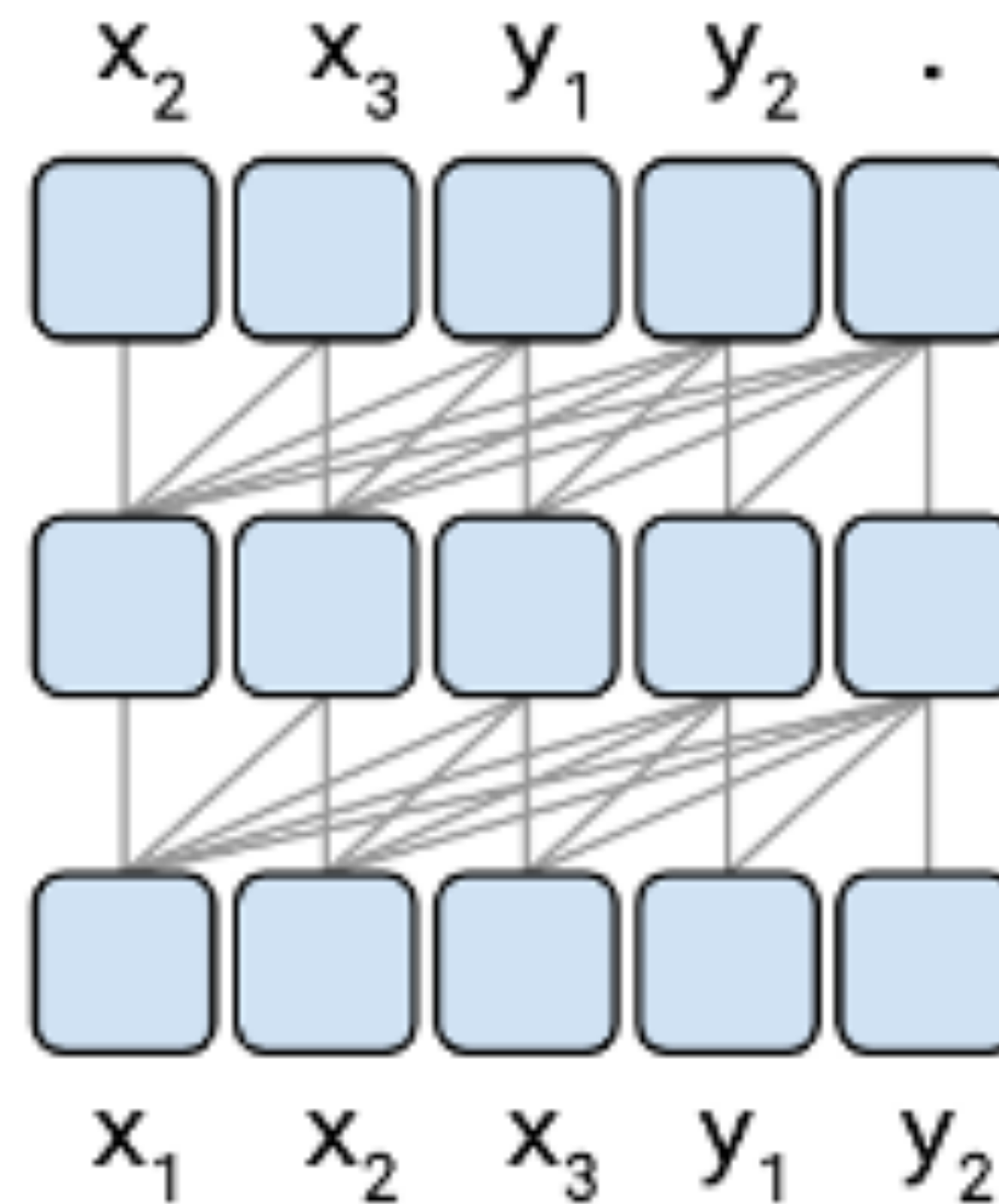
T5: Text to Text Transfer Transformer

Normally: Separate parameters for encoder/decoder



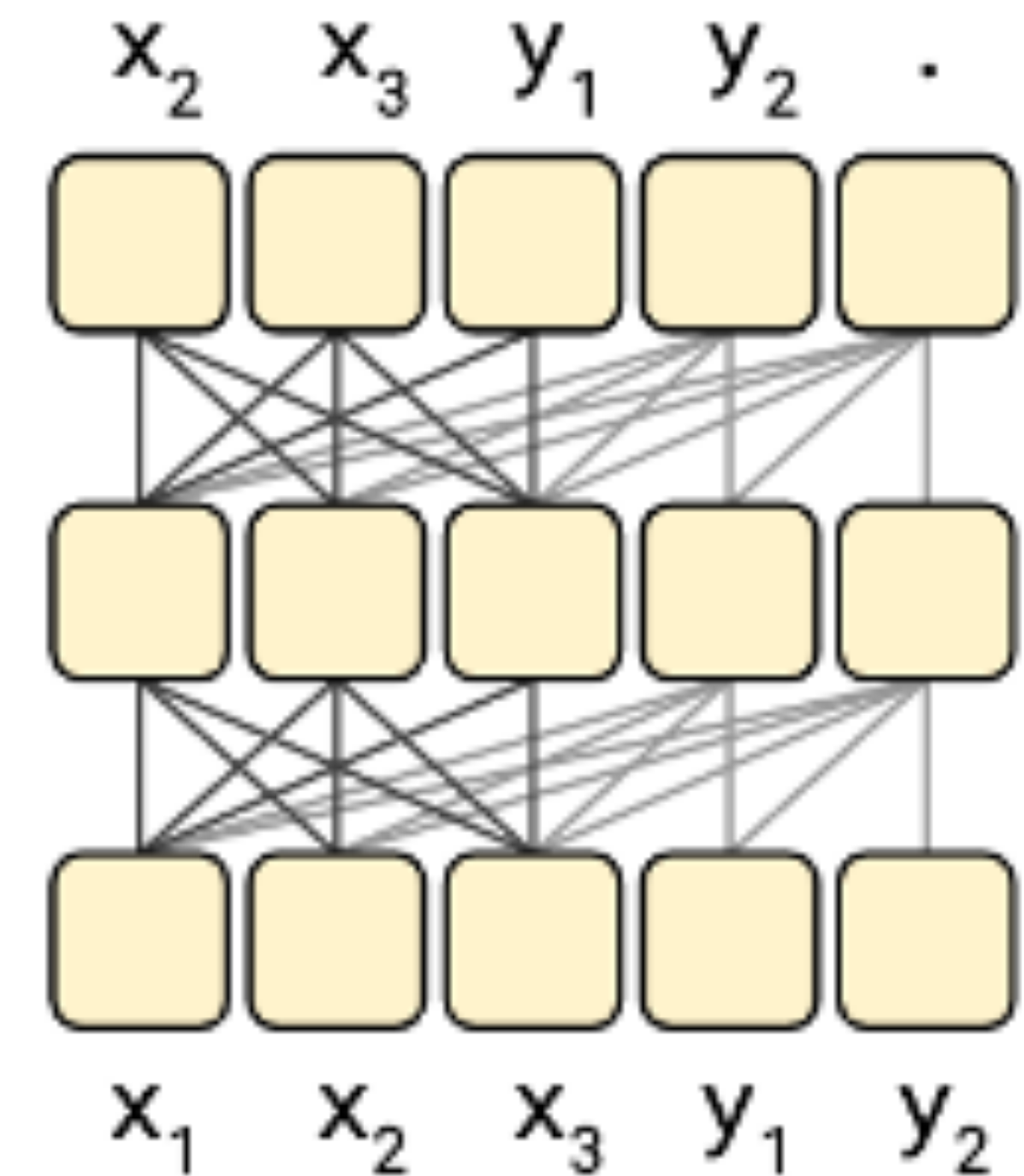
Causal masking only

Language model



Masking similar to encoder/decoder

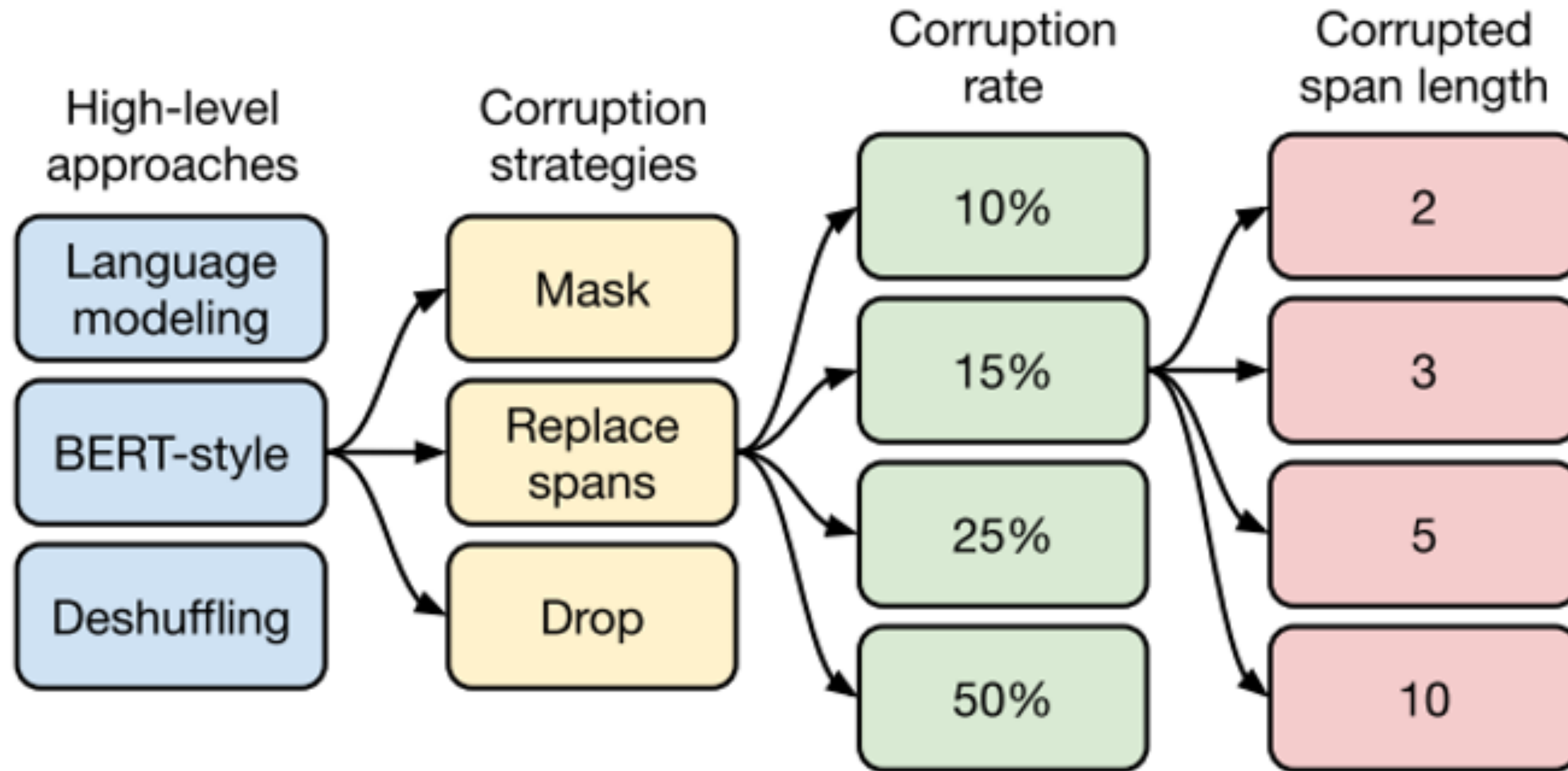
Prefix LM



Can force sharing of parameters for encoder/decoder Similar performance, less parameters

Exploring the Limits of Transfer Learning with a Unified Text-to-Text Transformer [Raffel et al, Google, JMLR 2020]

T5: Text to Text Transfer Transformer



Exploring the Limits of Transfer Learning with a Unified Text-to-Text Transformer [Raffel et al, Google, JMLR 2020]

T5 (use both encoder and decoder)

Span corruption works best

Replace different-length spans from the input with unique placeholders; decode out the spans that were removed!

Original text

Thank you ~~for inviting~~ me to your party ~~last~~ week.

