

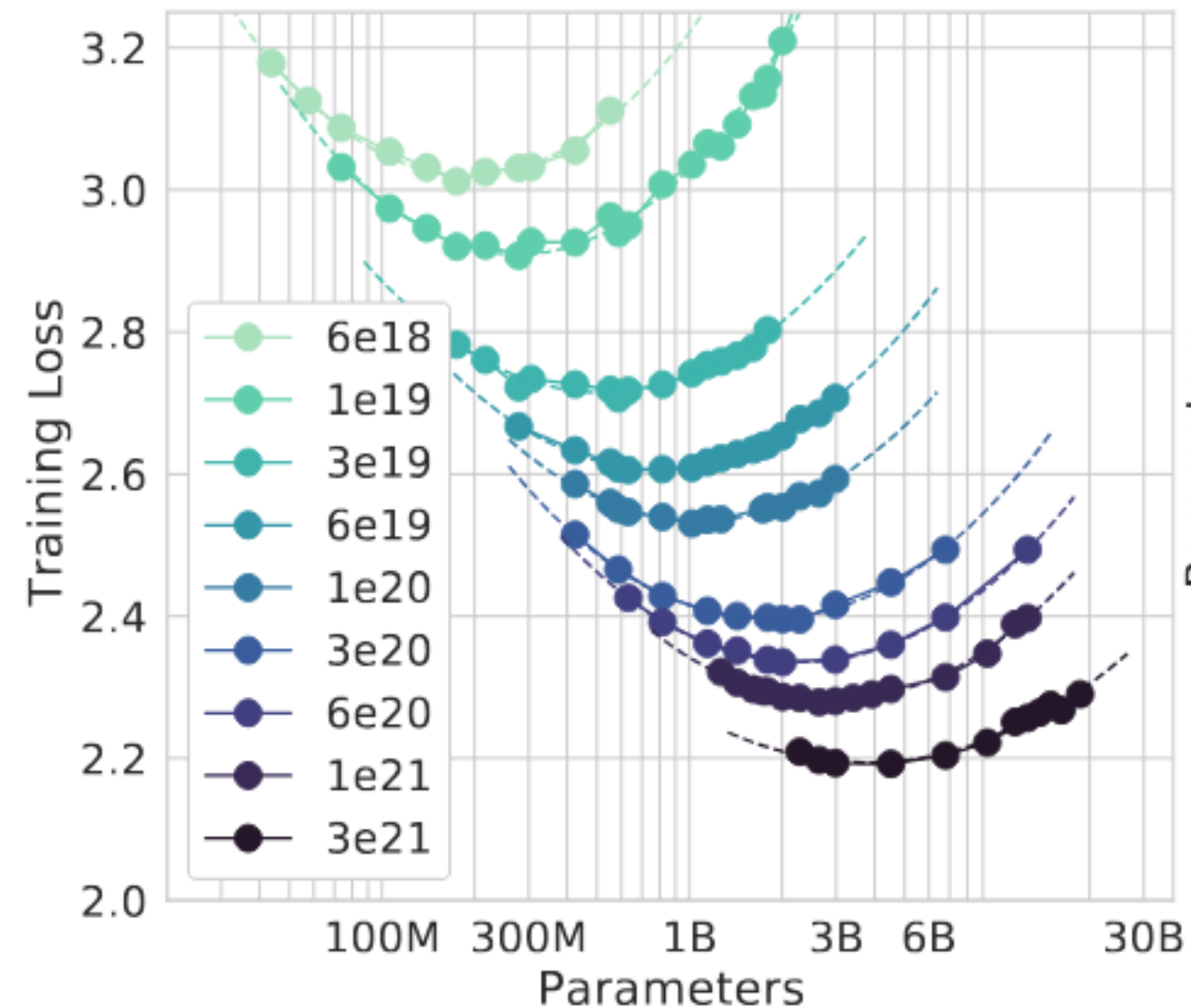
Causal Attention with Lookahead Keys

Zhuoqing Song, Peng Sun, Huizhuo Yuan, Quanquan Gu

October XX, 2025

Motivation

- Scaling law runs up against a practical bottleneck — high quality data
- How to improve models' performance under limited token budget?

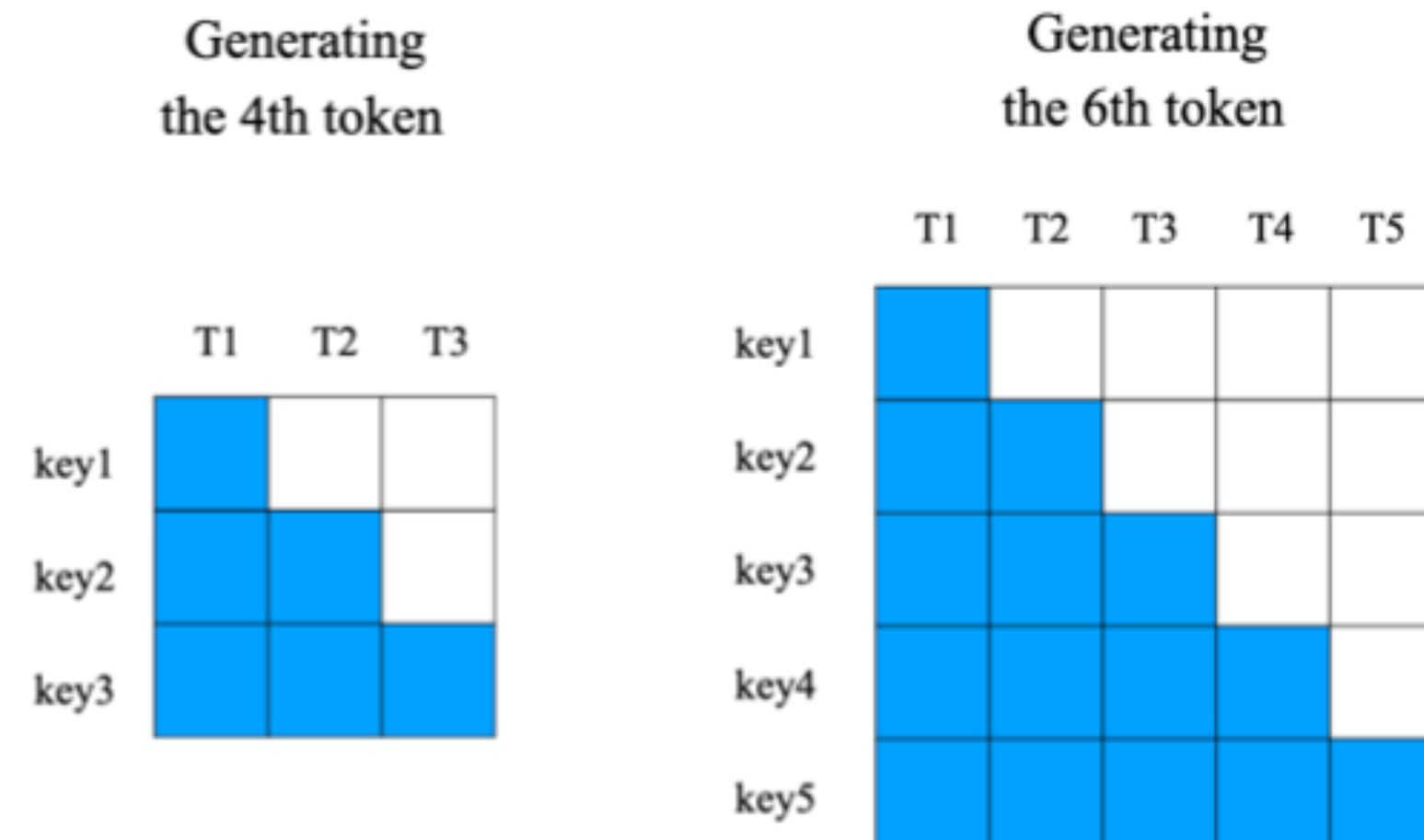


Shortcomings of Standard Causal Attention

- Standard causal attention in recurrent form

$$\text{causal-attention}(\mathbf{X}^t) = \text{softmax} \left(\frac{\mathbf{q}_t \mathbf{K}_t^\top}{\sqrt{d}} \right) \mathbf{V}_t = \frac{\sum_{s=1}^t \exp(\mathbf{q}_t \mathbf{k}_s^\top / \sqrt{d}) \mathbf{v}_s}{\sum_{s=1}^t \exp(\mathbf{q}_t \mathbf{k}_s^\top / \sqrt{d})} \in \mathbb{R}^{1 \times d}.$$

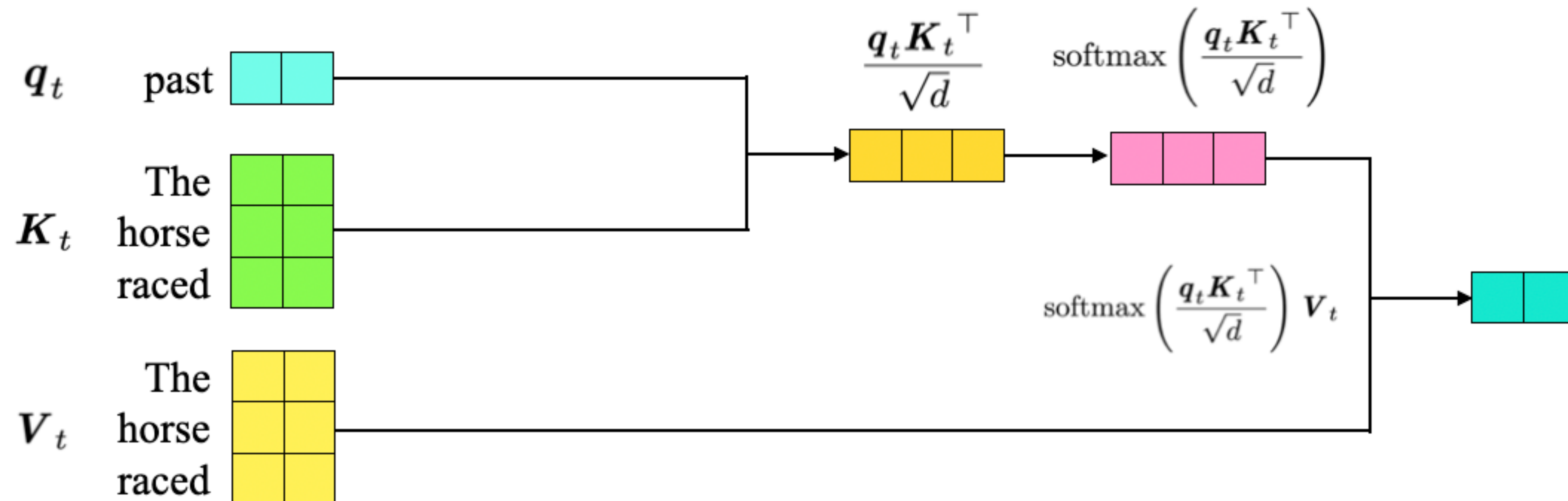
- Receptive fields of keys in standard causal attention



Standard Causal Attention in Recurrent Form

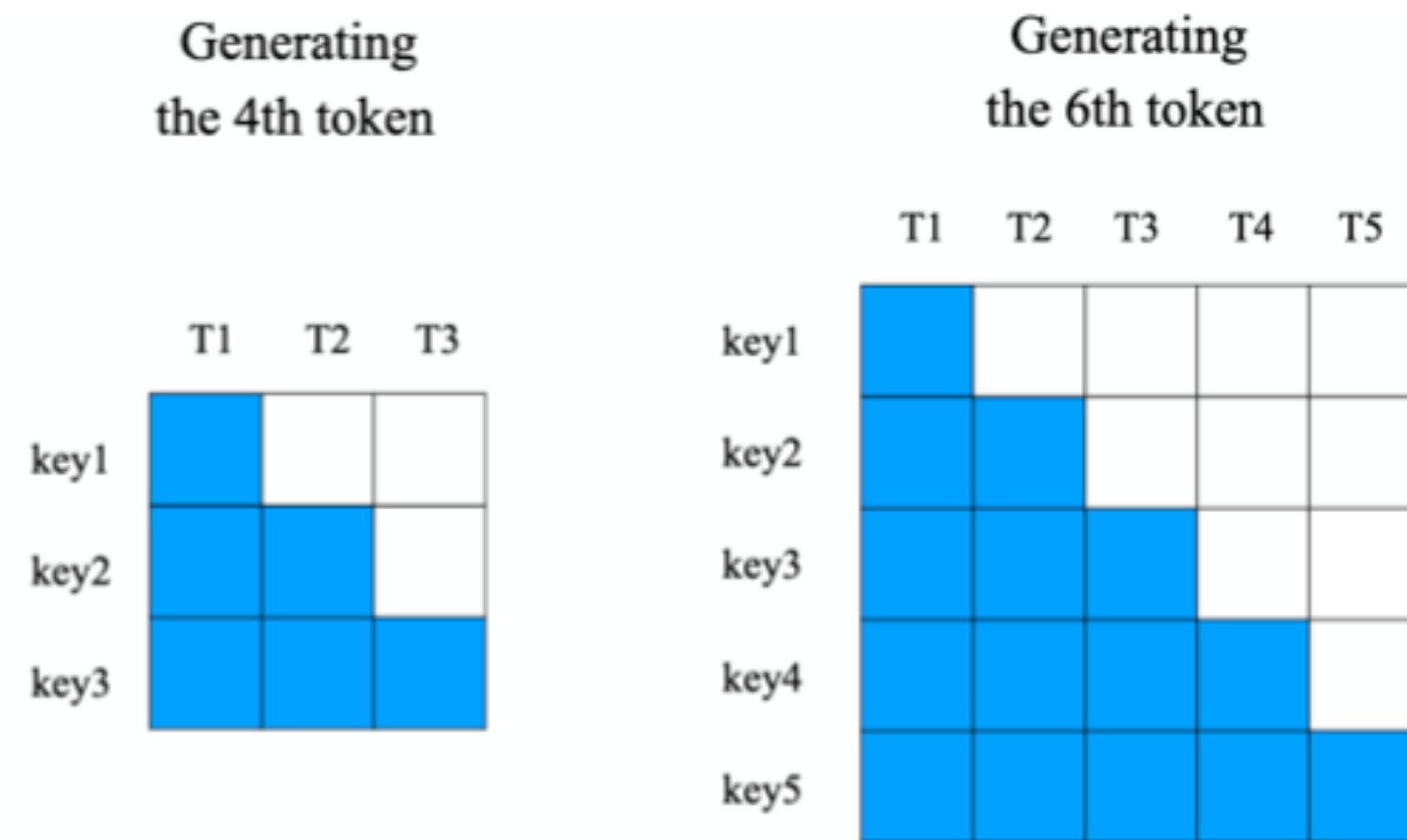
- Example: “The horse raced past the barn fell.”

