



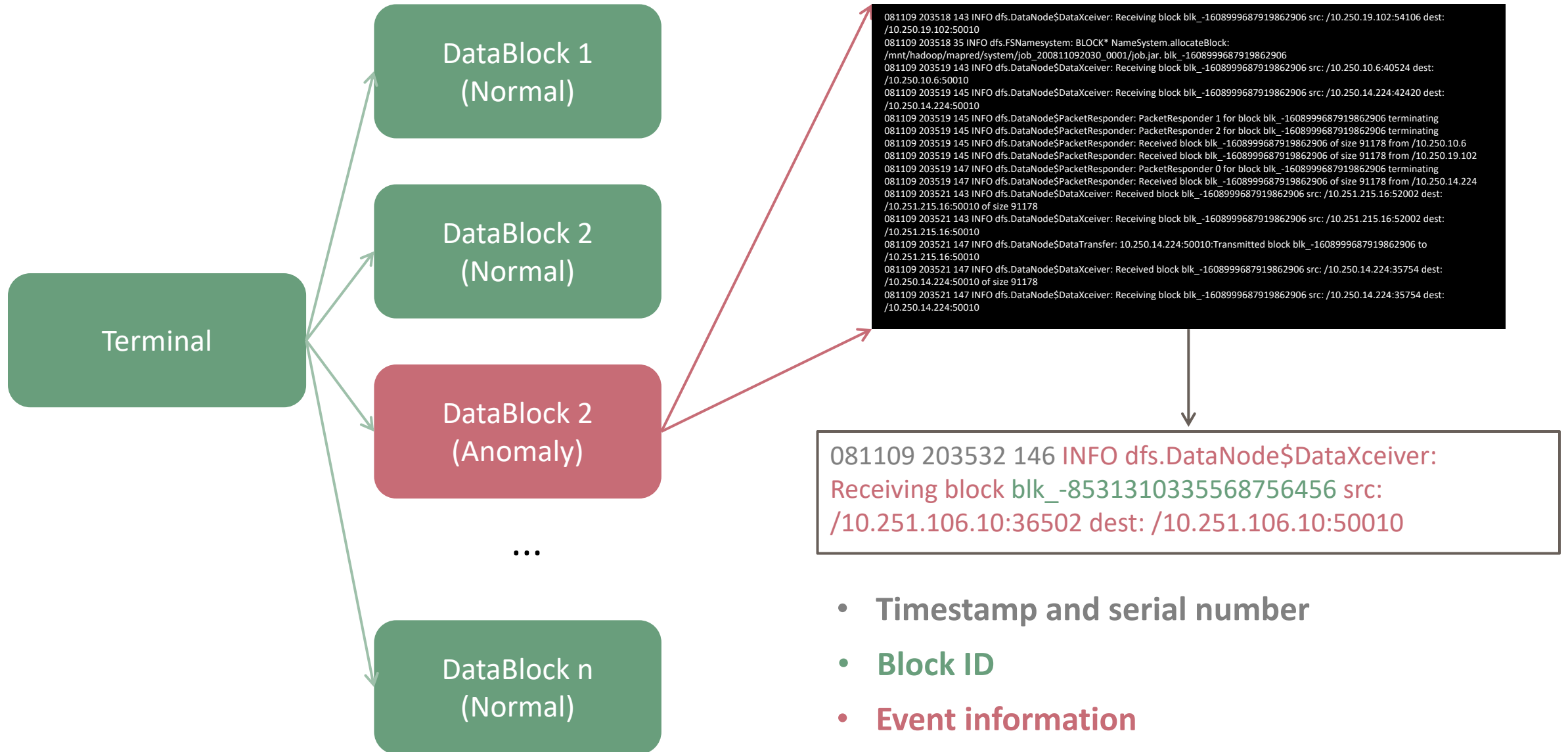
Advancing Log Anomaly Detection with ModernBERT: A Comparative Study of Novel Architectures and LogBERT

MINGTONG DU, HAORAN LI

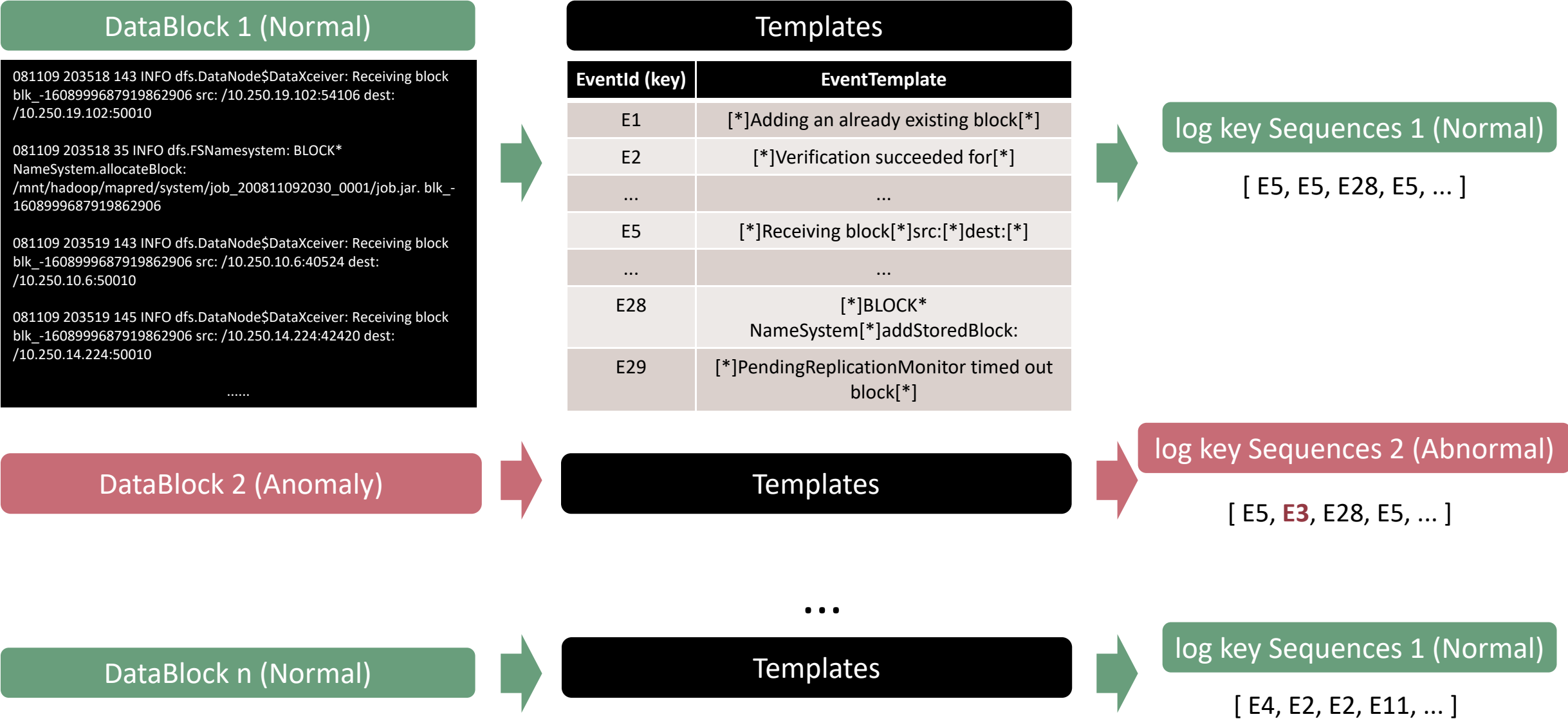
Overview

- Computer logs
- LogBERT (2021)
- Motivation
- Our Architecture 1: ModernLogBERT - BC (Binary Classification)
- Our Architecture 2: ModernLogBERT - MLWP (Masked Log Word Prediction)
- Experiments
- Conclusion

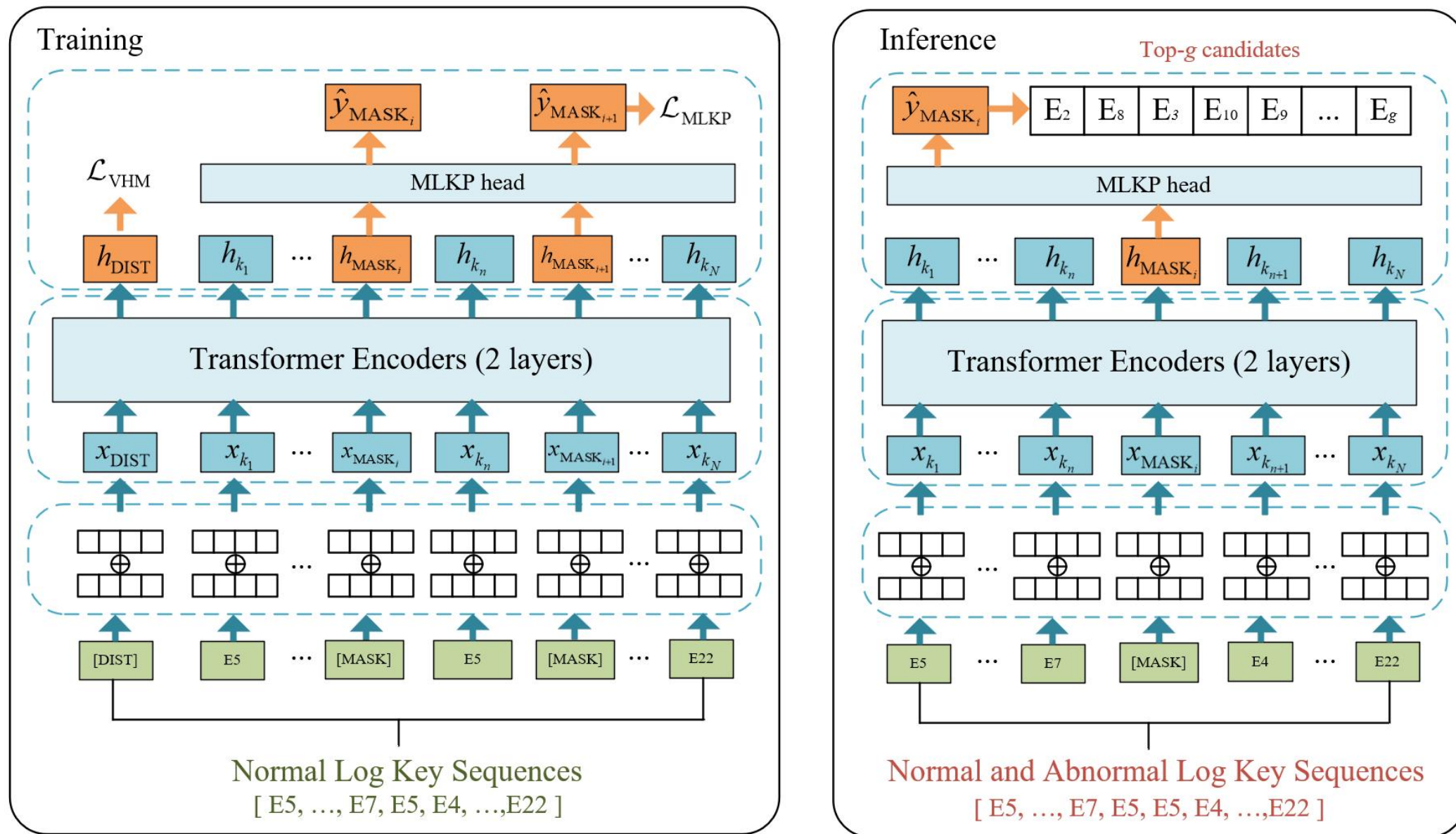
Computer logs



Computer logs — Logs to Event sequences



LogBERT: Log Anomaly Detection via BERT^[1](Guo, 2021)