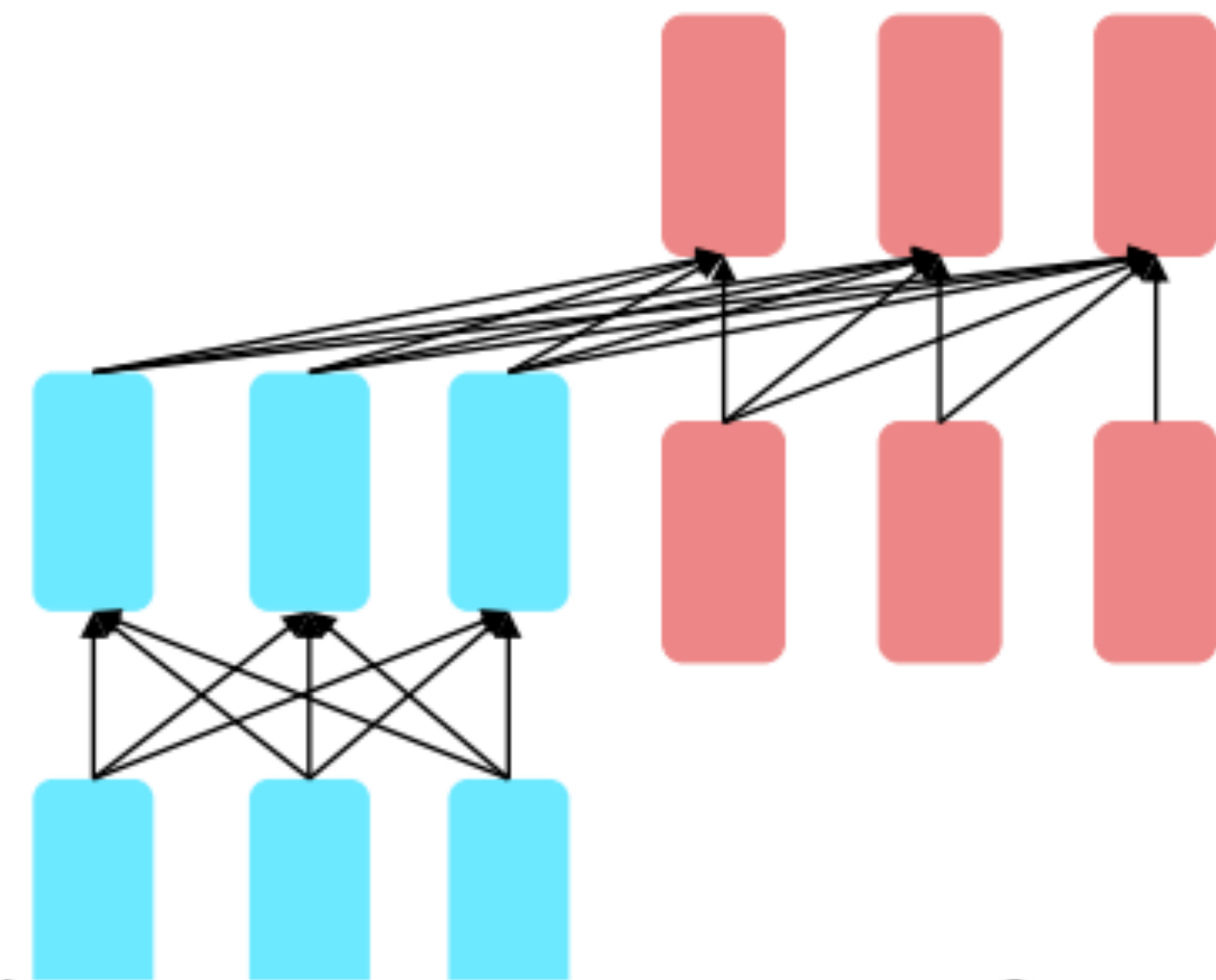
This is implemented in text preprocessing: it's still an objective that looks like **language modeling** at the decoder side.

Inputs

Thank you $\langle X \rangle$ me to your party $\langle Y \rangle$ week.

Targets

$\langle X \rangle$ for inviting $\langle Y \rangle$ last $\langle Z \rangle$



T5: Text to Text Transfer Transformer

Different corruption type

	Objective	GLUE	CNNDM	SQuAD	SGLUE	EnDe	EnFr	EnRo
Predict all	BERT-style (Devlin et al., 2018)	82.96	19.17	80.65	69.85	26.78	40.03	27.41
	MASS-style (Song et al., 2019)	82.32	19.16	80.10	69.28	26.79	39.89	27.55
Predict corrupted	★ Replace corrupted spans	83.28	19.24	80.88	71.36	26.98	39.82	27.65
	Drop corrupted tokens	84.44	19.31	80.52	68.67	27.07	39.76	27.82