Input Output



Shortcomings of Standard Causal Attention

- Receptive fields of keys in standard causal attention



- Causal mask blocks each token's access to its future information
 - Hurt natural language understanding (BERT vs. GPT)
 - Noise from precedent context
 - (Indirect evidence) diffusion LLMs are using bi-directional attention

Examples: Shortcomings of Causal Masking

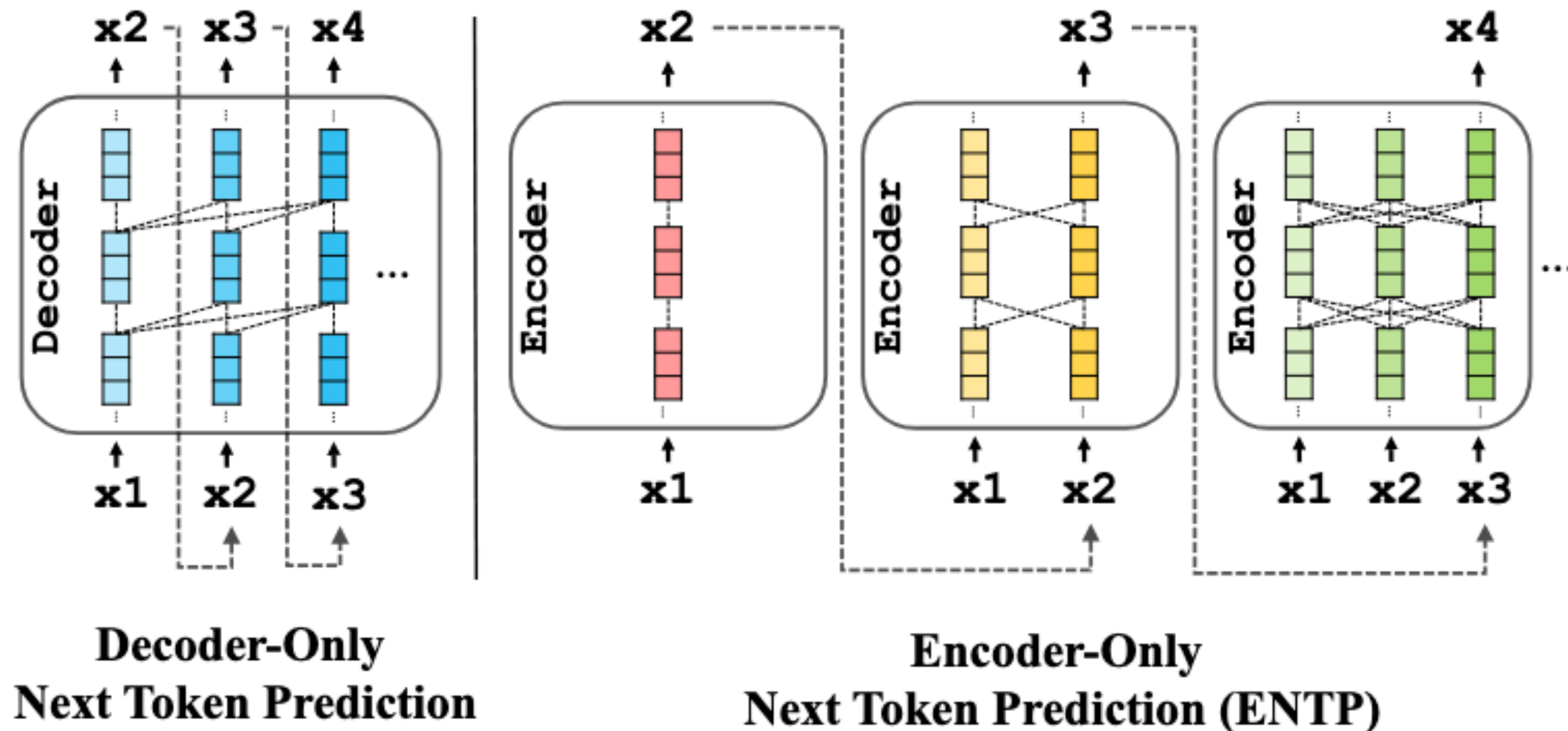
- Garden-path sentence
 - The old man the boat.
 - The complex houses married and single soldiers and their families.
- Question/focus at the end of inputs
- Variable assignment problem
 - $x = 1; y = 2; x = 3; y = 5; x = ?$

Related work: Sentence Embedding

- Backward Dependency Enhanced LLM
- Echo Embeddings
- Re-Reading

Related work: Encoder-only Next Token Prediction (ENTP)

- NTP by encoder-only Transformers (on KV cache, cubic training complexity)



Related work: Selective Attention

- Modify attention matrix, save KV cache

$$\text{SelectiveAttention}(Q, K, V) = \text{softmax}(\frac{QK^T}{\sqrt{d_k}} - F)V$$

$$F = \text{Accumulate}(\text{Constrain}(S))$$



Related work: PaTH Attention

- Data-dependent positional encodings

$$\mathbf{A}_{ij} \propto \exp\left(\mathbf{k}_j^\top \left(\prod_{s=j+1}^i \mathbf{H}_s\right) \mathbf{q}_i\right),$$

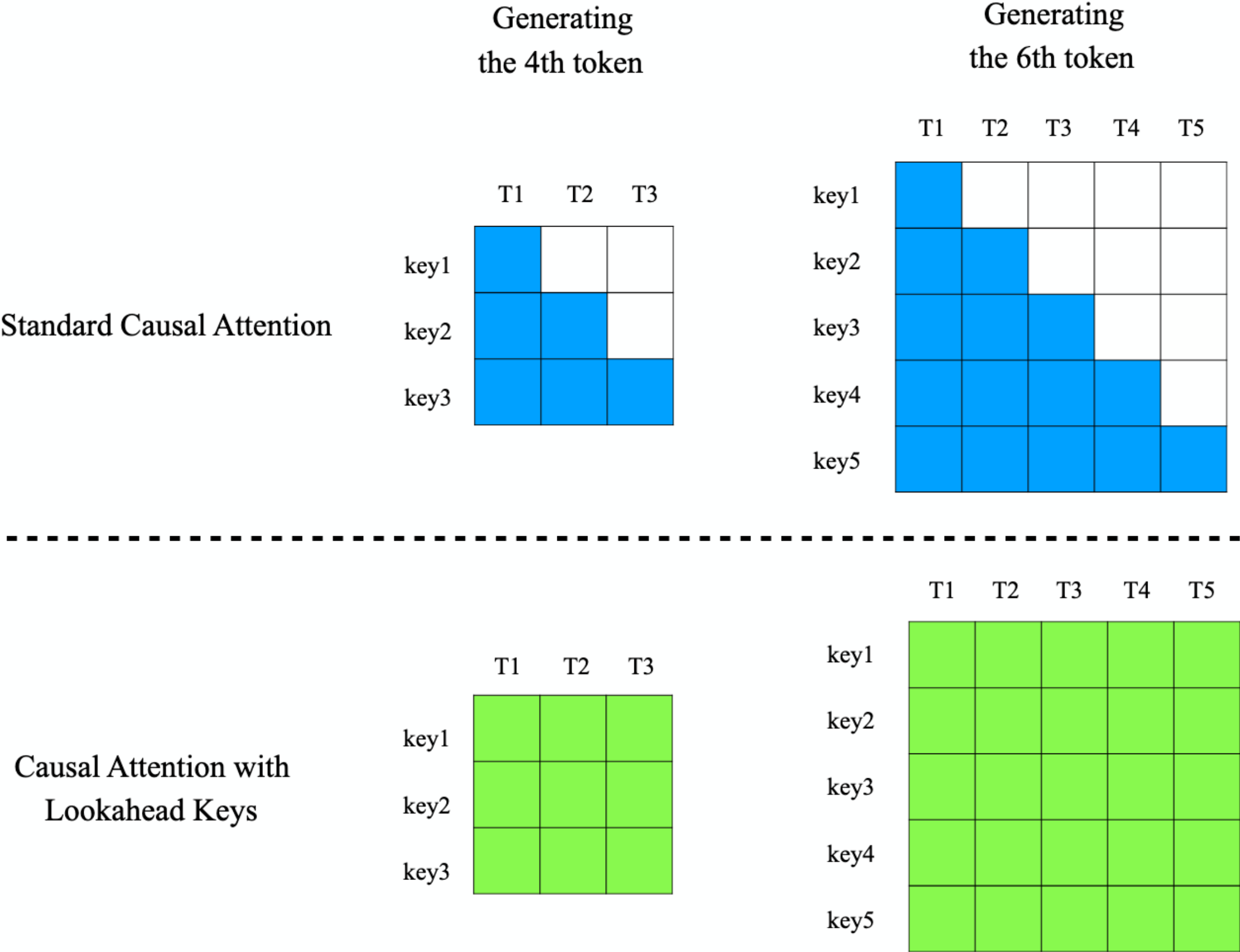
$$\mathbf{H}_t = \mathbf{I} - \beta_t \mathbf{w}_t \mathbf{w}_t^\top,$$

- Inference of PaTH

$$\mathbf{k}_i^{(t)} \leftarrow (\mathbf{I} - \beta_t \mathbf{w}_t \mathbf{w}_t^\top) \mathbf{k}_i^{(t-1)} \quad \text{for all } i < t,$$

Our Approach: Expanding Receptive Fields

- Should we expand receptive fields of Q, K or V?



Our Approach: Expanding Receptive Fields

- Should we expand receptive fields of Q, K or V?

The same reason with using KV cache instead of Q cache!

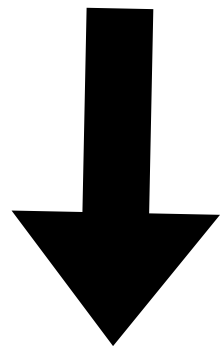
- Why are we updating keys rather than values?

Updating keys has an equivalent form which enables efficient parallel training.

If attentions scores are accurate ``enough'', information mix in values is unnecessary.