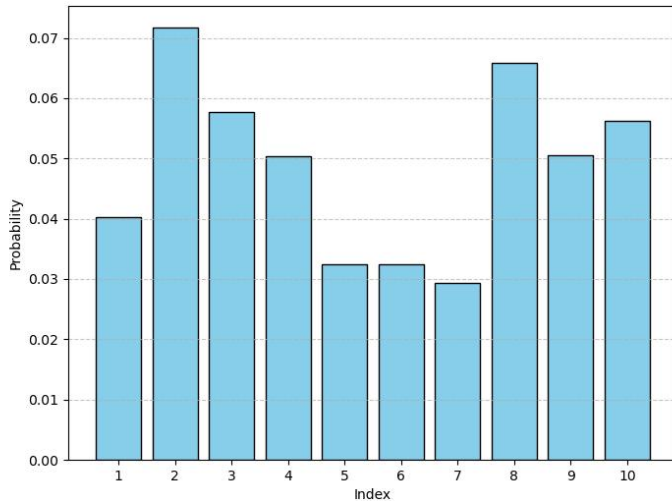
[1] Guo, Haixuan, Shuhan Yuan, and Xintao Wu. "Logbert: Log anomaly detection via bert." 2021 international joint conference on neural networks (IJCNN). IEEE, 2021.

LogBERT Task1: Masked Log Key Prediction

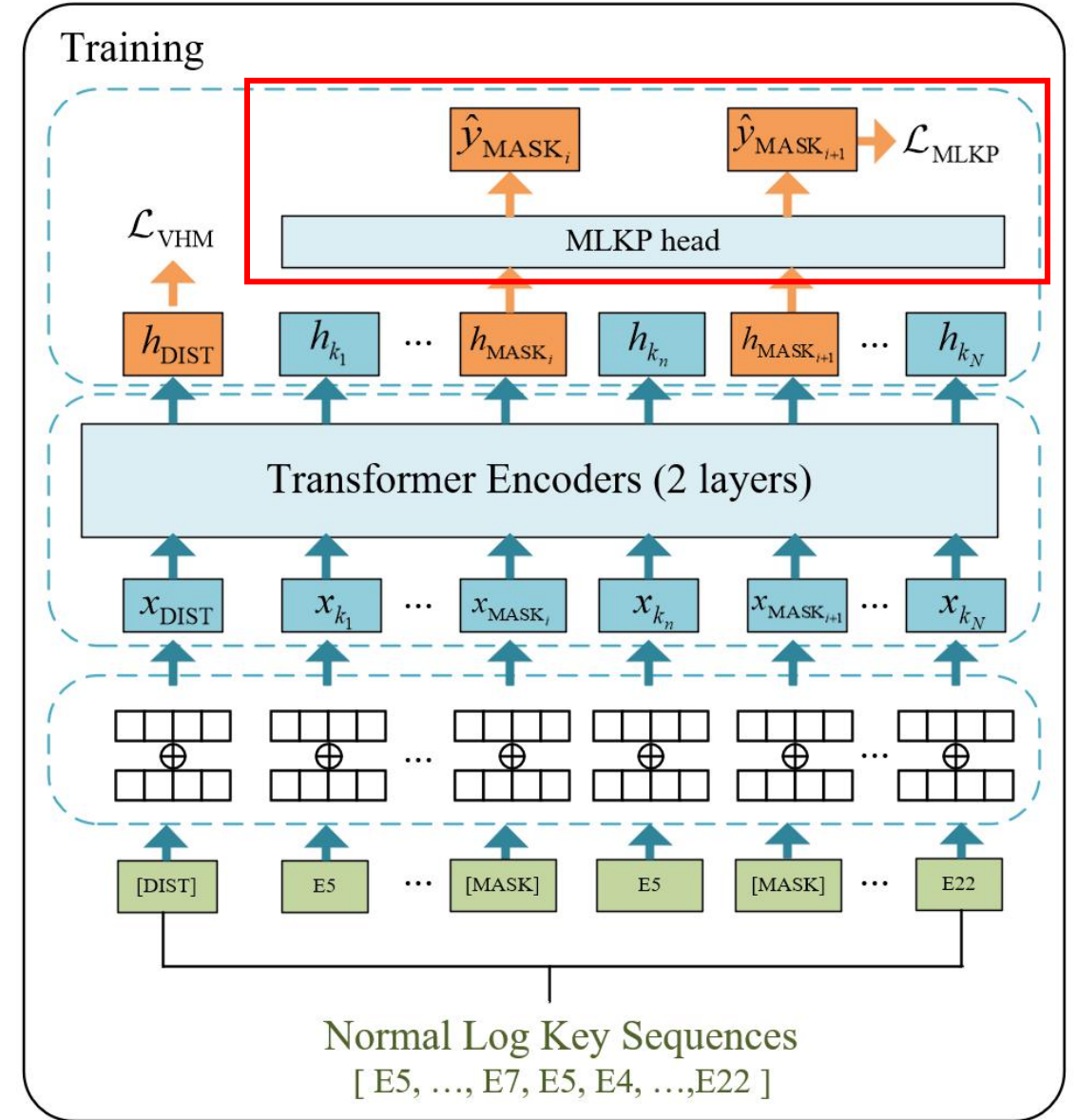
MLKP head:

$$\hat{y}_{\text{MASK}_i} = \text{Softmax}(W h_{\text{MASK}_i} + b)$$

- Inputs h_{MASK_i} : 1×768 ($1 \times \text{hidden_size}$)
- Weights W : 768×22 ($\text{hidden_size} \times \text{vocab_size}$)
- Output $\hat{y}_{\text{MASK}_i}$: 22×1 ($\text{vocab_size} \times 1$)



$$\text{Loss (Cross-Entropy): } \mathcal{L}_{\text{MLKP}} = -\frac{1}{N} \sum_{j=1}^N \sum_{i=1}^M y_{\text{MASK}_i}^j \log \hat{y}_{\text{MASK}_i}^j$$

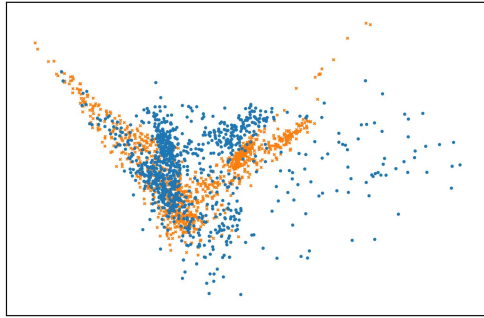


LogBERT Task2: Volume of Hypersphere Minimization

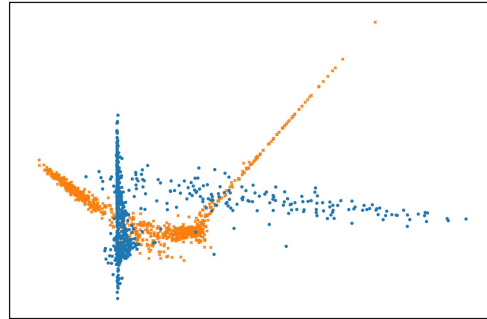
VHM Loss :

$$\mathcal{L}_{\text{VHM}} = -\frac{1}{N} \sum_{j=1}^N \|h_{\text{DIST}}^j - c\|^2$$

- Where c is the center of all h_{DIST}^j : $c = \frac{1}{N} \sum_{j=1}^N h_{\text{DIST}}^j$