Different corruption rate

Corruption rate	GLUE	CNNDM	SQuAD	SGLUE	EnDe	EnFr	EnRo
10%	82.82	19.00	80.38	69.55	26.87	39.28	27.44
★ 15%	83.28	19.24	80.88	71.36	26.98	39.82	27.65
25%	83.00	19.54	80.96	70.48	27.04	39.83	27.47
50%	81.27	19.32	79.80	70.33	27.01	39.90	27.49

Exploring the Limits of Transfer Learning with a Unified Text-to-Text Transformer [Raffel et al, Google, JMLR 2020]

T5 (use both encoder and decoder)

[Raffel et al., 2018](#) found encoder-decoders to work better than decoders for their tasks, and span corruption (denoising) to work better than language modeling.

Architecture	Objective	Params	Cost	GLUE	CNNDM	SQuAD	SGLUE	EnDe	EnFr	EnRo
★ Encoder-decoder	Denoising	$2P$	M	83.28	19.24	80.88	71.36	26.98	39.82	27.65
Enc-dec, shared	Denoising	P	M	82.81	18.78	80.63	70.73	26.72	39.03	27.46
Enc-dec, 6 layers	Denoising	P	$M/2$	80.88	18.97	77.59	68.42	26.38	38.40	26.95
Language model	Denoising	P	M	74.70	17.93	61.14	55.02	25.09	35.28	25.86
Prefix LM	Denoising	P	M	81.82	18.61	78.94	68.11	26.43	37.98	27.39
Encoder-decoder	LM	$2P$	M	79.56	18.59	76.02	64.29	26.27	39.17	26.86
Enc-dec, shared	LM	P	M	79.60	18.13	76.35	63.50	26.62	39.17	27.05
Enc-dec, 6 layers	LM	P	$M/2$	78.67	18.26	75.32	64.06	26.13	38.42	26.89
Language model	LM	P	M	73.78	17.54	53.81	56.51	25.23	34.31	25.38
Prefix LM	LM	P	M	79.68	17.84	76.87	64.86	26.28	37.51	26.76