Our Approach: Expanding Receptive Fields

- Expanding receptive fields of keys
- Maintaining efficient training and inference



- Update keys during inference
- Without materializing keys at each position during training

CASTLE in Recurrent Form (Overview)

CASTLE

- Causal-key attention score

$$\mathbf{s}_t^C = \frac{\mathbf{q}_t^C \mathbf{K}_t^{C\top}}{\sqrt{d}} \in \mathbb{R}^{1 \times t}$$

- Lookahead-key attention score

$$\mathbf{s}_t^U = \frac{\mathbf{q}_t^C \mathbf{U}^{t\top}}{\sqrt{d}} \in \mathbb{R}^{1 \times t}$$

- Attention weights

$$\mathbf{p}_t = \text{softmax}(\mathbf{s}_t^C - \text{SiLU}(\mathbf{s}_t^U)) \in \mathbb{R}^{1 \times t}$$

- Outputs

$$\text{attention}(\mathbf{X}^t) = \mathbf{p}_t \mathbf{V}_t^C \in \mathbb{R}^{1 \times d}$$

Standard Causal Attention

- Causal-key attention score

$$\mathbf{s}_t = \frac{\mathbf{q}_t \mathbf{K}_t^\top}{\sqrt{d}} \in \mathbb{R}^{1 \times t}$$

- Attention weights

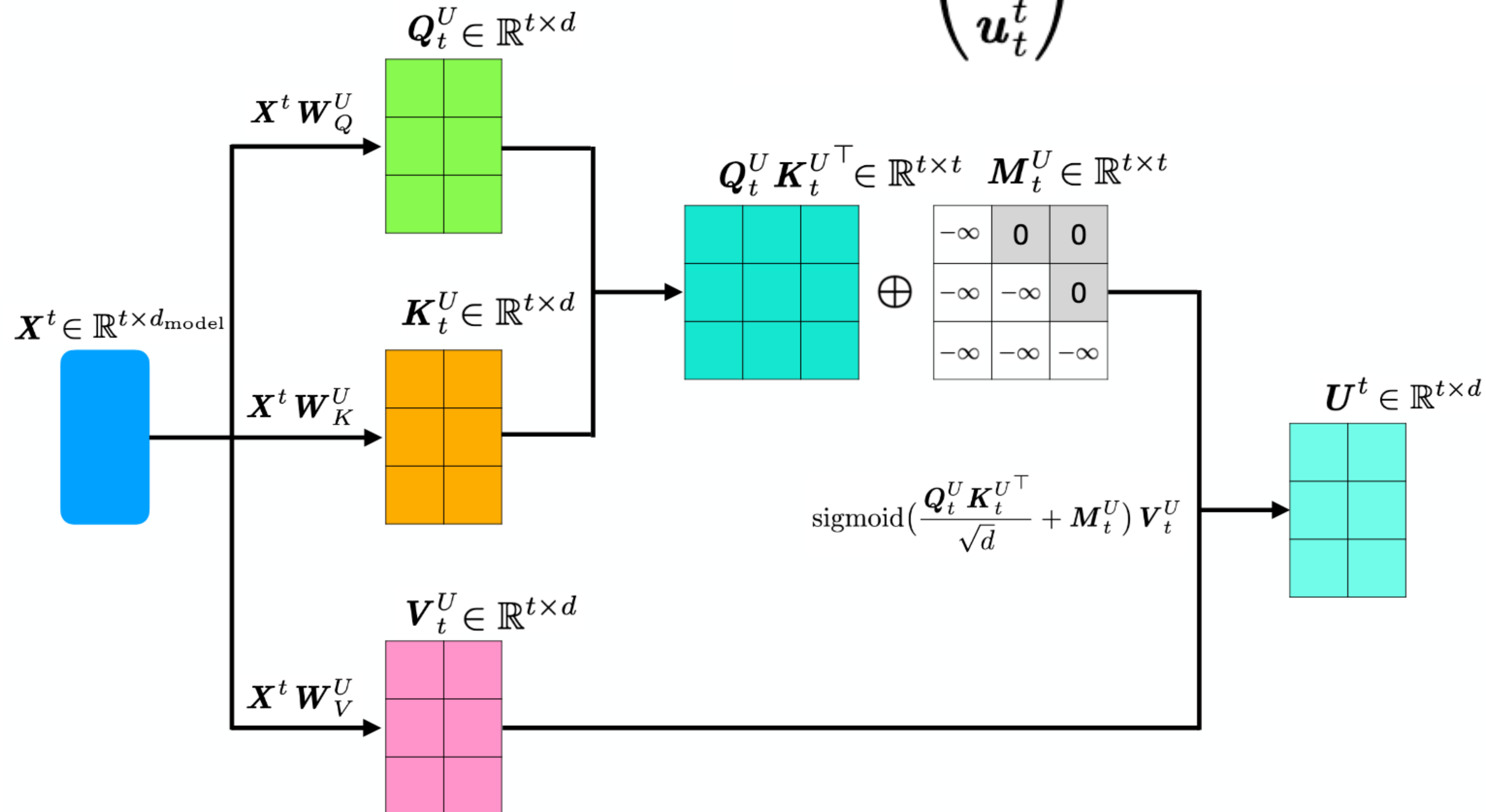
$$\mathbf{p}_t = \text{softmax}(\mathbf{s}_t) \in \mathbb{R}^{1 \times t}$$

- Outputs

$$\text{attention}(\mathbf{X}^t) = \mathbf{p}_t \mathbf{V}_t \in \mathbb{R}^{1 \times d}$$

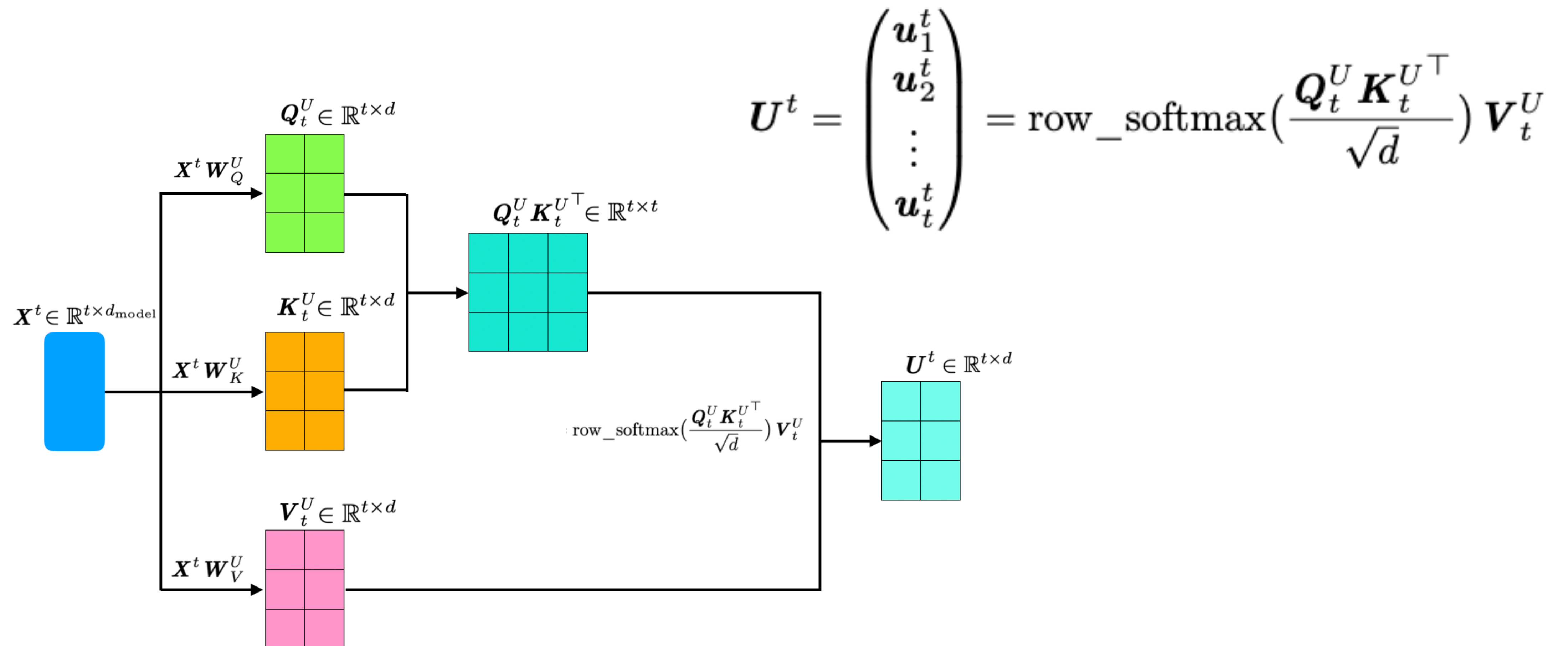
Formal Definition of Lookahead Keys

- Similar to attention mechanism $U^t = \begin{pmatrix} u_1^t \\ u_2^t \\ \vdots \\ u_t^t \end{pmatrix} = \text{sigmoid}\left(\frac{Q_t^U K_t^{U^\top}}{\sqrt{d}} + M_t^U\right) V_t^U \in \mathbb{R}^{t \times d};$