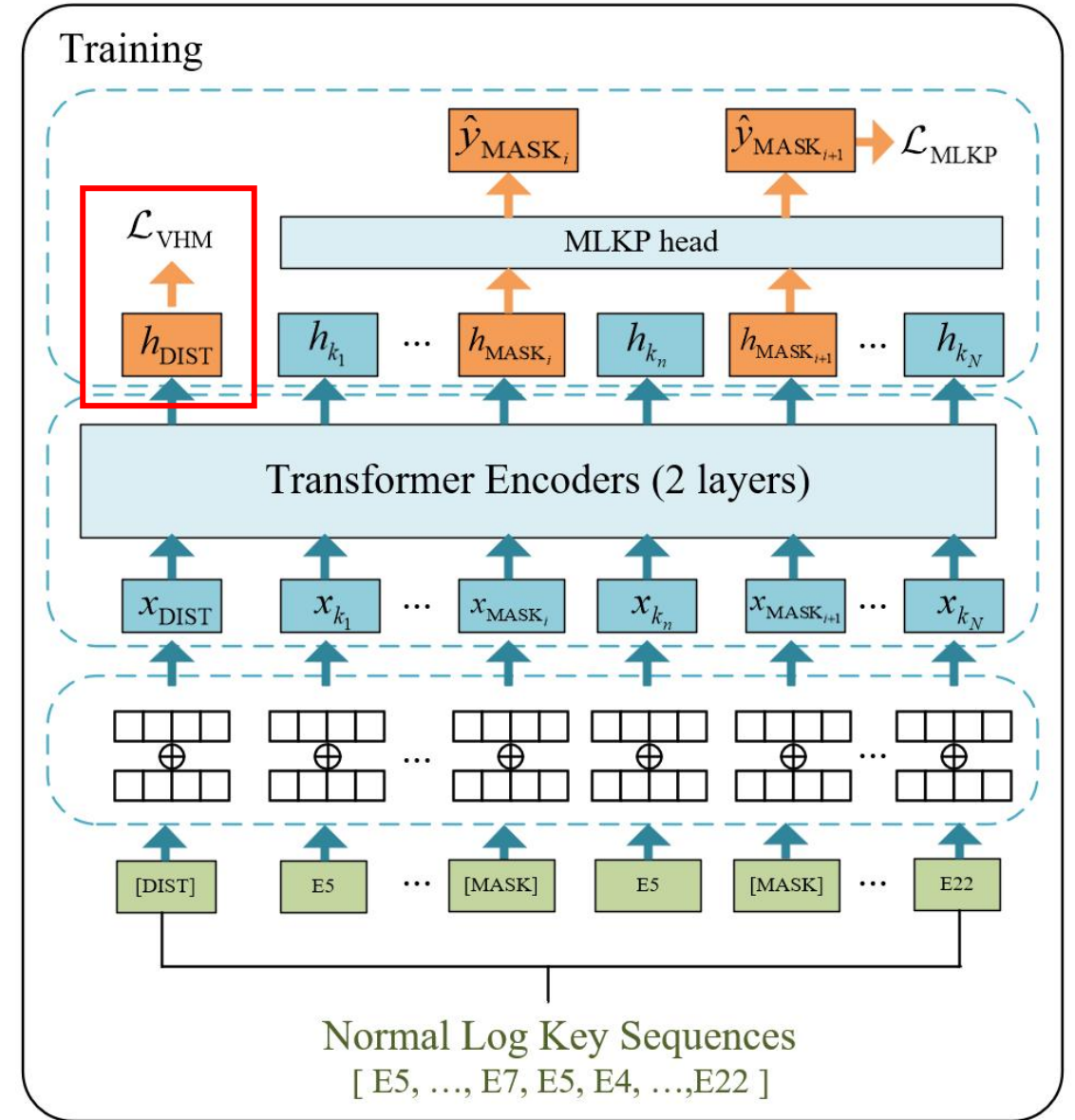
With VHM loss



Without VHM loss

Fig 1. h_{DIST} mapped into a two dimensional space

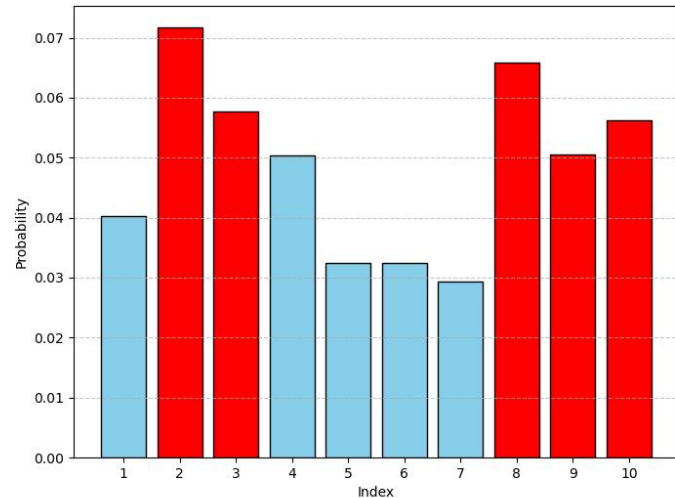
Backpropagation: $\mathcal{L} = \mathcal{L}_{\text{MLKP}} + \alpha \mathcal{L}_{\text{VHM}}$



LogBERT Anomaly Detection

$$\hat{y}_{\text{MASK}_i}^j = \text{Softmax}(W h_{\text{MASK}_i}^j + b)$$

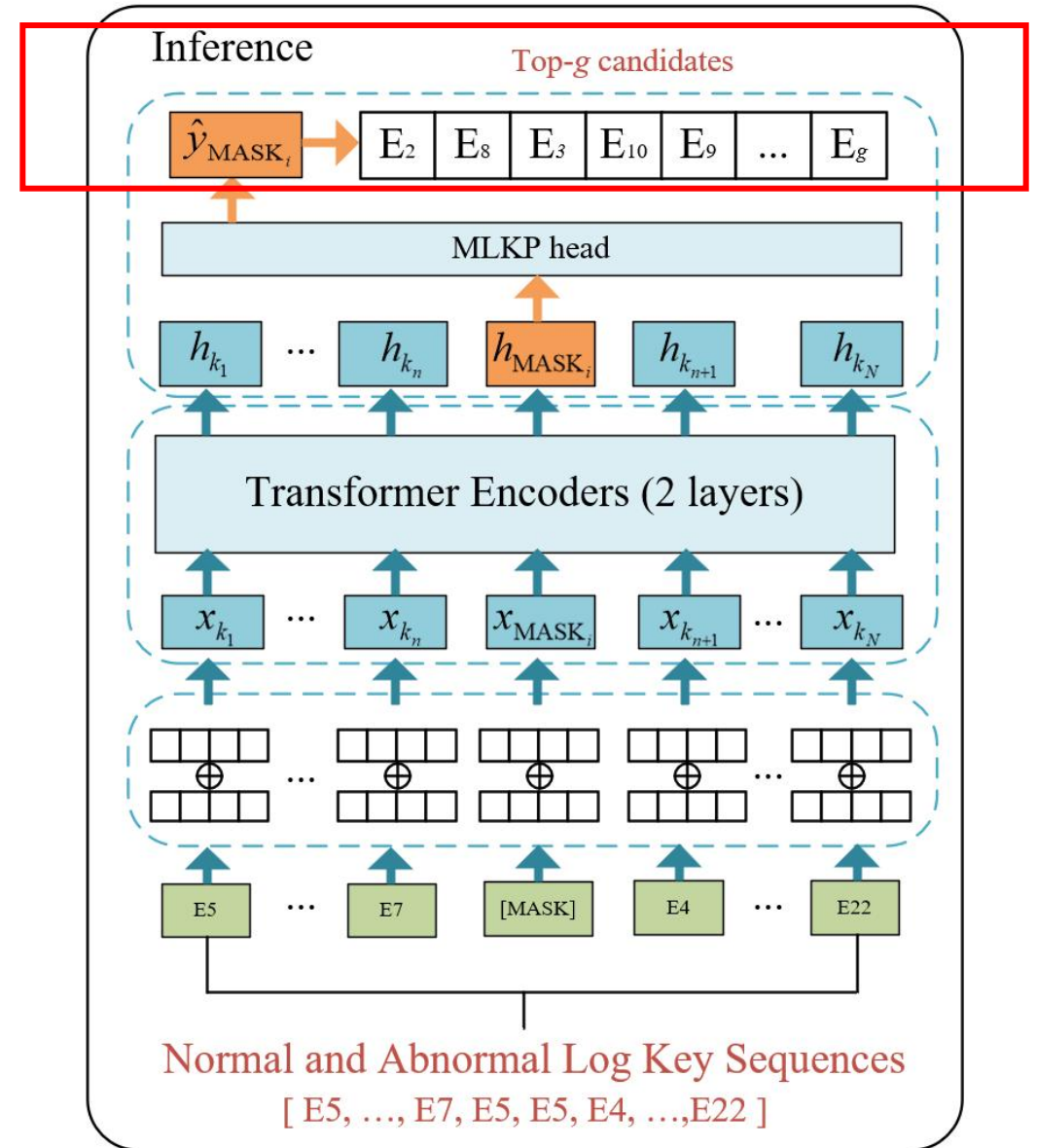
Probability for each token:



$g=5$

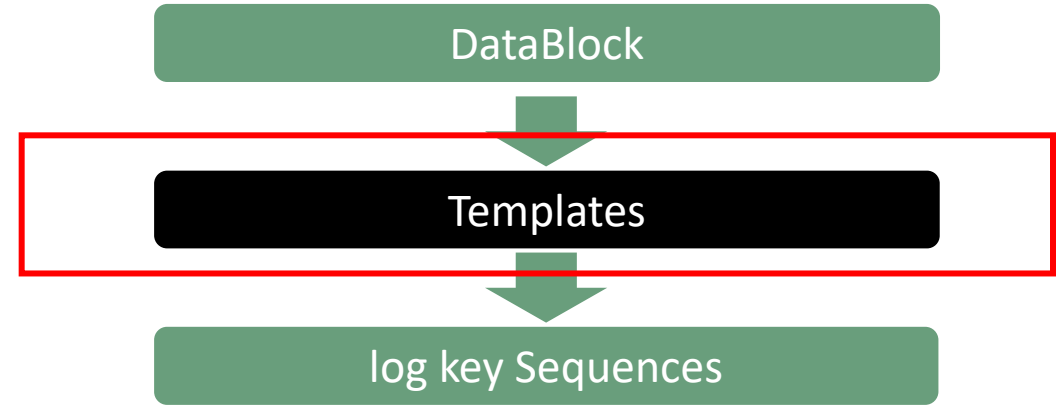
Threshold g and r :

- g : Determine whether the true log key is in the top- g highest likelihoods candidate set.
- r : If there are more than r log keys not in the candidate set, the sequence will be abnormal



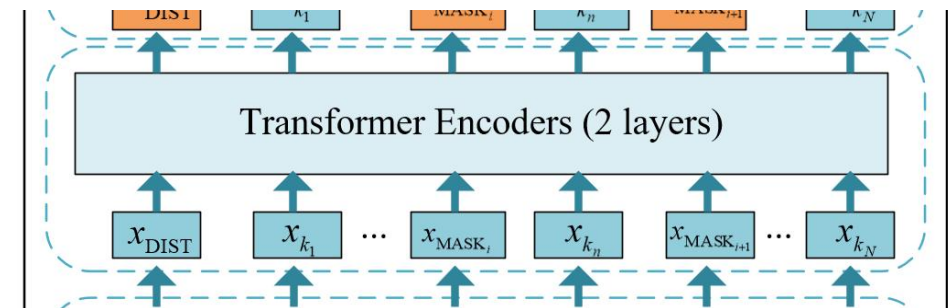
Motivation

- Templates Rely too much on the dataset and have no generalization ability
- The log information itself contains certain semantics
- The ability of two-layer transformer to extract semantic information is insufficient



```
081109 203532 146 INFO dfs.DataNode$DataXceiver:
Receiving block blk_-8531310335568756456 src:
/10.251.106.10:36502 dest: /10.251.106.10:50010
```

- **Event information**



Data Preprocessing

DataBlock

```
081109 203518 143 INFO dfs.DataNode$DataXceiver: Receiving block  
blk_-1608999687919862906 src: /10.250.19.102:54106 dest:  
/10.250.19.102:50010  
  
081109 203518 35 INFO dfs.FSNamesystem: BLOCK*  
NameSystem.allocateBlock:  
/mnt/hadoop/mapred/system/job_200811092030_0001/job.jar. blk_-  
1608999687919862906  
  
081109 203519 143 INFO dfs.DataNode$DataXceiver: Receiving block  
blk_-1608999687919862906 src: /10.250.10.6:40524 dest:  
/10.250.10.6:50010  
  
081109 203519 145 INFO dfs.DataNode$DataXceiver: Receiving block  
blk_-1608999687919862906 src: /10.250.14.224:42420 dest:  
/10.250.14.224:50010  
  
.....
```

```
081109 203532 146 INFO dfs.DataNode$DataXceiver:  
Receiving block blk_-8531310335568756456 src:  
/10.251.106.10:36502 dest: /10.251.106.10:50010
```

