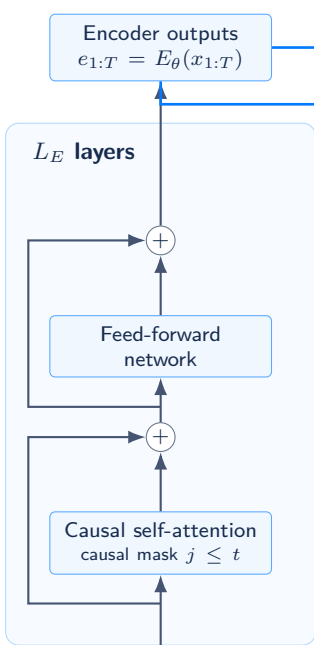


Causal encoder

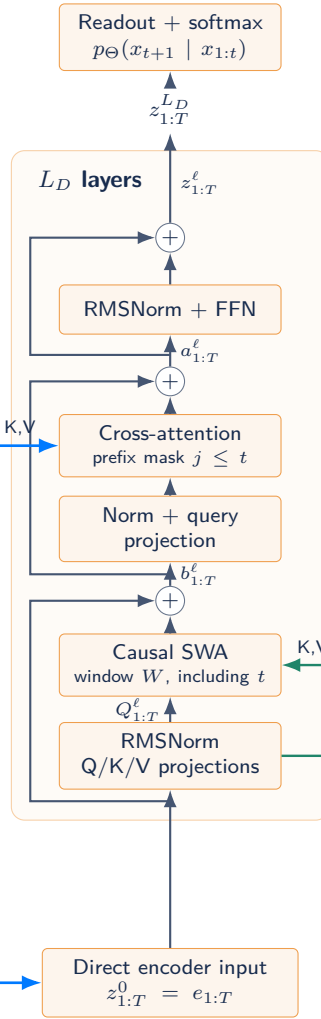
Known sequence $x_{1:T}$

Causal masking allows parallel encoder prefill across positions.



RLT-0 decoder

Parallel positions within each layer



One prediction per position t

Readout + softmax
 $p_{\Theta}(x_{t+1} \mid x_{1:t})$

Training / prefill
Positions in parallel;
layers in sequence.

Separate KV per layer;
window at position t :
 $j \leq t$ and $t - j < W$

For continuation
save $C_T^{D,\ell}$
last $W - 1$ entries

Layer ℓ SWA KV
 $\{(k_j^{D,\ell}, v_j^{D,\ell})\}_{j=1}^T$

Generation remains
one token at a time.

Attention includes
output projection.
+ denotes a residual addition.
Keys include positional
transformations.