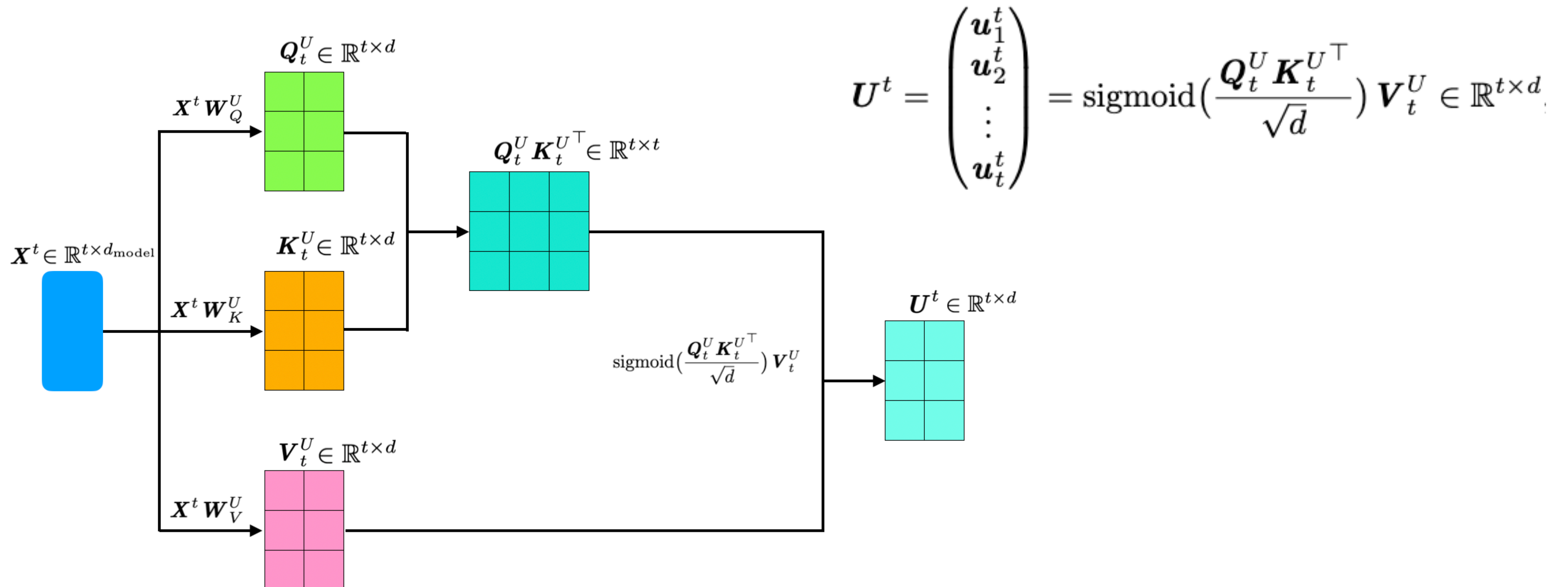
Failed Attempt of Lookahead keys

- Encode all information from token 1 to token t: numerically unstable



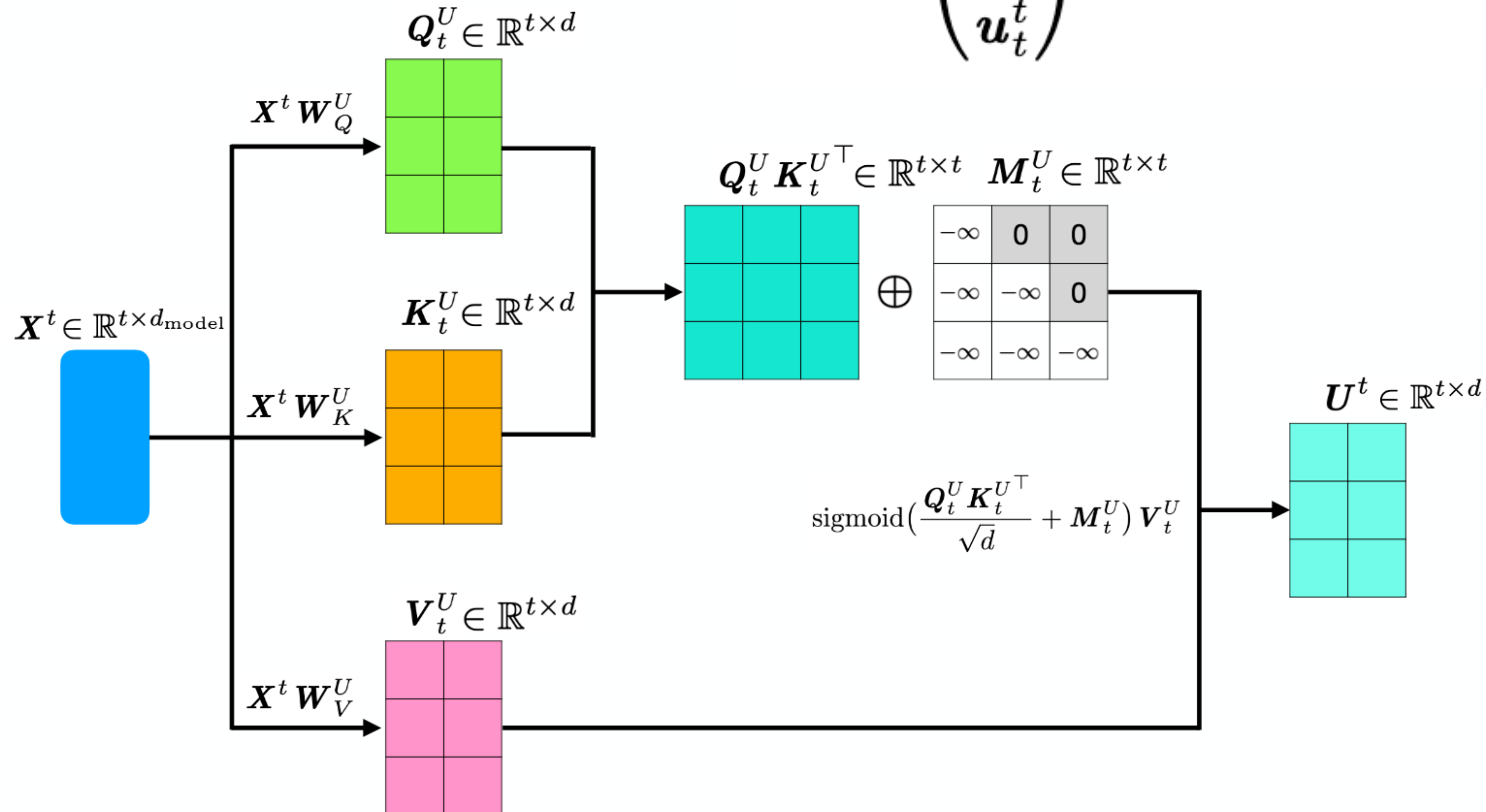
Failed Attempt of Lookahead keys

- Replace softmax by sigmoid: training instability (blow up easily)



Formal Definition of Lookahead Keys

- Similar to attention mechanism $U^t = \begin{pmatrix} u_1^t \\ u_2^t \\ \vdots \\ u_t^t \end{pmatrix} = \text{sigmoid}\left(\frac{Q_t^U K_t^{U^\top}}{\sqrt{d}} + M_t^U\right) V_t^U \in \mathbb{R}^{t \times d};$



Definition of Causal Keys

- Causal keys are the same as the keys in standard causal attention

$$\mathbf{K}_t^C = \begin{pmatrix} \mathbf{k}_1^C \\ \mathbf{k}_2^C \\ \vdots \\ \mathbf{k}_t^C \end{pmatrix} = \mathbf{X}^t \mathbf{W}_K^C \in \mathbb{R}^{t \times d}.$$

Ablation Studies on Causal Keys

- Whether we need causal keys?
 - Aligning parameters

CASTLE TYPE	n_{params}	n_{layers}	d_{model}	n_{heads}	d
CASTLE	120M	12	768	6	64
CASTLE w/o causal keys	120M	12	768	7	64

- Training & validation loss

	Train		Eval	
	Loss	PPL	Loss	PPL
CASTLE	2.913	18.417	2.920	18.541
CASTLE w/o causal keys	3.006	20.213	3.021	20.505

Comparison between Causal Keys and Lookahead Keys

	x1	x2	x3	x4	x5
causal key 1					
causal key 2					
causal key 3					
causal key 4					
causal key 5					

	T1	T2	T3	T4	T5
causal key 1					
causal key 2					
causal key 3					
causal key 4					
causal key 5					

	x1	x2	x3	x4	x5
lookahead key 1					
lookahead key 2					
lookahead key 3					
lookahead key 4					
lookahead key 5					

	T1	T2	T3	T4	T5
lookahead key 1					
lookahead key 2					
lookahead key 3					
lookahead key 4					
lookahead key 5					

CASTLE in Recurrent Form

- Causal-key attention score

$$\mathbf{s}_t^C = \frac{\mathbf{q}_t^C \mathbf{K}_t^{C\top}}{\sqrt{d}} \in \mathbb{R}^{1 \times t}$$

- Lookahead-key attention score

$$\mathbf{s}_t^U = \frac{\mathbf{q}_t^C \mathbf{U}^{t\top}}{\sqrt{d}} \in \mathbb{R}^{1 \times t},$$

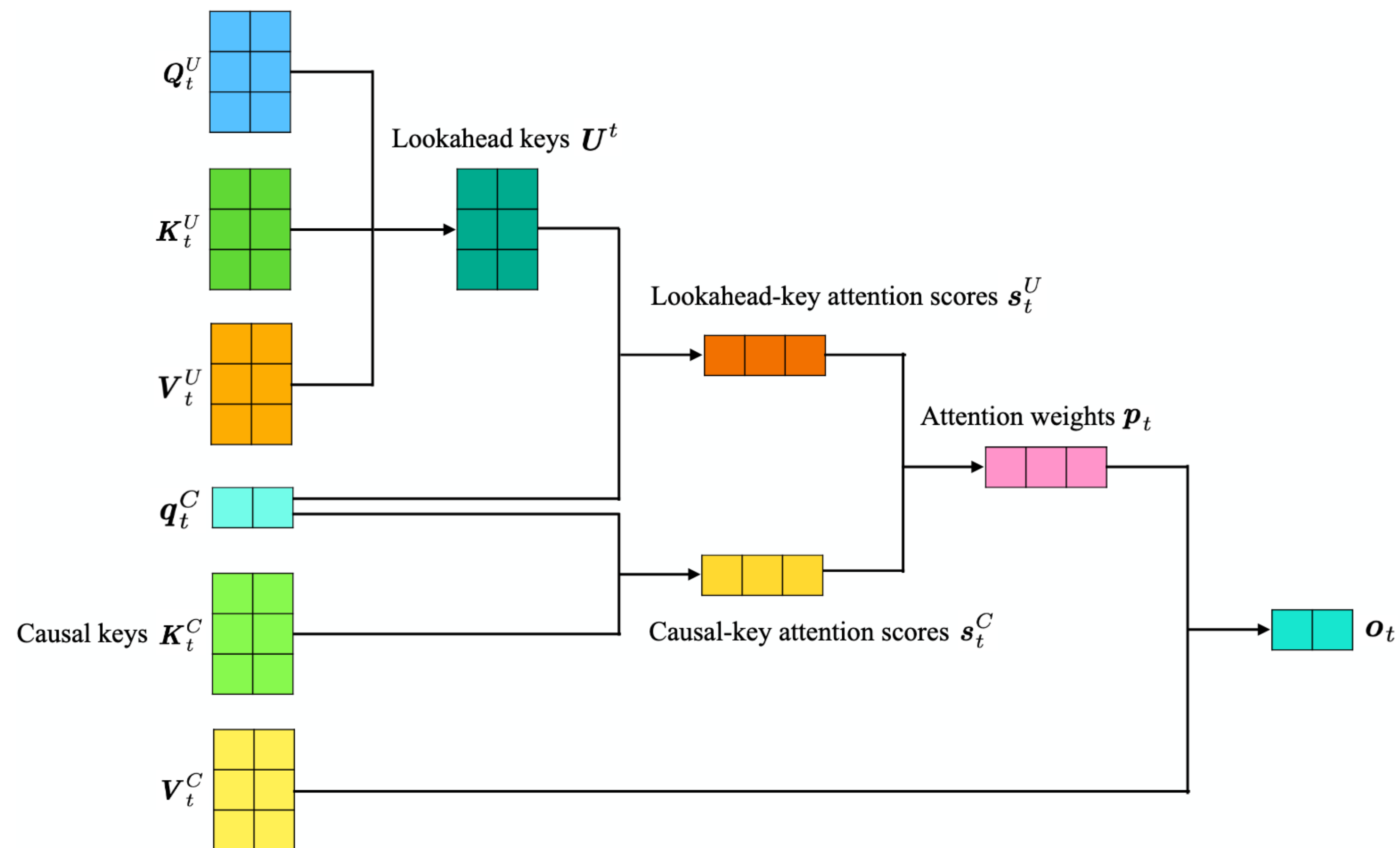
- Attention weights

$$\mathbf{p}_t = \text{softmax}(\mathbf{s}_t^C - \text{SiLU}(\mathbf{s}_t^U)) \in \mathbb{R}^{1 \times t},$$

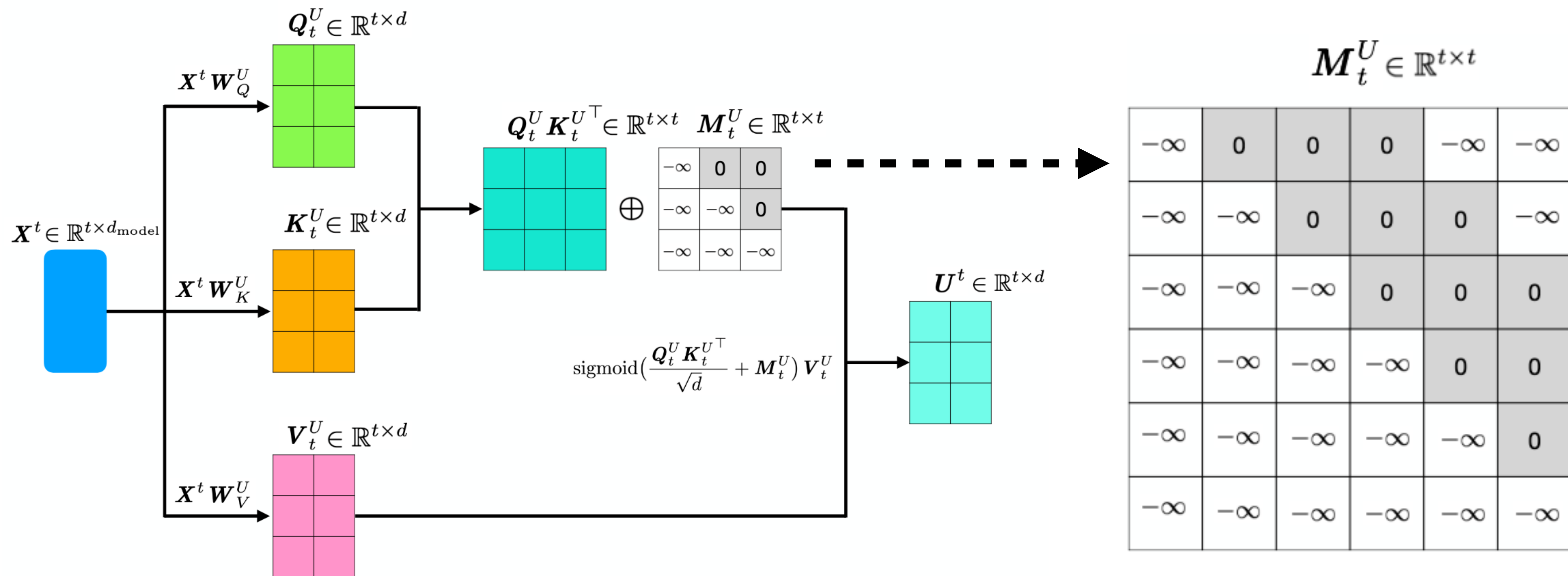
- outputs

$$\text{attention}(\mathbf{X}^t) = \mathbf{p}_t \mathbf{V}_t^C \in \mathbb{R}^{1 \times d}.$$

Illustration for CASTLE in Recurrent Form



A variant of CASTLE: CASTLE-SWL



Ablation Studies on Sliding Window Sizes in CASTLE-SWL

- Not sensitive in tested range [128, 1024]

		Train		Eval	
	n_{params}	Loss	PPL	Loss	PPL
Baseline-S	160M	2.892	18.037	2.901	18.197
CASTLE-SWL128-S	160M	2.883	17.859	2.889	17.971
CASTLE-SWL256-S	160M	2.888	17.952	2.895	18.092
CASTLE-SWL512-S	160M	2.885	17.910	<u>2.892</u>	<u>18.023</u>
CASTLE-SWL1024-S	160M	<u>2.885</u>	<u>17.901</u>	2.892	18.031

		Train		Eval	
	n_{params}	Loss	PPL	Loss	PPL
Baseline-XL	1.310B	2.548	12.779	2.543	12.723
CASTLE-SWL128-XL	1.304B	2.517	12.393	2.513	12.339
CASTLE-SWL256-XL	1.304B	2.514	12.353	2.510	12.300
CASTLE-SWL512-XL	1.304B	2.506	12.255	2.503	12.217
CASTLE-SWL1024-XL	1.304B	<u>2.514</u>	<u>12.351</u>	<u>2.509</u>	<u>12.294</u>

$$\mathbf{M}_t^U \in \mathbb{R}^{t \times t}$$

$-\infty$	0	0	0	$-\infty$	$-\infty$
$-\infty$	$-\infty$	0	0	0	$-\infty$
$-\infty$	$-\infty$	$-\infty$	0	0	0
$-\infty$	$-\infty$	$-\infty$	$-\infty$	0	0
$-\infty$	$-\infty$	$-\infty$	$-\infty$	$-\infty$	0
$-\infty$	$-\infty$	$-\infty$	$-\infty$	$-\infty$	$-\infty$