- Timestamp and serial number
- Block ID
- Event information

Remove Dynamic
Information

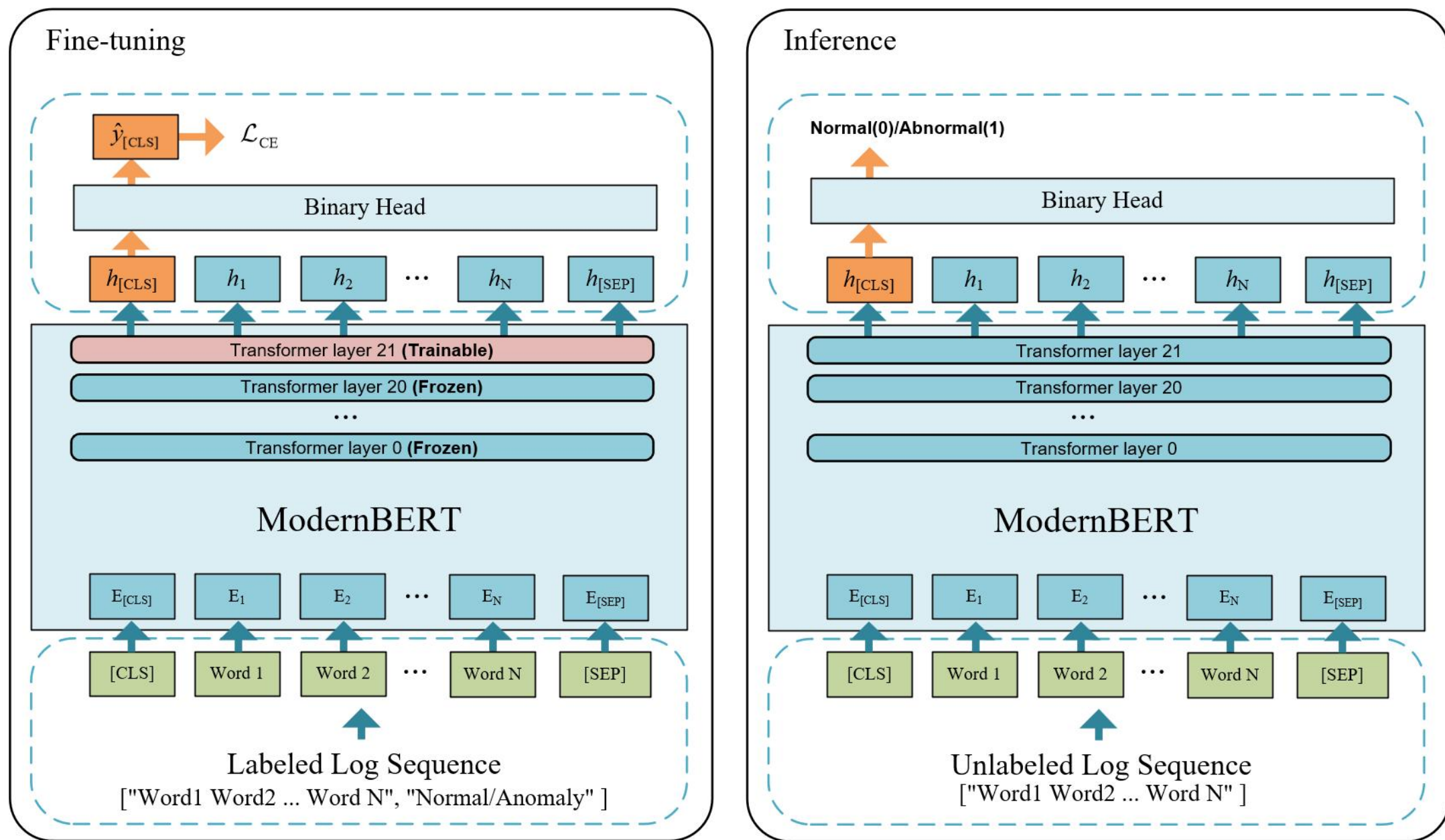
Log Sequences

```
INFO dfs.DataNode$DataXceiver: Receiving block src: dest:  
  
INFO dfs.FSNamesystem: BLOCK* NameSystem.allocateBlock:  
  
INFO dfs.DataNode$DataXceiver: Receiving block src dest  
  
INFO dfs.DataNode$DataXceiver: Receiving block src: dest:  
  
.....
```

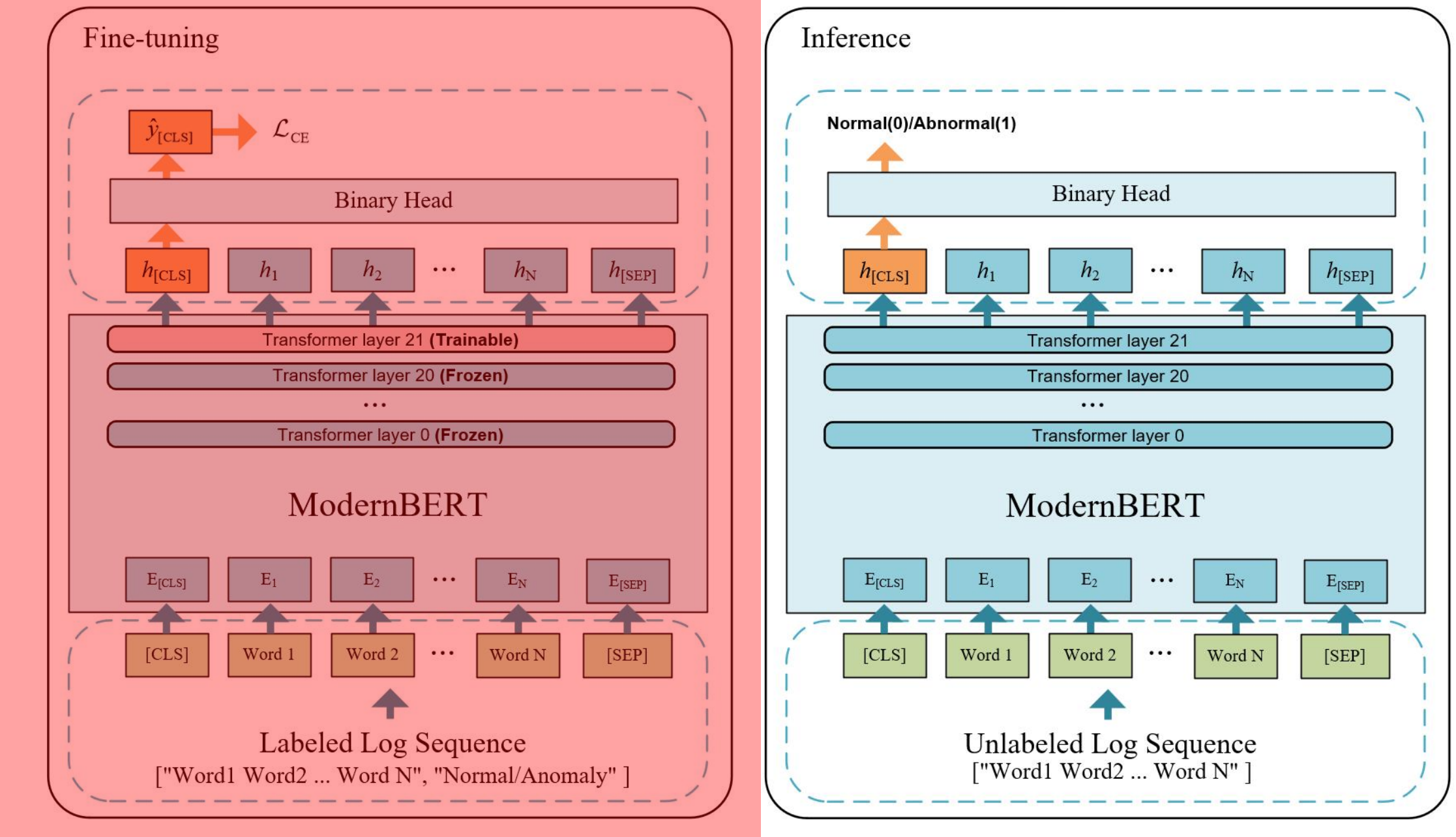
```
INFO dfs.DataNode$DataXceiver: Receiving block src:  
dest:
```

- Event information

Our Architecture 1: ModernLogBERT - BC (Binary Classification)



Our Architecture 1: ModernLogBERT - BC (Binary Classification)



ModernLogBERT - BC Task: Binary Classification

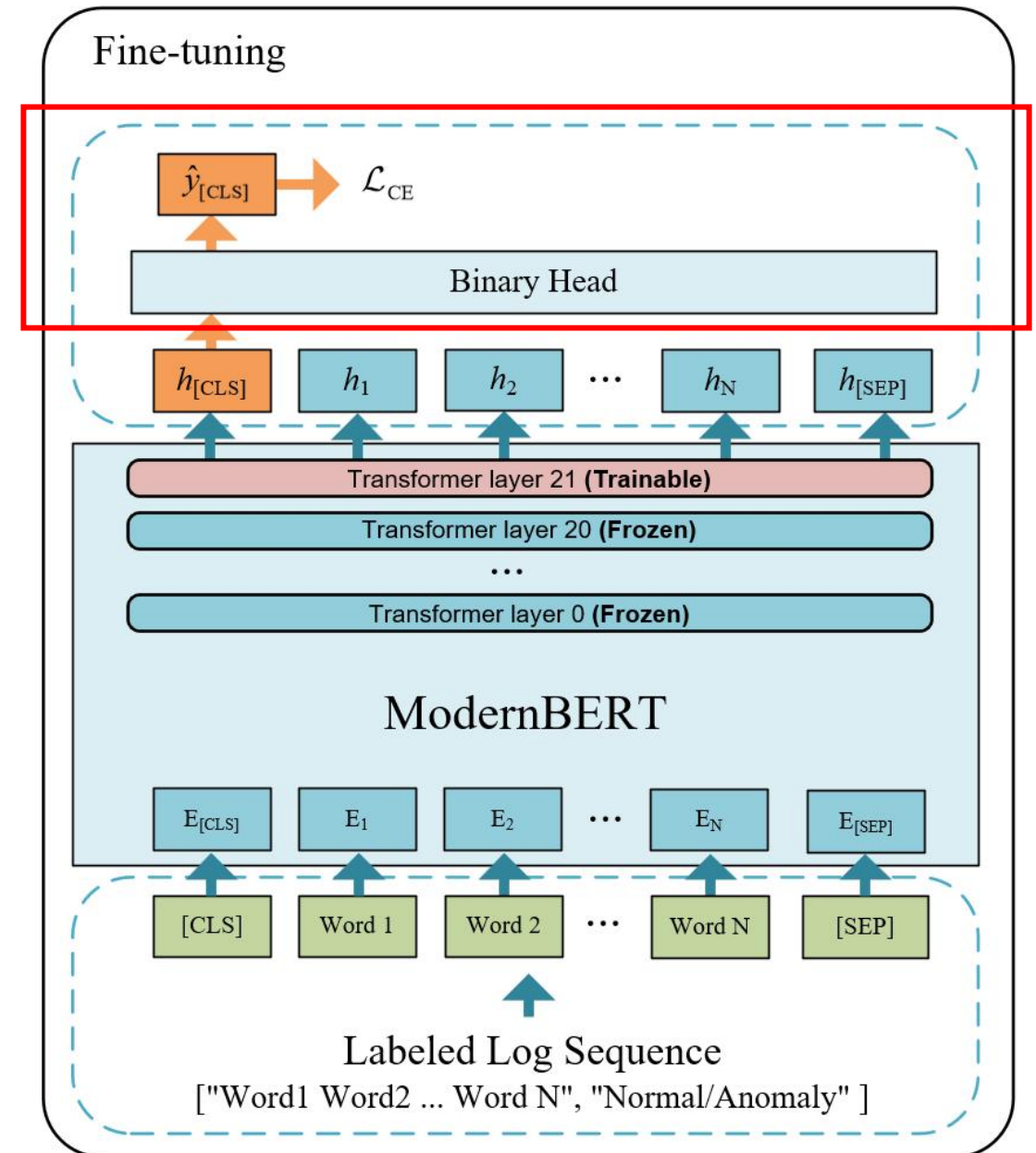
Binary classification head:

$$\hat{y}_{[\text{CLS}]} = \text{Softmax}(Wh_{[\text{CLS}]} + b)$$

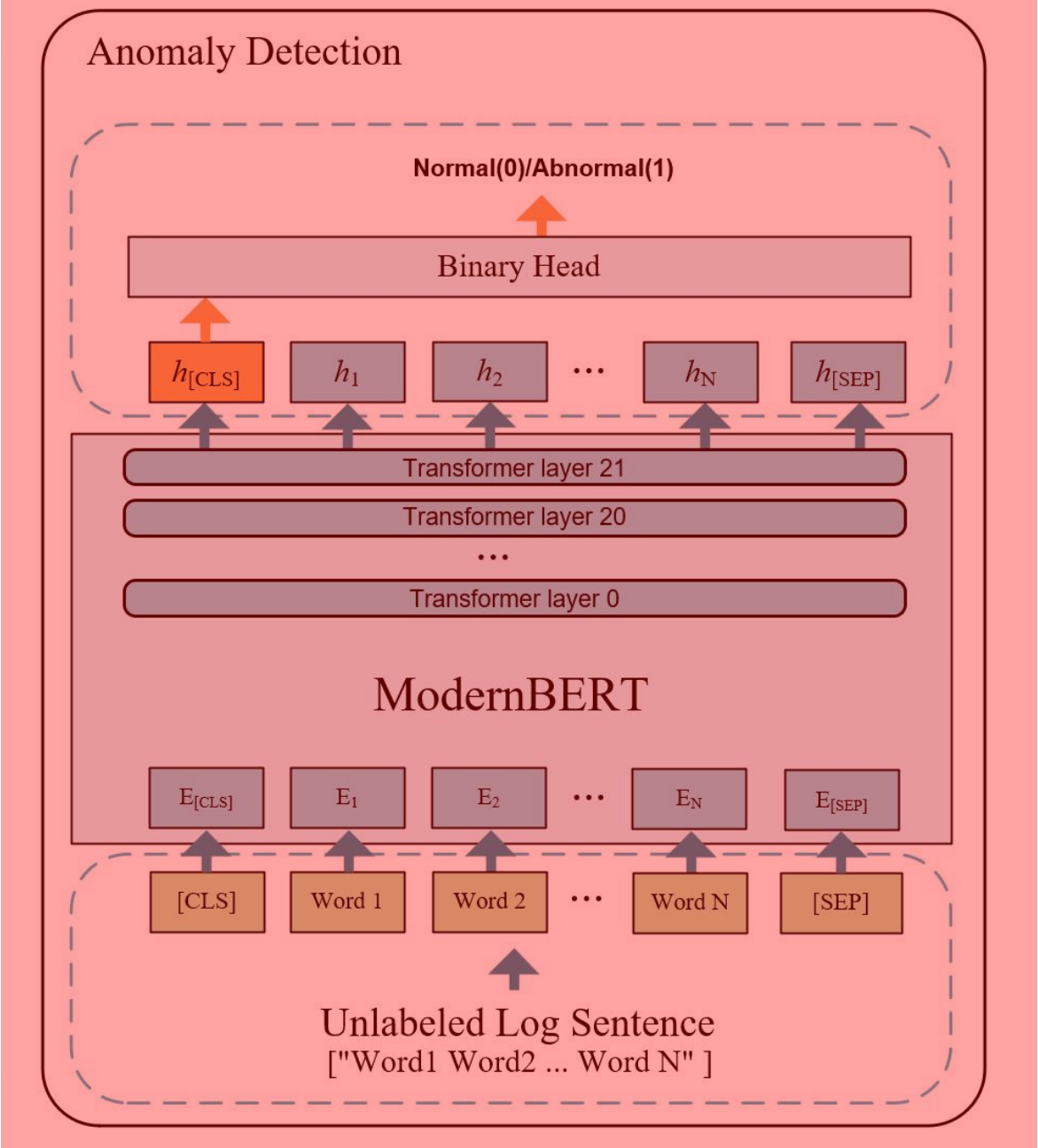
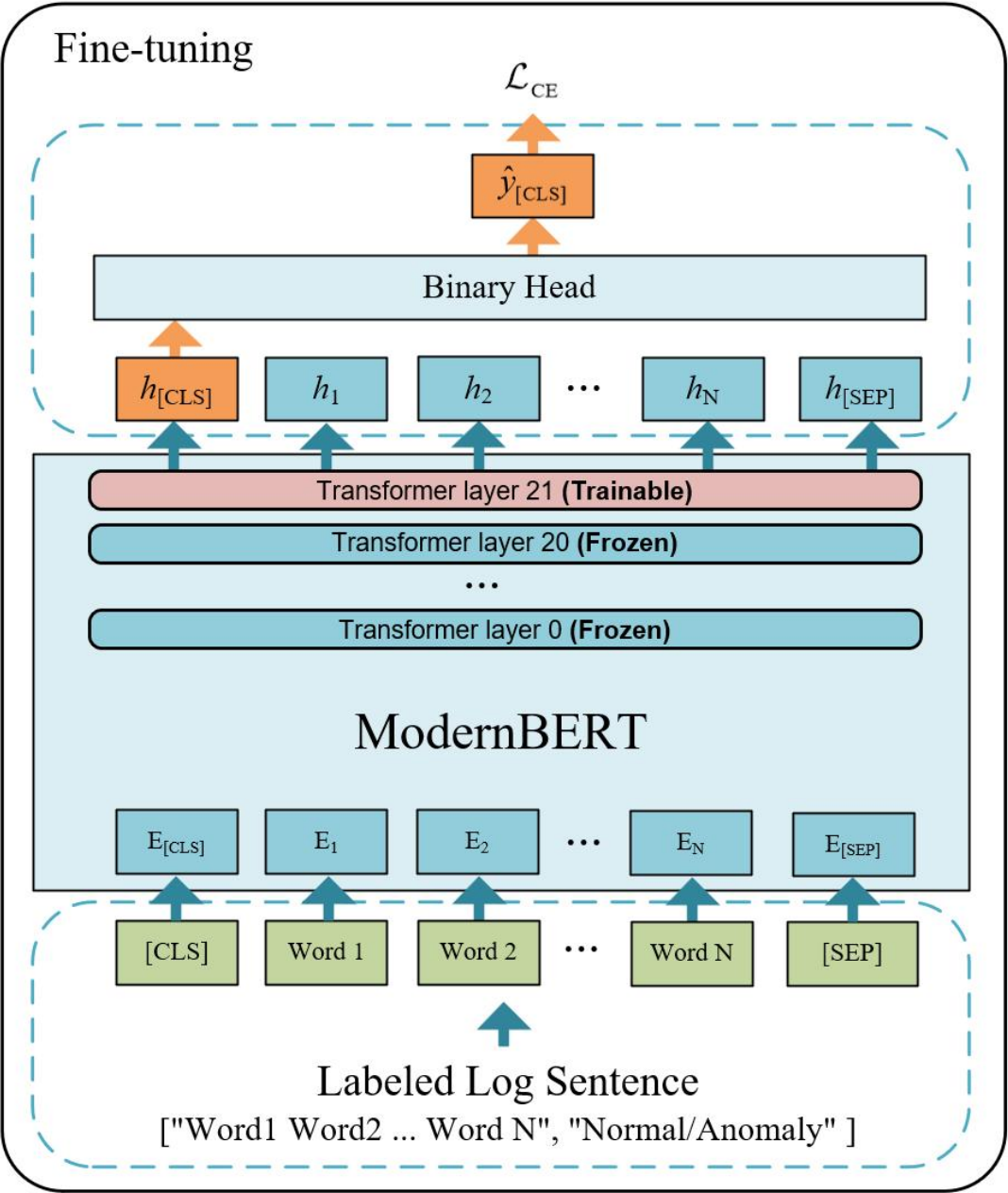
- Inputs $h_{[\text{CLS}]}$: 1×768 ($1 \times \text{hidden_size}$)
- Weights W : 768×2 ($\text{hidden_size} \times \text{label_size}$)
- Output $\hat{y}_{[\text{CLS}]}$: 2×1 ($\text{label_size} \times 1$)

Loss (Cross-Entropy):

$$\mathcal{L}_{\text{CE}} = \frac{1}{N} \sum_{j=1}^N y_{[\text{CLS}]}^j \log \hat{y}_{[\text{CLS}]}^j$$



Our Architecture 1: ModernLogBERT - BC (Binary Classification)



ModernLogBERT - BC Anomaly Detection

Binary classification head:

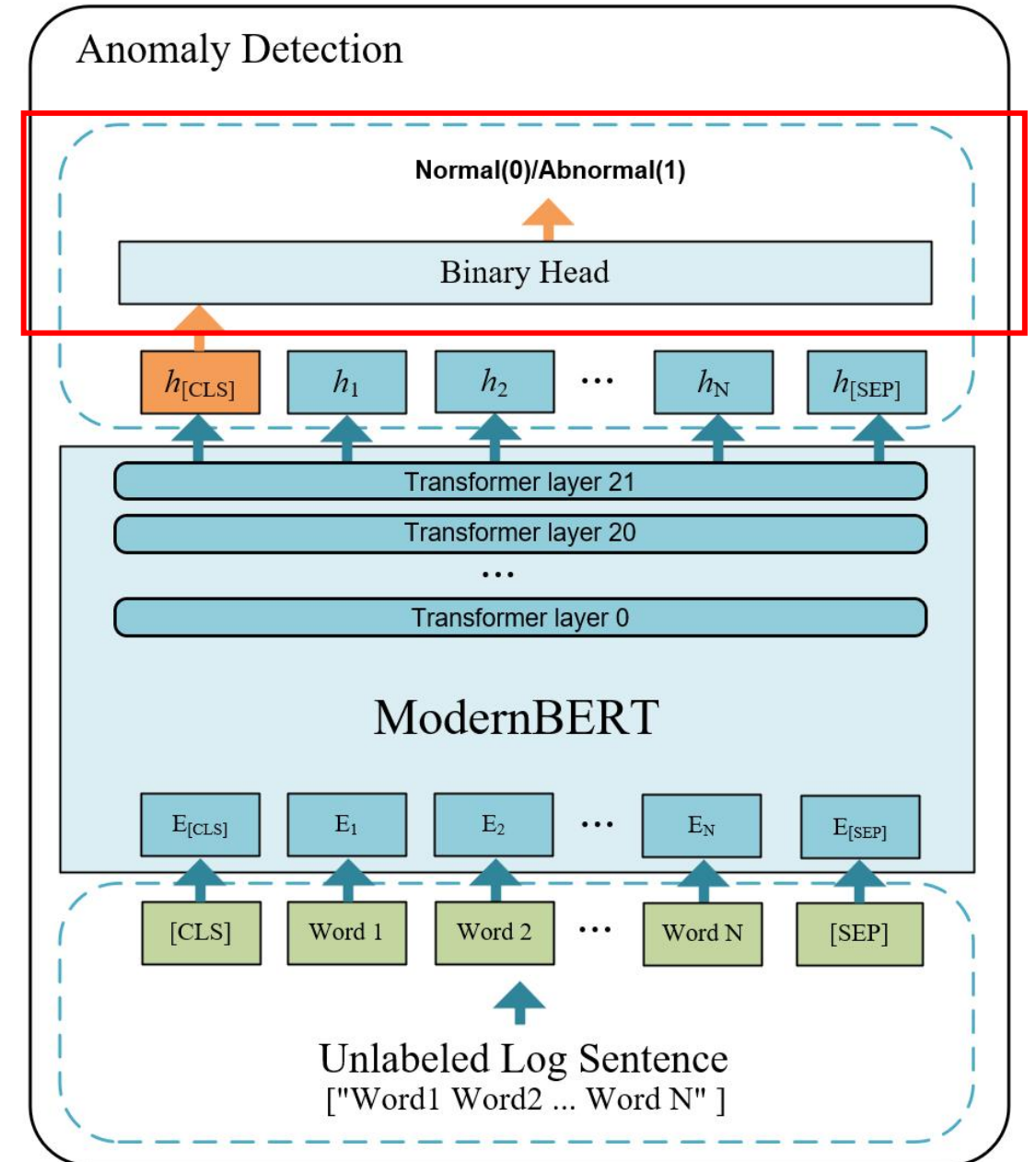
$$\hat{y}_{[\text{CLS}]} = \text{Softmax}(Wh_{[\text{CLS}]} + b)$$