T5 summary

Raffel+ 2019

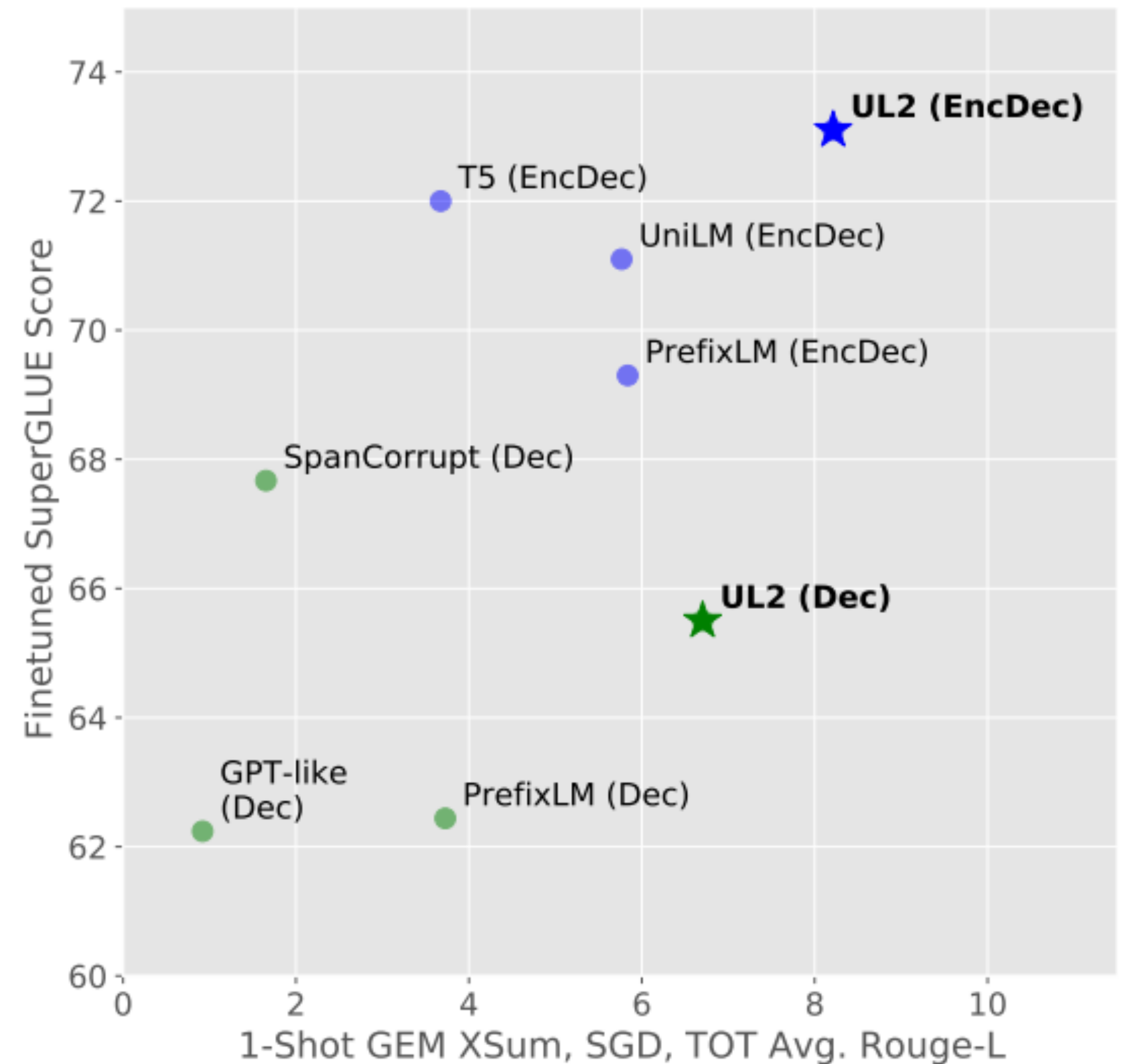
<https://arxiv.org/abs/1910.10683>

- Ablation study on many aspects of pre-training and fine-tuning
 - Model size (bigger is better; 11B parameters)
 - Amount of training data (more is better)
 - Domain / cleanliness of training data [-ve]
 - Pre-training objective (e.g. span length of masked text) [-ve]