Ablation Studies on SiLU in softmax

- What roles does SiLU play in the combining step? Training stability
- Removing SiLU will cause blowing up when training XL models for both CASTLE and CASTLE-SWL
- Reducing lr can stabilize training but degrade performance severely

	Learning Rate	Train		Eval	
		Loss	PPL	Loss	PPL
CASTLE-SWL-XL	2×10^{-4}	2.506	12.255	2.503	12.217
CASTLE-SWL-XL w/o SiLU	1×10^{-4}	2.523	12.468	2.520	12.424
CASTLE-SWL-XL w/o SiLU	5×10^{-5}	2.571	13.084	2.571	13.075

Efficient Parallel Training

- Revisit the definition of lookahead keys

$$U^t = \begin{pmatrix} \mathbf{u}_1^t \\ \mathbf{u}_2^t \\ \vdots \\ \mathbf{u}_t^t \end{pmatrix} = \text{sigmoid}\left(\frac{\mathbf{Q}_t^U \mathbf{K}_t^{U^\top}}{\sqrt{d}} + \mathbf{M}_t^U\right) \mathbf{V}_t^U \in \mathbb{R}^{t \times d}, \quad \mathbf{s}_t^U = \frac{\mathbf{q}_t^C \mathbf{U}^{t^\top}}{\sqrt{d}} \in \mathbb{R}^{1 \times t}, \quad \mathbf{p}_t = \text{softmax}\left(\mathbf{s}_t^C - \text{SiLU}\left(\mathbf{s}_t^U\right)\right) \in \mathbb{R}^{1 \times t}$$

$$\text{attention}\left(\mathbf{X}^t\right) = \mathbf{p}_t \mathbf{V}_t^C \in \mathbb{R}^{1 \times d}.$$

- Equivalent mathematical formulation
 - Lookahead-keys attention scores

$$\mathbf{S}^U = \left(\frac{\mathbf{Q}^C \mathbf{V}^{U^\top}}{\sqrt{d}} \odot \widetilde{\mathbf{M}}^C \right) \left(\text{sigmoid} \left(\frac{\mathbf{Q}^U \mathbf{K}^{U^\top}}{\sqrt{d}} + \mathbf{M}^U \right) \right)^\top$$

- Parallel form of attention outputs

$$\text{Attention}(\mathbf{X}^L) = \text{row_softmax} \left(\frac{\mathbf{Q}^C \mathbf{K}^{C^\top}}{\sqrt{d}} + \mathbf{M}^C - \text{SiLU}(\mathbf{S}^U) \right) \mathbf{V}^C.$$

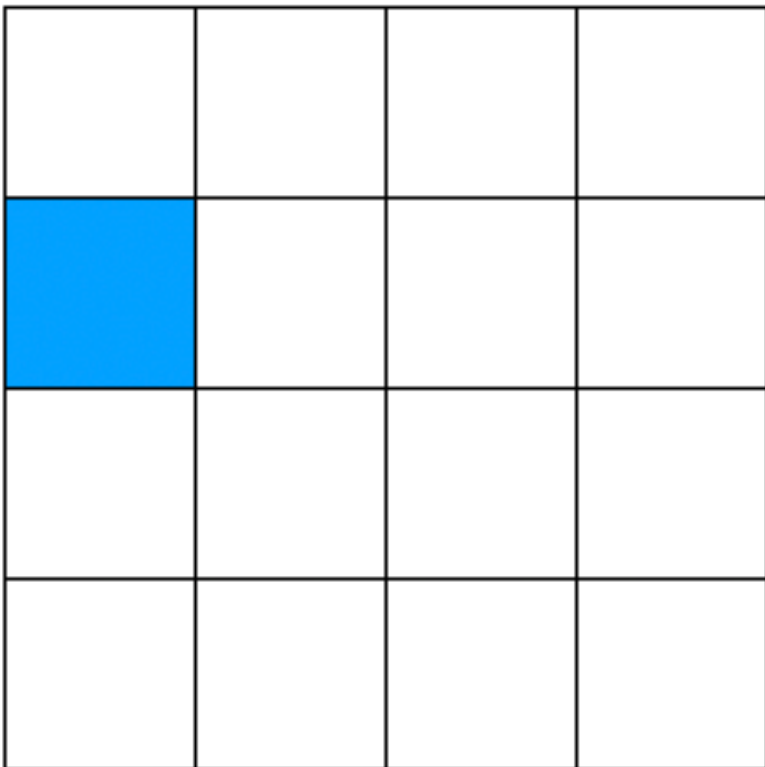
Efficient Parallel Training

$$S^U$$

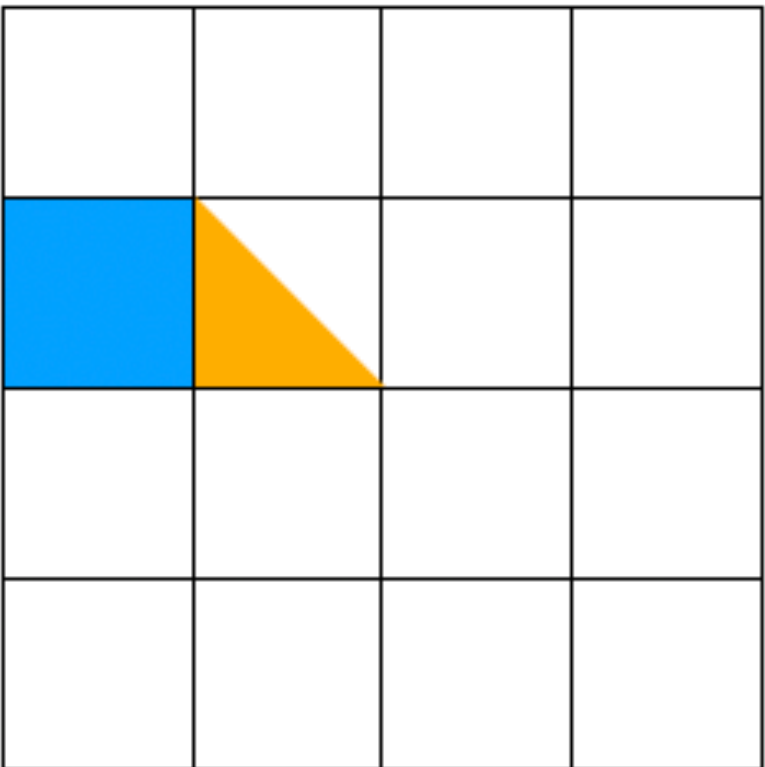
$$\frac{Q^C V^{U\top}}{\sqrt{d}} \odot \widetilde{M}^C$$

$$\left(\text{sigmoid} \left(\frac{Q^U K^{U\top}}{\sqrt{d}} + M^U \right) \right)^\top$$

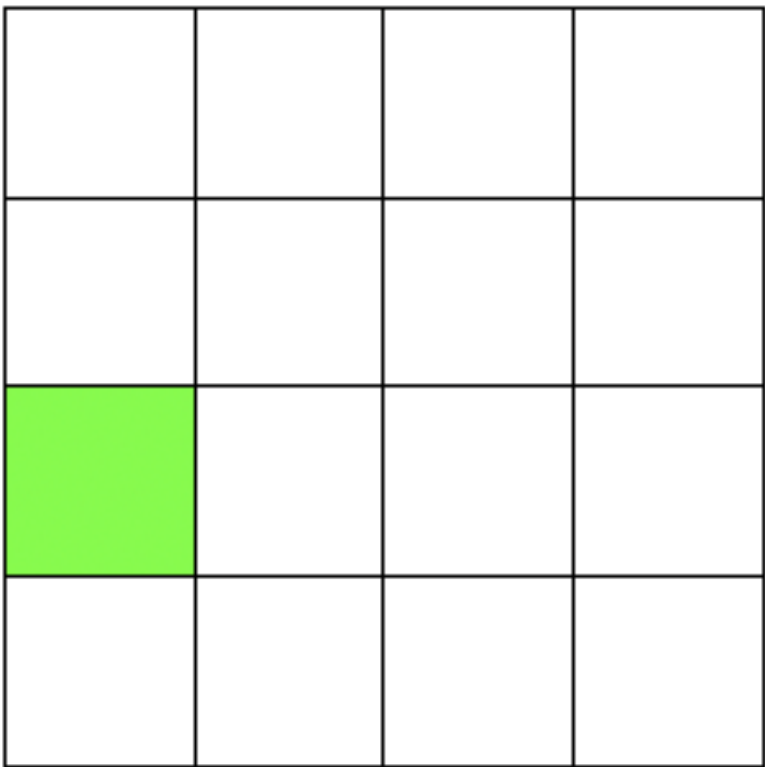
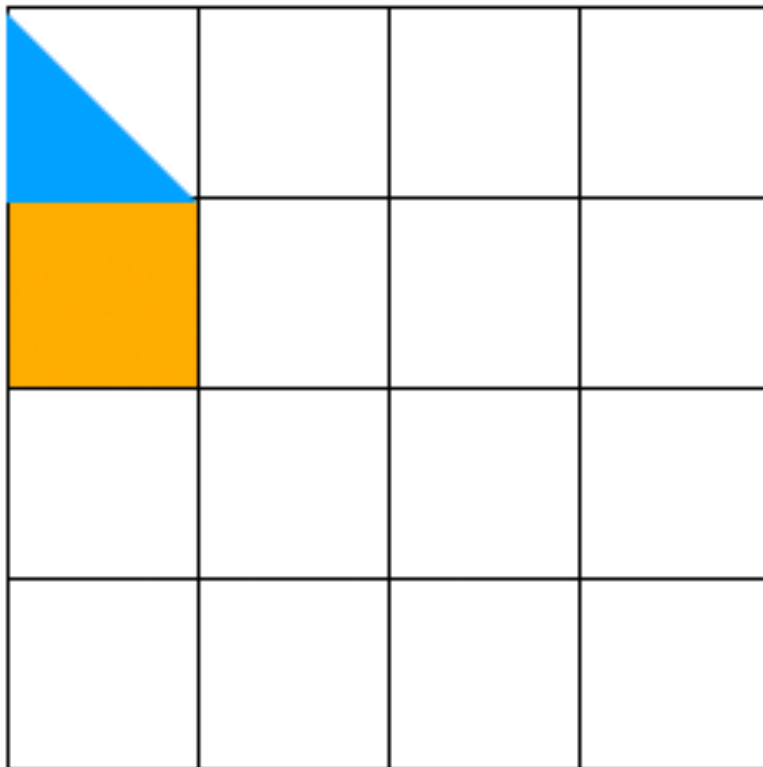
$$V_{T_1,:}^\top \quad V_{T_2,:}^\top \quad V_{T_3,:}^\top \quad V_{T_4,:}^\top$$