- Output $\hat{y}_{[\text{CLS}]}$: 2×1 (label_size $\times$ 1)

$$\hat{y}_{[\text{CLS}]} = [\hat{y}_{[0]}, \hat{y}_{[1]}]$$

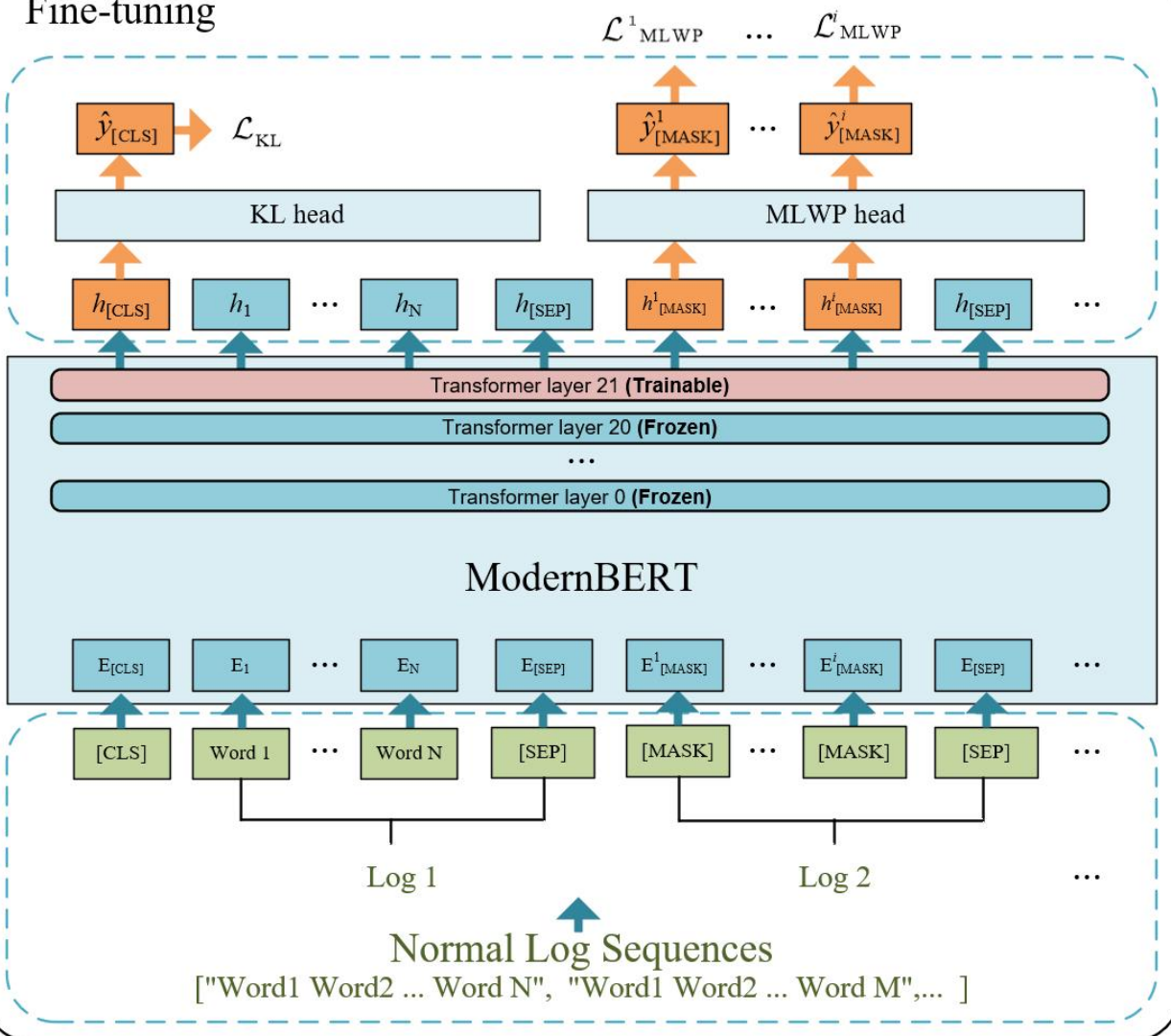
if $\hat{y}_{[0]} \geq \hat{y}_{[1]}$: Normal

if $\hat{y}_{[0]} < \hat{y}_{[1]}$: Anomalous

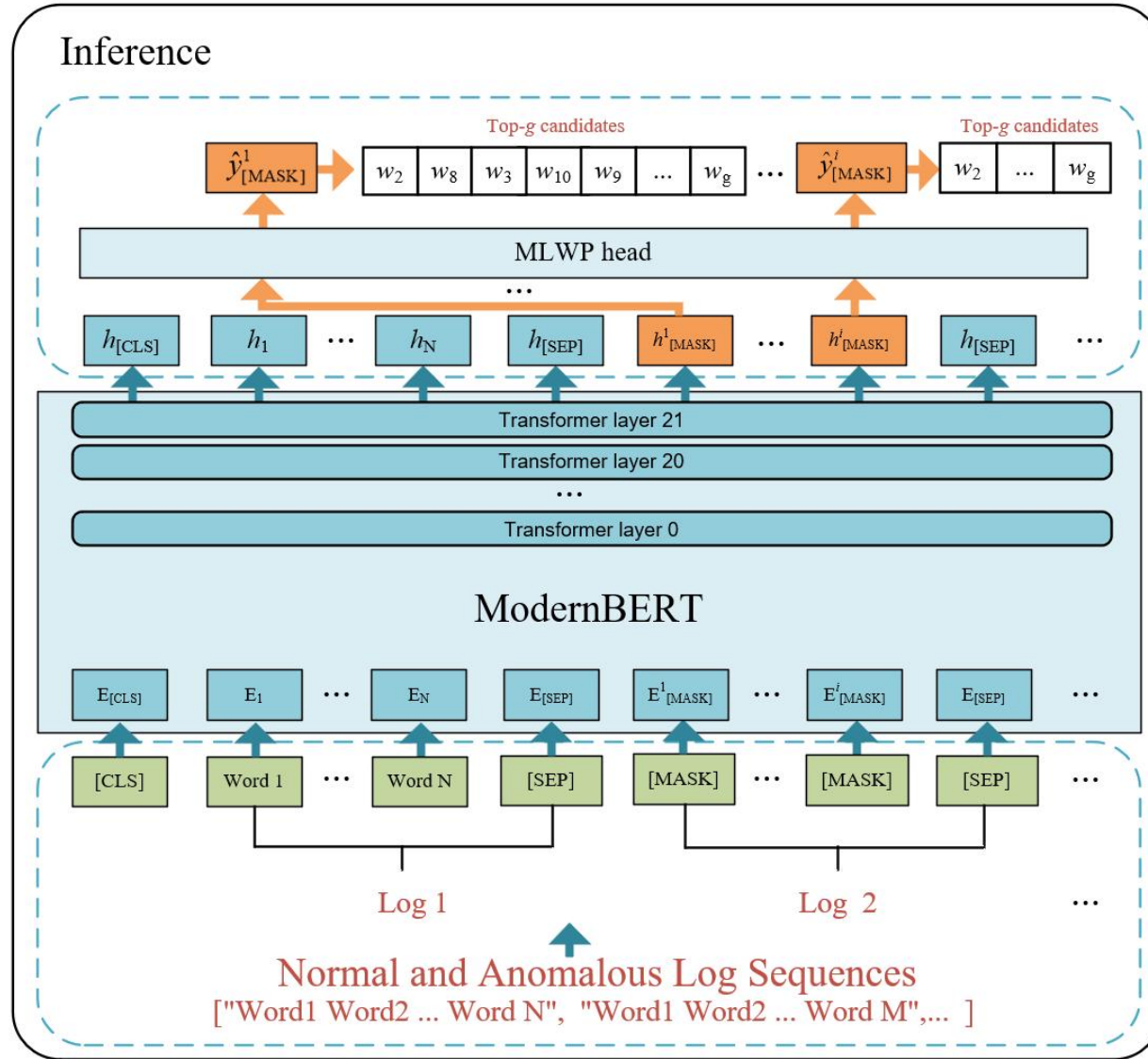


Our Architecture 2: ModernLogBERT - MLWP (Masked Log Word Prediction)