 - Ensemble models [-ve]
 - Fine-tuning recipe (e.g. only allow top k layers to fine-tune) [-ve]
 - Multi-task training [-ve]

Pre-training objectives

Pre-training objectives

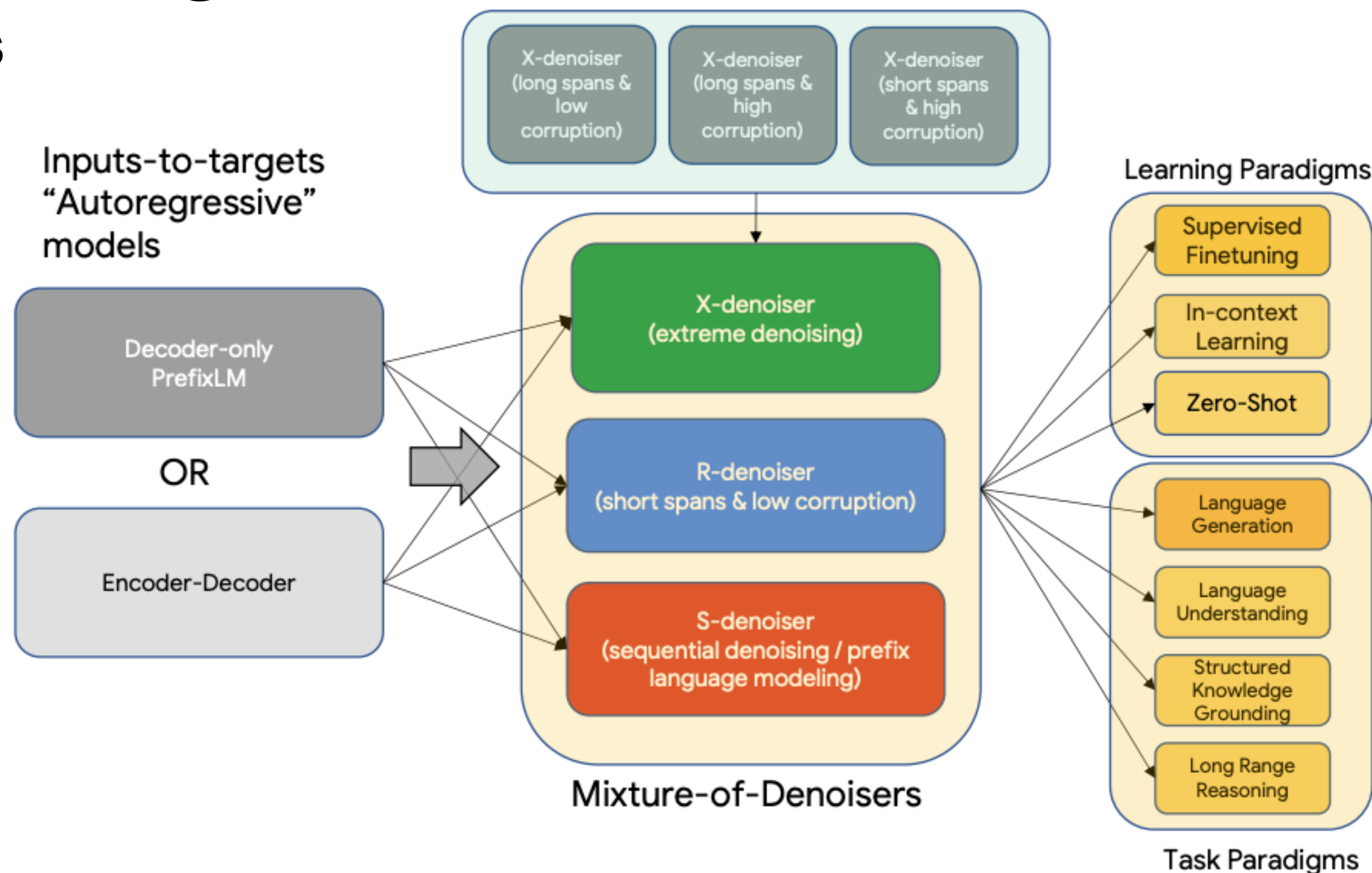
- Pre-training objectives can be decoupled from architectures
- Can train using corruption using Decoder only or Encoder-Decoder objective