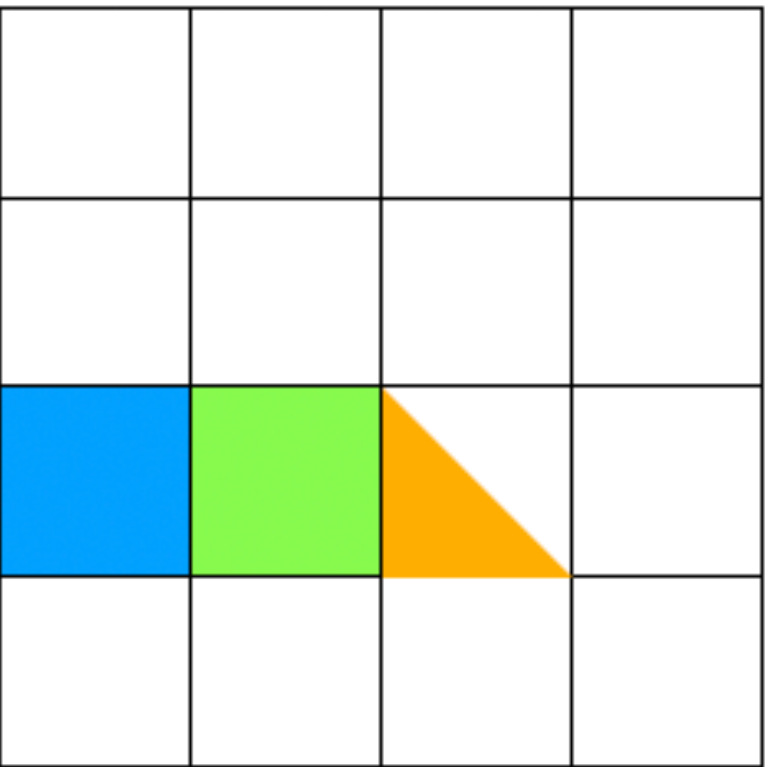
$$=$$



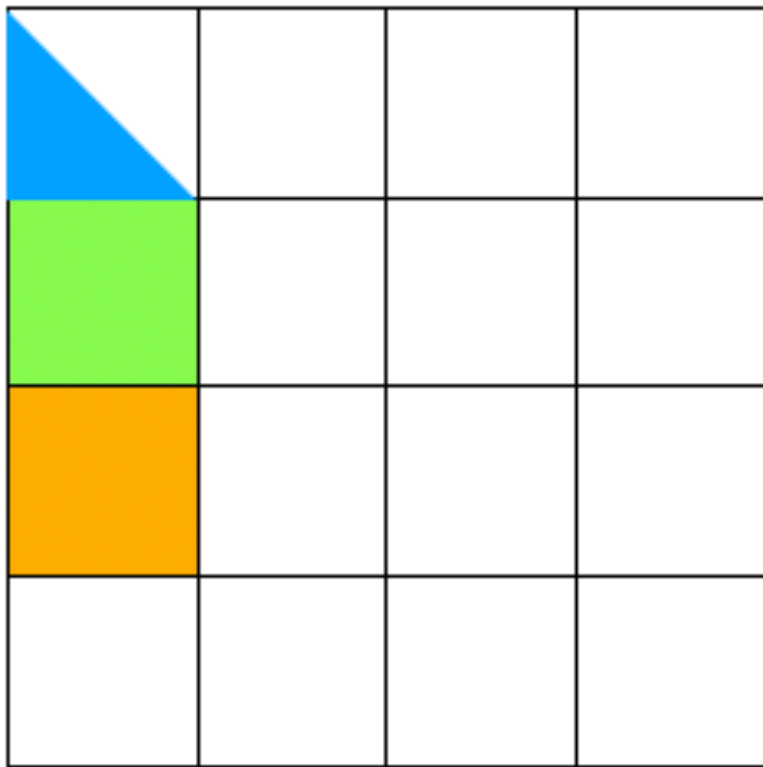
$$@$$



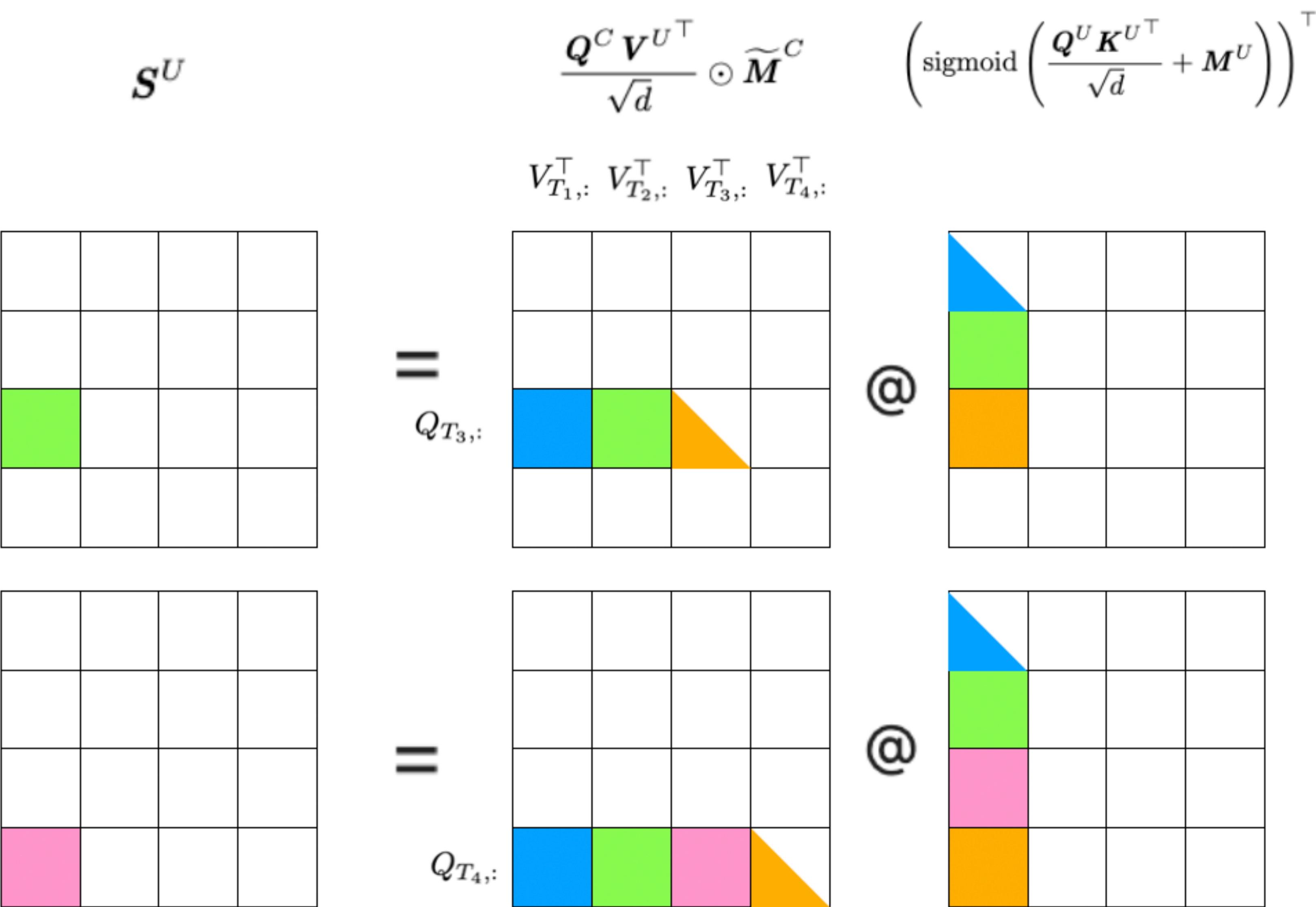
$$=$$



$$@$$



Efficient Parallel Training



Efficient Parallel Training

- Lookahead-key attention scores

$$\mathbf{s}_{T_{j+k}, T_j}^U = \mathbf{Q}_{T_{j+k}, :}^C \sum_{i=j}^{j+k-1} \frac{\mathbf{V}_{T_i, :}^U}{\sqrt{d}} \left(\text{sigmoid} \left(\frac{\mathbf{Q}_{T_j, :}^U \mathbf{K}_{T_i, :}^{U \top}}{\sqrt{d}} + \mathbf{M}_{T_j, T_i}^U \right) \right)^\top$$

$$+ \left(\frac{\mathbf{Q}_{T_{j+k}, :}^C \mathbf{V}_{T_{j+k}, :}^U}{\sqrt{d}} \odot \widetilde{\mathbf{M}}_{T_{j+k}, T_{j+k}}^C \right) \left(\text{sigmoid} \left(\frac{\mathbf{Q}_{T_j, :}^U \mathbf{K}_{T_{j+k}, :}^{U \top}}{\sqrt{d}} + \mathbf{M}_{T_j, T_{j+k}}^U \right) \right)^\top$$
- Utilizing masked low-rank structure

- Update auxiliary variable

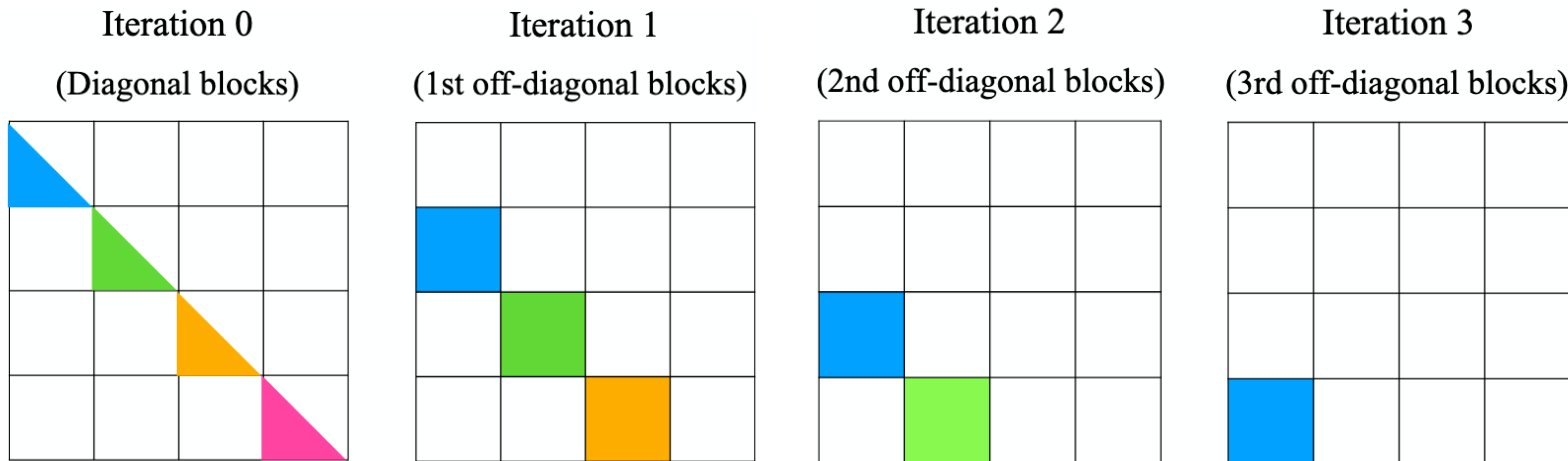
$$\mathbf{D}_{:, T_j}^{(k)} = \mathbf{D}_{:, T_j}^{(k-1)} + \mathbf{V}_{T_{j+k-1}, :}^U \left(\text{sigmoid} \left(\frac{\mathbf{Q}_{T_j, :}^U \mathbf{K}_{T_{j+k-1}, :}^{U \top}}{\sqrt{d}} + \mathbf{M}_{T_j, T_{j+k-1}}^U \right) \right)^\top$$

- Compute lookahead-key attention scores

$$\mathbf{s}_{T_{j+k}, T_j}^U = \frac{\mathbf{Q}_{T_{j+k}, :}^C \mathbf{D}_{:, T_j}^{(k)}}{\sqrt{d}} + \left(\frac{\mathbf{Q}_{T_{j+k}, :}^C \mathbf{V}_{T_{j+k}, :}^U}{\sqrt{d}} \odot \widetilde{\mathbf{M}}_{T_{j+k}, T_{j+k}}^C \right) \left(\text{sigmoid} \left(\frac{\mathbf{Q}_{T_j, :}^U \mathbf{K}_{T_{j+k}, :}^{U \top}}{\sqrt{d}} + \mathbf{M}_{T_j, T_{j+k}}^U \right) \right)^\top$$

Efficient Parallel Training

- Parallel scheme of forward pass



Efficient Inference with UQ-KV Cache

- Updating step: store U, Q

$$\mathbf{U}^t = \begin{pmatrix} \mathbf{U}^{t-1} + \text{sigmoid}\left(\frac{\mathbf{Q}_{t-1}^U \mathbf{k}_t^{U^\top}}{\sqrt{d}}\right) \mathbf{v}_t^U \\ \mathbf{0}^{1 \times d} \end{pmatrix}$$

- Combining step: store K, V

$$\mathbf{s}_t^C = \frac{\mathbf{q}_t^C \mathbf{K}_t^{C^\top}}{\sqrt{d}} \in \mathbb{R}^{1 \times t}$$

$$\mathbf{s}_t^U = \frac{\mathbf{q}_t^C \mathbf{U}^{t^\top}}{\sqrt{d}} \in \mathbb{R}^{1 \times t},$$

$$\mathbf{p}_t = \text{softmax}(\mathbf{s}_t^C - \text{SiLU}(\mathbf{s}_t^U)) \in \mathbb{R}^{1 \times t},$$

$$\text{attention}(\mathbf{X}^t) = \mathbf{p}_t \mathbf{V}_t^C \in \mathbb{R}^{1 \times d}.$$

Experiment Setup

- Model Architecture

$$\mathbf{Y}^{(l)} = \text{MultiHead-Attention}(\text{RMSNorm}(\mathbf{X}^{(l)}))$$