Fine-tuning

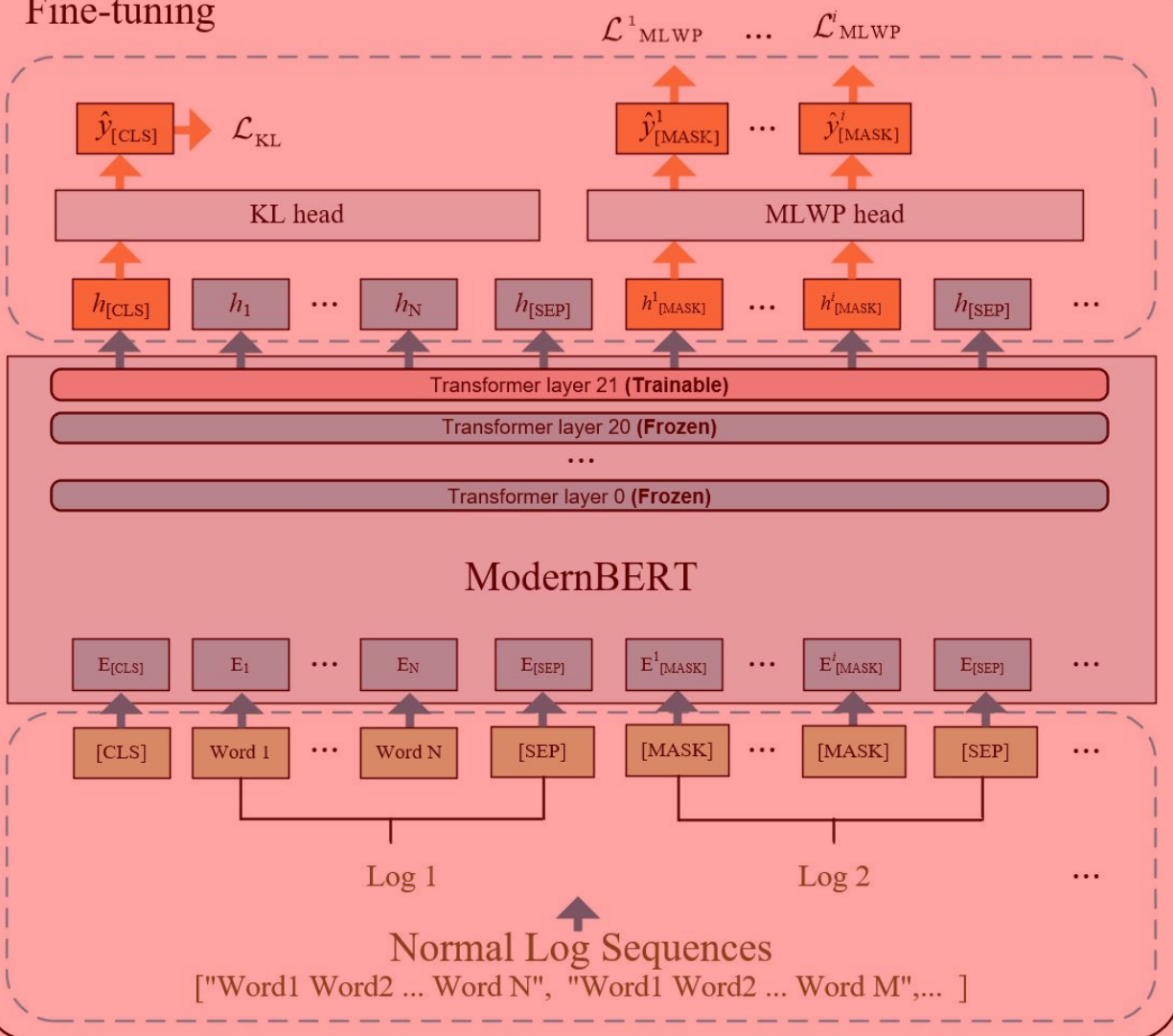


Inference

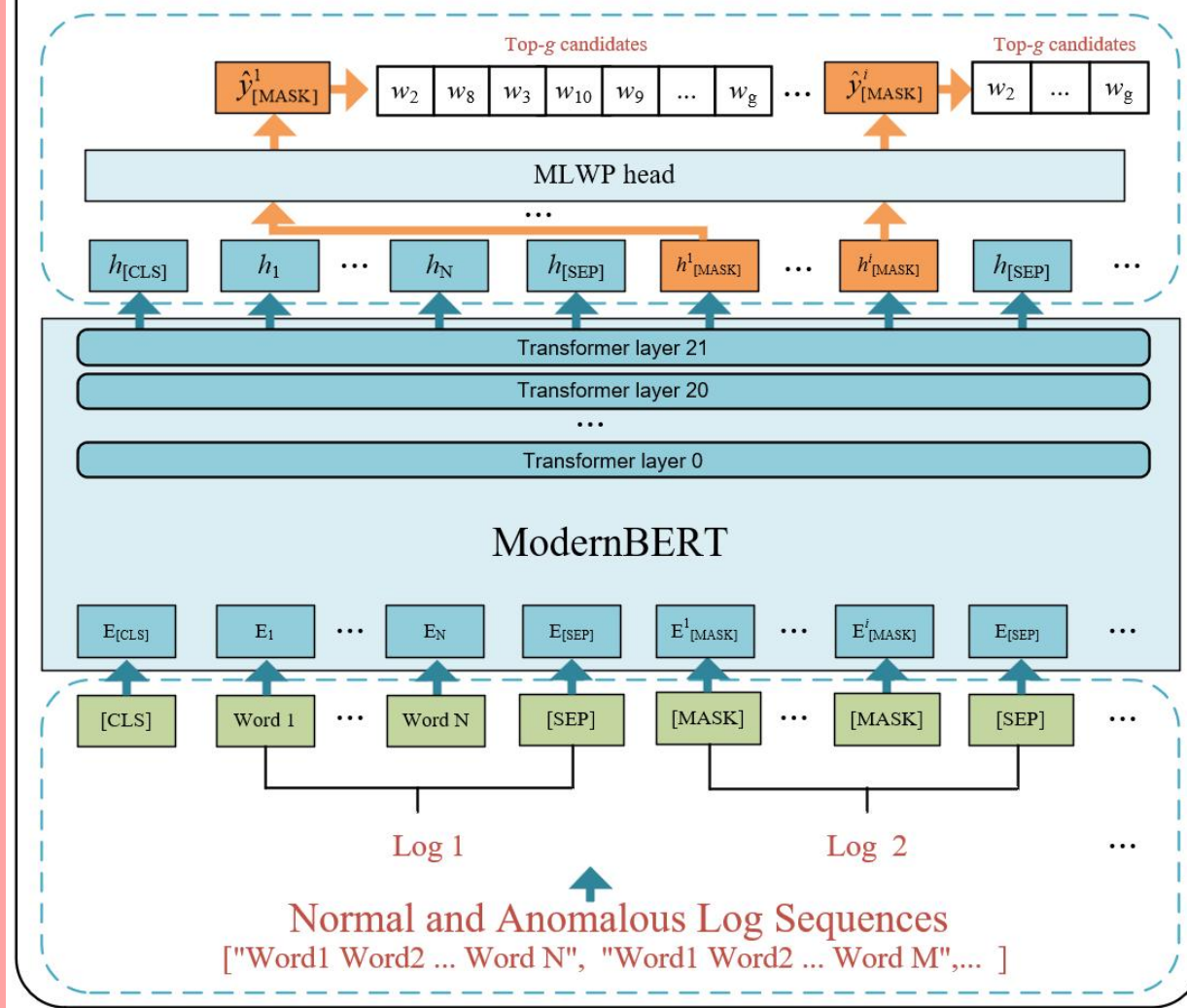


Our Architecture 2: ModernLogBERT - MLWP (Masked Log Word Prediction)

Fine-tuning



Inference

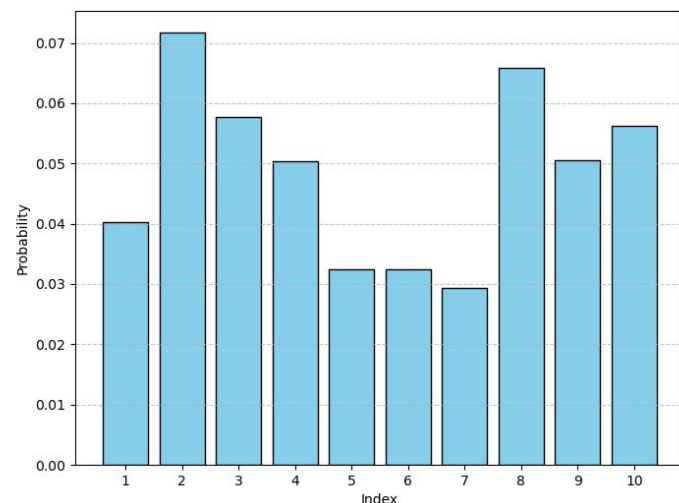


ModernLogBERT - MLWP Task1: Masked Log Word Prediction

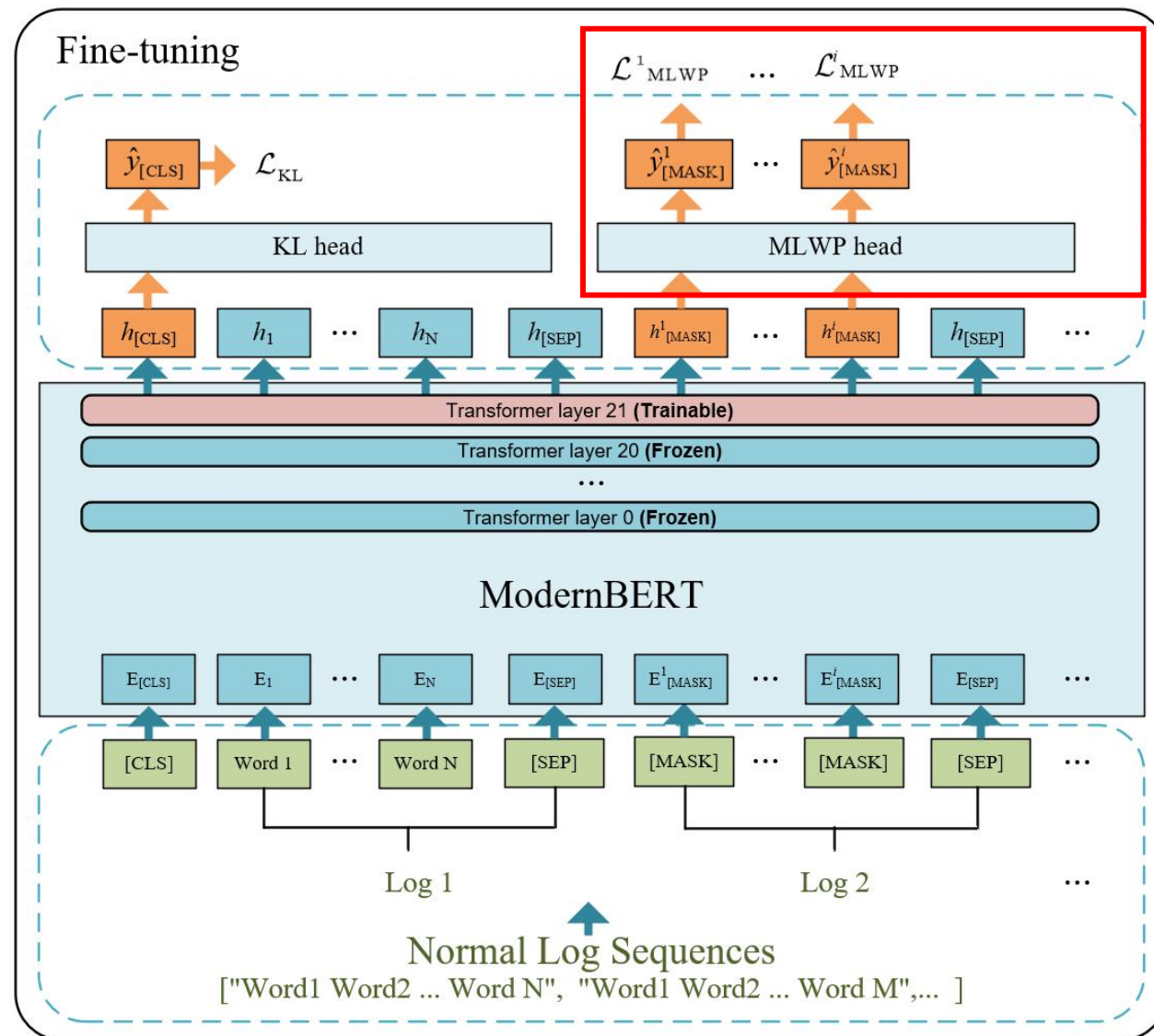
MLWP head:

$$\hat{y}_{\text{MASK}_i} = \text{Softmax}(W h_{\text{MASK}_i} + b)$$

- Inputs h_{MASK_i} : $1 \times \text{hidden_size}$
- Weights W : $\text{hidden_size} \times \text{vocab_size}$
- Output $\hat{y}_{\text{MASK}_i}$: $\text{vocab_size} \times 1$