Unified pre-training

Mixture-of-denoisers

- Train to perform **next token prediction** with different noise schemes
- $\text{SpanCorrupt}(\mu, r, n)$:
 - μ mean span length,
 - r corruption rate,
 - n number of corrupted spans
- Architecture agnostic



Unified pre-training

Denoisers

- **Mode switching** with paradigm token: [R], [S], [X] to tell the model which model it is in

Regular Denoising
Short spans (2-5 tokens), Low corruption (15% masking)

R-Denoising

Inputs:

[R]	He dealt in archetypes before anyone knew such things existed, and his	3	to take an emotion or a situation	5	it to the limit helped create a cadre of plays that have been endlessly	4	– and copied. Apart from this, Romeo and Juliet inspired Malorie Blackman's Noughts	5	there are references to Hamlet in Lunar Park by Bret Easton Ellis	2	The Tempest was the cue for The Magus by John Fowles.
-----	--	---	-----------------------------------	---	---	---	---	---	---	---	---

Target:

	3	<S>	5	<S>	4	<S>	5
<S>	2	<E>					

S-Denoising

Inputs:

[S]	He dealt in archetypes before anyone knew such things existed, and his ability to take an emotion or a situation and push it to the limit helped create a cadre of plays that have been endlessly staged – and copied. Apart from this, Romeo and Juliet	
-----	--	--

Target:

	
	95
	<E>

X-Denoising

Inputs:

[X]	He dealt in archetypes be	16			
	things existed, and his ability to take an emotion or a situation	32			
	plays that have been endlessly staged – and copied. Apart from	24	Malorie Blackman's Noughts & Crosses, there are references to Hamlet in Lunar	24	
	Tempest was the cue for The Magus by John Fowles.				

Target:

	16	<S>	
	32	<S>	
	24	<S>	
	24	<S>	<E>

Inputs:

[X]	He dealt in archetypes	3	anyone knew such things existed, a	3	ability to take an	5					
	situation and push it to the limit helped	4	cadre of plays	4	been endlessly staged – and	5					
	Apart from this, Romeo and Juliet inspired Malorie Blackman's	5	Crosses,	3	are references to Hamlet in	3	Park by Bret Easton	2	and	4	
	4	was the	2	for The	4	by John	5				

Target:

	3	<S>	3	<S>	5	<S>	4	<S>	
4	<S>	5	<S>	5	<S>	3	<S>		
3	<S>	2	<S>	4	<S>	4	<S>	2	<S>
4	<S>	5	<E>						

Grayed out areas are masked out
Next token prediction