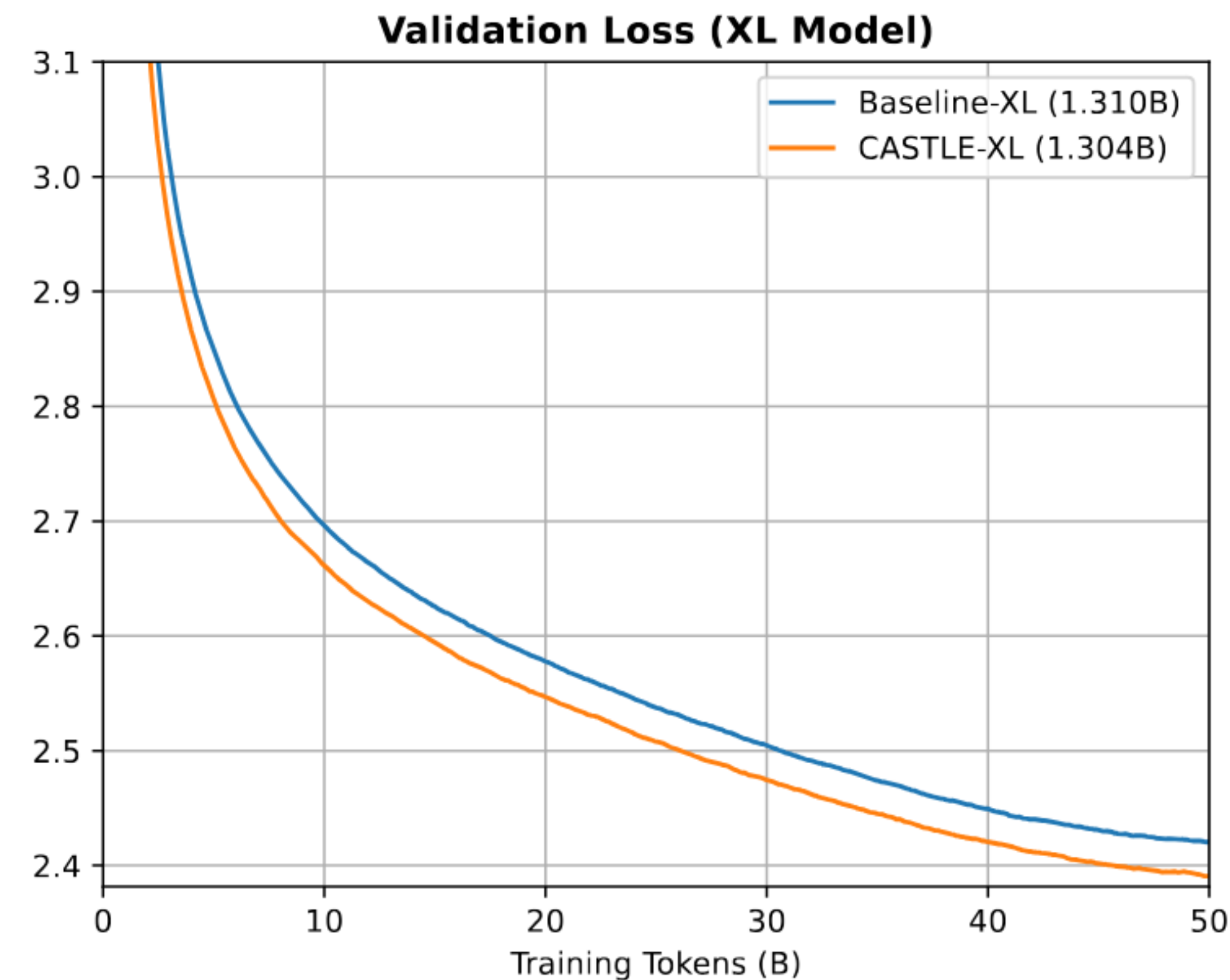
$$\mathbf{X}^{(l+1)} = \text{SwiGLU}(\text{RMSNorm}(\mathbf{Y}^{(l)})),$$
- Model configuration and training recipe

Model Name	n_{params}	n_{layers}	d_{model}	n_{heads}	d	Batch Size	Learning Rate
Baseline-S	160M	12	896 (=14 * 64)	14	64	0.5M	6.0×10^{-4}
CASTLE-S	160M	12	896	8	64		
Baseline-M	353M	24	1024 (=16 * 64)	16	64	0.5M	3.0×10^{-4}
CASTLE-M	351M	24	1024	9	64		
Baseline-L	756M	24	1536 (=16 * 96)	16	96	0.5M	2.5×10^{-4}
CASTLE-L	753M	24	1536	9	96		
Baseline-XL	1.310B	24	2048 (=16 * 128)	16	128	0.5M	2.0×10^{-4}
CASTLE-XL	1.304B	24	2048	9	128		

Experimental Results

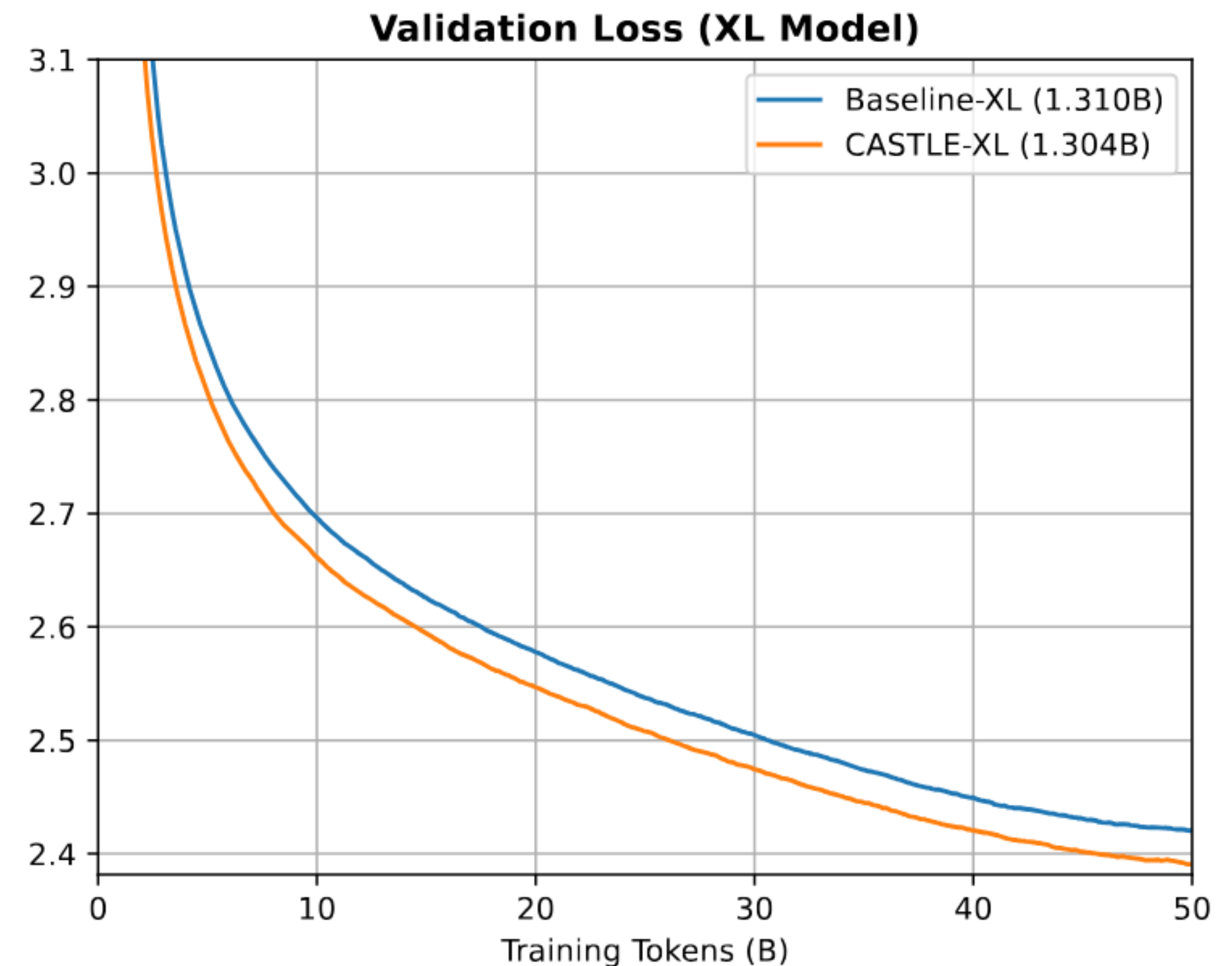
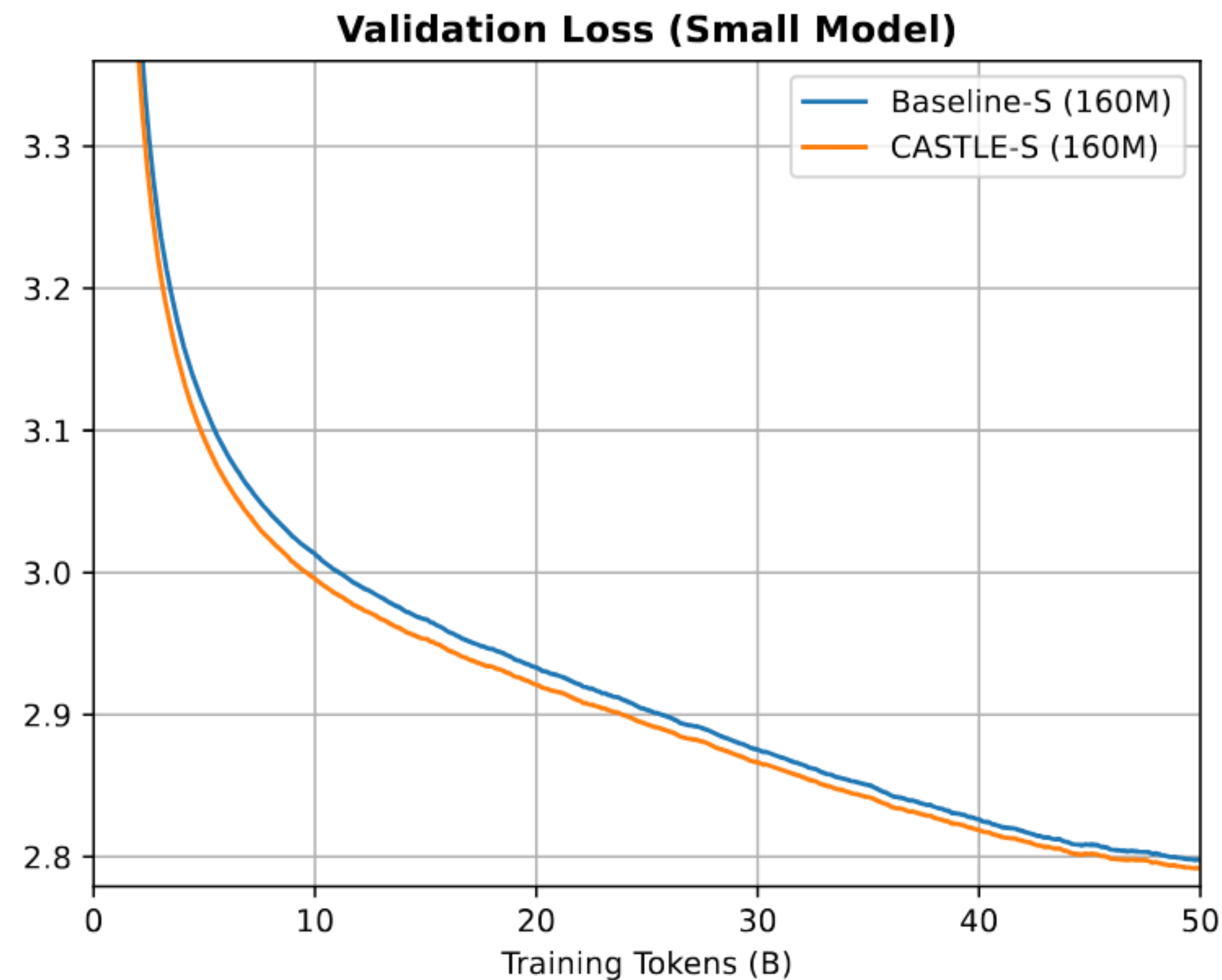
- Training & Validation Loss

	Train		Eval	
	Loss	PPL	Loss	PPL
Baseline-S	2.795	16.364	2.798	16.411
CASTLE-S	2.789	16.259	2.792	16.315
Baseline-M	2.641	14.030	2.639	14.004
CASTLE-M	2.616	13.684	2.615	13.665
Baseline-L	2.513	12.346	2.507	12.269
CASTLE-L	2.476	11.897	2.472	11.840
Baseline-XL	2.430	11.360	2.426	11.309
CASTLE-XL	2.401	11.031	2.391	10.922



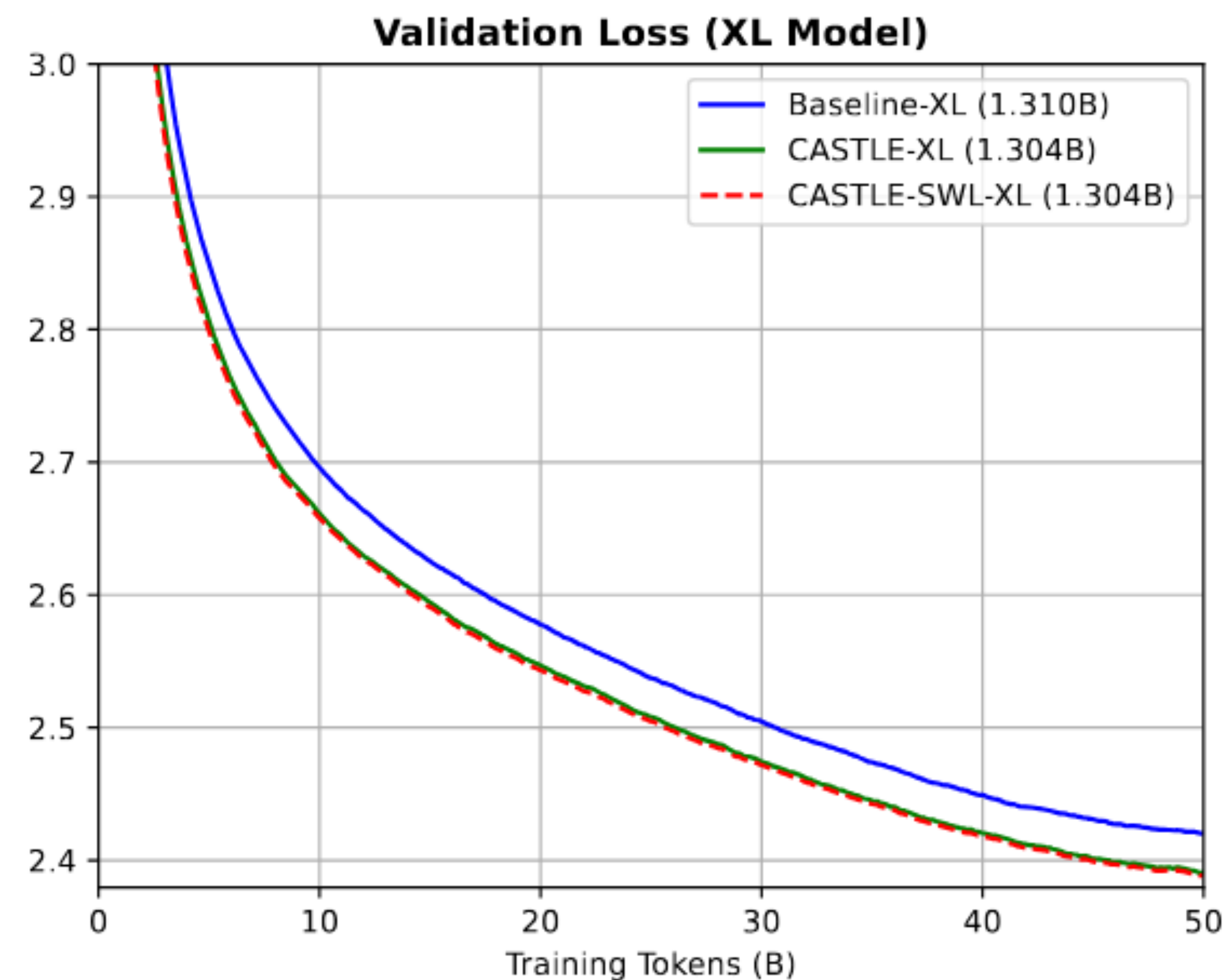
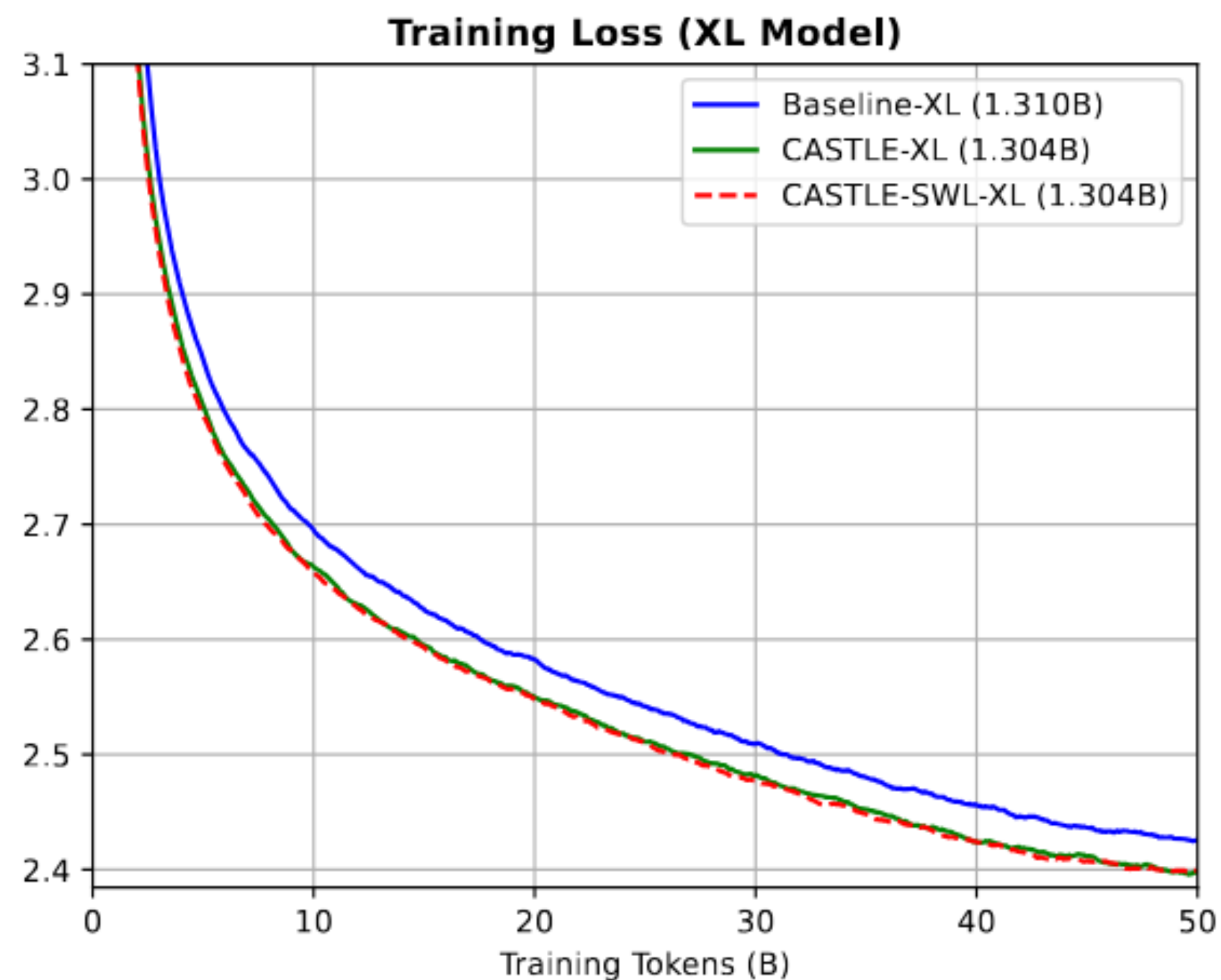
Experimental Results

- Comparison between loss curves of small models and XL models



Experimental Results: CASTLE vs. CASTLE-SWL

- CASTLE-SWL matches CASTLE on model scales (XL model as example here)



Experimental Results

- Downstream Tasks (0-shot)

Model Name	ARC-C	ARC-E	BoolQ	Hella.	MMLU	OBQA	PIQA	Wino.	Avg.
Baseline-S	26.71	<u>54.76</u>	52.51	35.78	<u>22.89</u>	30.40	63.98	52.57	42.45
CASTLE-S	26.19	56.69	59.85	36.28	23.00	<u>31.60</u>	<u>64.25</u>	<u>52.25</u>	43.76
CASTLE-SWL-S	<u>26.62</u>	54.12	<u>59.60</u>	<u>35.97</u>	22.87	32.60	64.31	50.51	<u>43.33</u>
Baseline-M	28.58	60.90	53.61	43.01	23.21	<u>33.40</u>	<u>67.95</u>	50.91	45.20
CASTLE-M	30.20	<u>61.36</u>	58.01	<u>43.24</u>	25.34	34.60	<u>67.95</u>	<u>52.64</u>	46.67
CASTLE-SWL-M	<u>29.52</u>	61.41	<u>57.03</u>	43.76	<u>23.29</u>	33.00	68.12	53.83	<u>46.24</u>
Baseline-L	<u>32.59</u>	<u>65.07</u>	57.49	47.45	23.57	32.60	<u>70.51</u>	50.75	47.50
CASTLE-L	32.34	65.15	<u>57.65</u>	<u>47.87</u>	24.51	<u>35.60</u>	70.78	<u>53.51</u>	<u>48.43</u>
CASTLE-SWL-L	32.76	63.51	60.95	48.50	<u>23.72</u>	36.00	70.02	54.85	48.79
Baseline-XL	33.79	66.62	<u>61.04</u>	51.40	26.72	36.20	72.58	54.06	50.30
CASTLE-XL	<u>35.32</u>	<u>67.51</u>	62.81	52.15	23.74	<u>37.00</u>	70.67	56.59	<u>50.72</u>
CASTLE-SWL-XL	36.43	69.07	60.24	<u>51.99</u>	<u>24.47</u>	37.40	<u>71.27</u>	<u>55.09</u>	50.74

Experimental Results

- Downstream tasks (5-shot)

Model Name	ARC-C	ARC-E	BoolQ	Hella.	MMLU	OBQA	PIQA	Wino.	Avg.
Baseline-S	25.68	<u>54.97</u>	<u>56.09</u>	33.81	25.54	28.20	63.98	52.57	42.60
CASTLE-S	<u>26.02</u>	54.25	57.13	<u>35.24</u>	25.22	<u>29.80</u>	<u>64.53</u>	50.99	<u>42.90</u>
CASTLE-SWL-S	27.39	56.19	<u>56.09</u>	35.46	<u>25.34</u>	30.20	64.85	<u>51.22</u>	43.34
Baseline-M	31.06	62.46	48.47	42.83	<u>25.22</u>	<u>33.00</u>	68.39	51.78	45.40
CASTLE-M	<u>32.17</u>	64.06	54.62	<u>43.47</u>	<u>25.22</u>	33.80	69.48	<u>52.49</u>	<u>46.91</u>
CASTLE-SWL-M	32.85	<u>63.85</u>	<u>54.19</u>	44.18	26.43	33.80	<u>69.04</u>	53.43	47.22
Baseline-L	33.36	63.64	<u>59.24</u>	46.16	26.82	33.40	<u>69.53</u>	54.06	48.28
CASTLE-L	37.37	67.89	50.95	<u>47.71</u>	<u>26.11</u>	<u>34.20</u>	70.18	54.06	<u>48.56</u>
CASTLE-SWL-L	<u>36.26</u>	<u>65.53</u>	60.58	48.55	24.70	34.80	69.10	<u>53.67</u>	49.15
Baseline-XL	35.58	65.78	61.07	50.84	26.71	36.20	<u>71.27</u>	52.72	50.02
CASTLE-XL	39.08	70.24	62.60	<u>51.63</u>	24.16	37.40	71.00	58.33	<u>51.80</u>
CASTLE-SWL-XL	<u>38.99</u>	<u>70.08</u>	<u>61.74</u>	52.35	<u>25.85</u>	<u>37.20</u>	72.52	<u>56.75</u>	51.93

Ablation Studies on Number of Keys

- Is the improvement from more keys?

- Model configurations

Model Name	n_{params}	n_{layers}	d_{model}	n_{heads}	$n_{\text{LookaheadKeys}} + n_{\text{CausalKeys}}$	d
Baseline-M	353M	24	1024 (=16 * 64)	16	16	64
CASTLE-M	351M	24	1024	9	18	64
CASTLE-M-16	340M	24	1024	8	16	64
Baseline-XL	1.310B	24	2048 (=16 * 128)	16	16	128
CASTLE-XL	1.304B	24	2048	9	18	128
CASTLE-XL-16	1.260B	24	2048	8	16	128

- Training & Validation loss

	n_{params}	Train		Eval	
		Loss	PPL	Loss	PPL
Baseline-M	353M	2.740	15.483	2.742	15.523
CASTLE-M	351M	2.709	15.018	2.711	15.039
CASTLE-M-16	340M	<u>2.714</u>	<u>15.093</u>	<u>2.716</u>	<u>15.126</u>
Baseline-XL	1.310B	2.548	12.779	2.543	12.723
CASTLE-XL	1.304B	2.507	12.267	2.503	12.219
CASTLE-XL-16	1.260B	<u>2.514</u>	<u>12.349</u>	<u>2.511</u>	<u>12.316</u>

	n_{params}	Train		Eval	
		Loss	PPL	Loss	PPL
Baseline-M	353M	2.740	15.483	2.742	15.523
CASTLE-SWL-M	351M	2.710	15.036	2.713	15.068
CASTLE-SWL-M-16	340M	<u>2.716</u>	<u>15.117</u>	<u>2.718</u>	<u>15.150</u>
Baseline-XL	1.310B	2.548	12.779	2.543	12.723
CASTLE-SWL-XL	1.304B	2.506	12.255	2.503	12.217
CASTLE-SWL-XL-16	1.260B	<u>2.513</u>	<u>12.339</u>	<u>2.508</u>	<u>12.276</u>

Thank you!