Loss (Cross-Entropy): $\mathcal{L}_{\text{MLKP}} = -\sum_{i=1}^M y_{\text{MASK}_i} \log \hat{y}_{\text{MASK}_i}$



ModernLogBERT - MLWP Task2: Kullback-Leibler Divergence Minimization

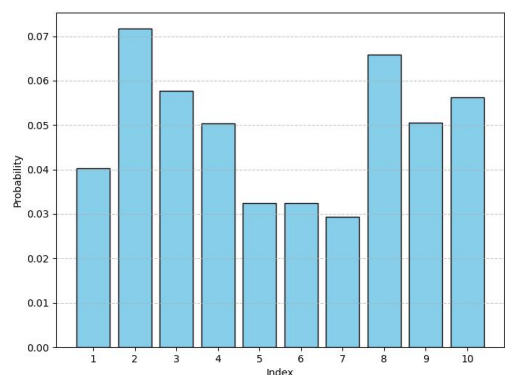
KL head :

$$\hat{y}_{[\text{CLS}]} = \text{Softmax}(Wh_{[\text{CLS}]} + b)$$

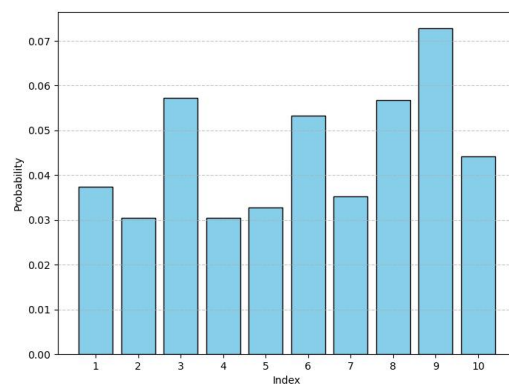
- Inputs $h_{[\text{CLS}]}$: $1 \times \text{hidden_size}$
- Weights W : $\text{hidden_size} \times \text{vocab_size}$
- Output $\hat{y}_{[\text{CLS}]}$: $\text{vocab_size} \times 1$

KL Div Loss:

$$\mathcal{L}_{\text{KL}} = \frac{1}{N} \sum_{j=1}^N y_{[\text{CLS}]}^j (\log y_{[\text{CLS}]}^j - \hat{y}_{[\text{CLS}]}^j)$$



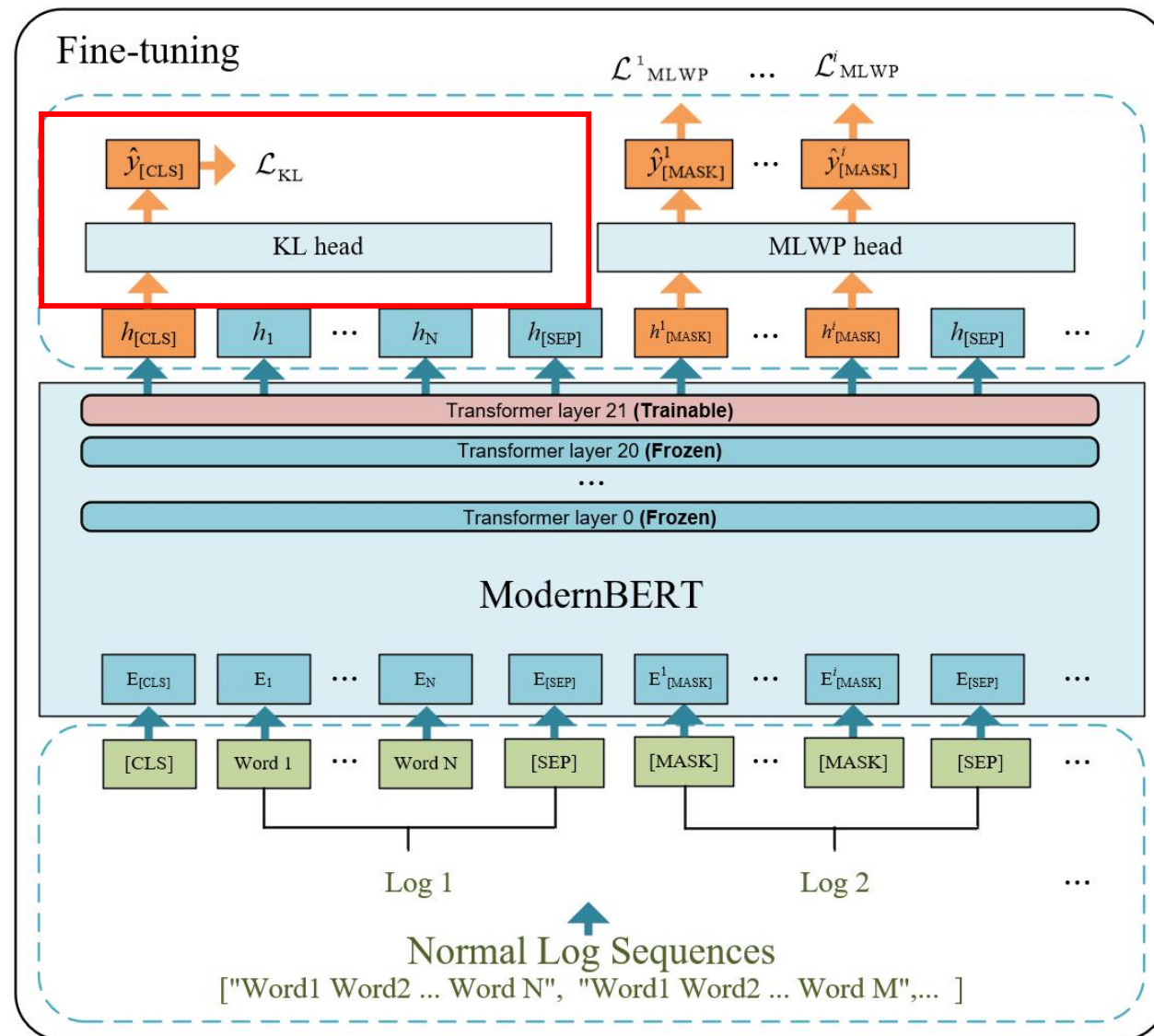
Prediction Distribution



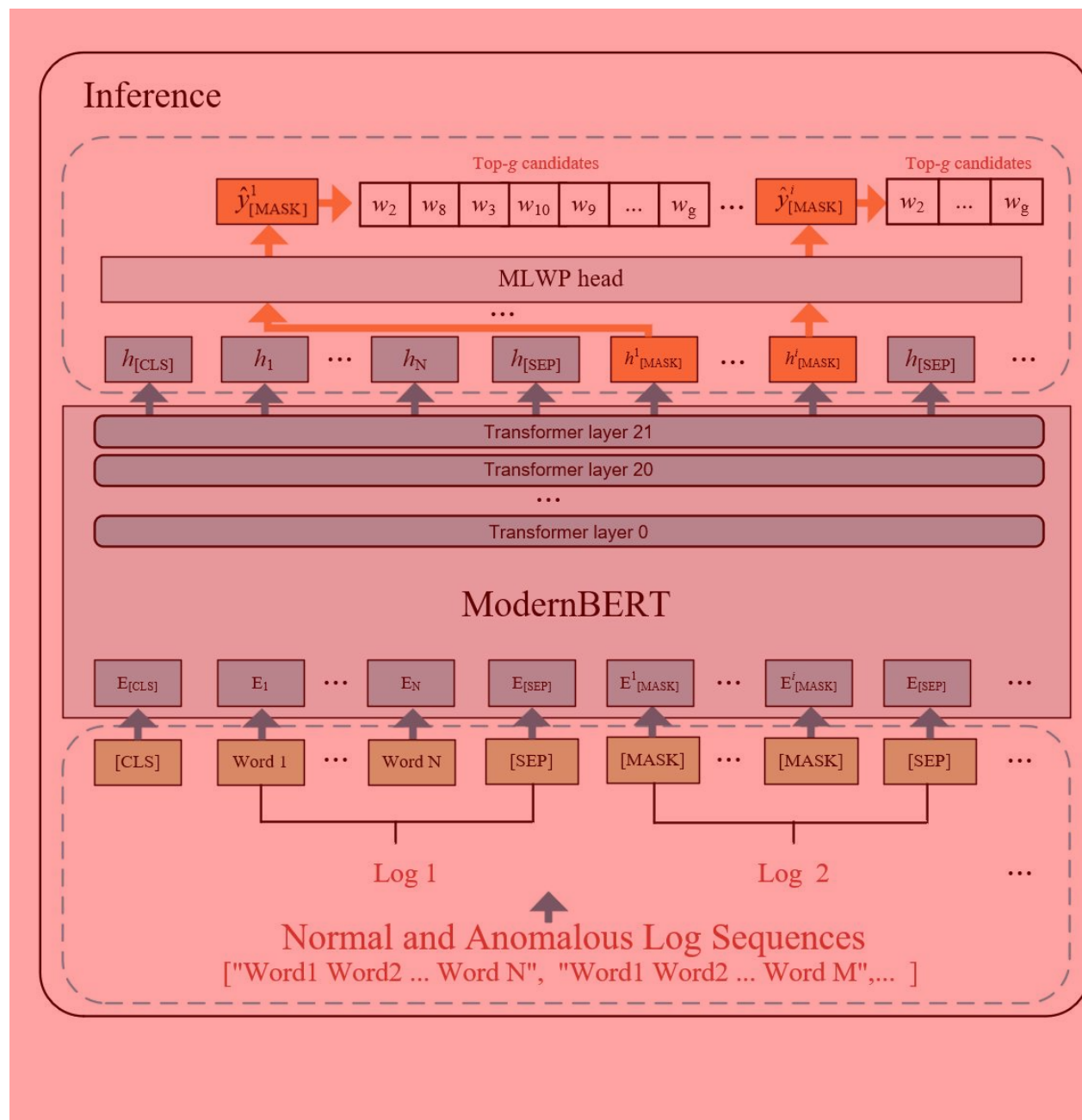
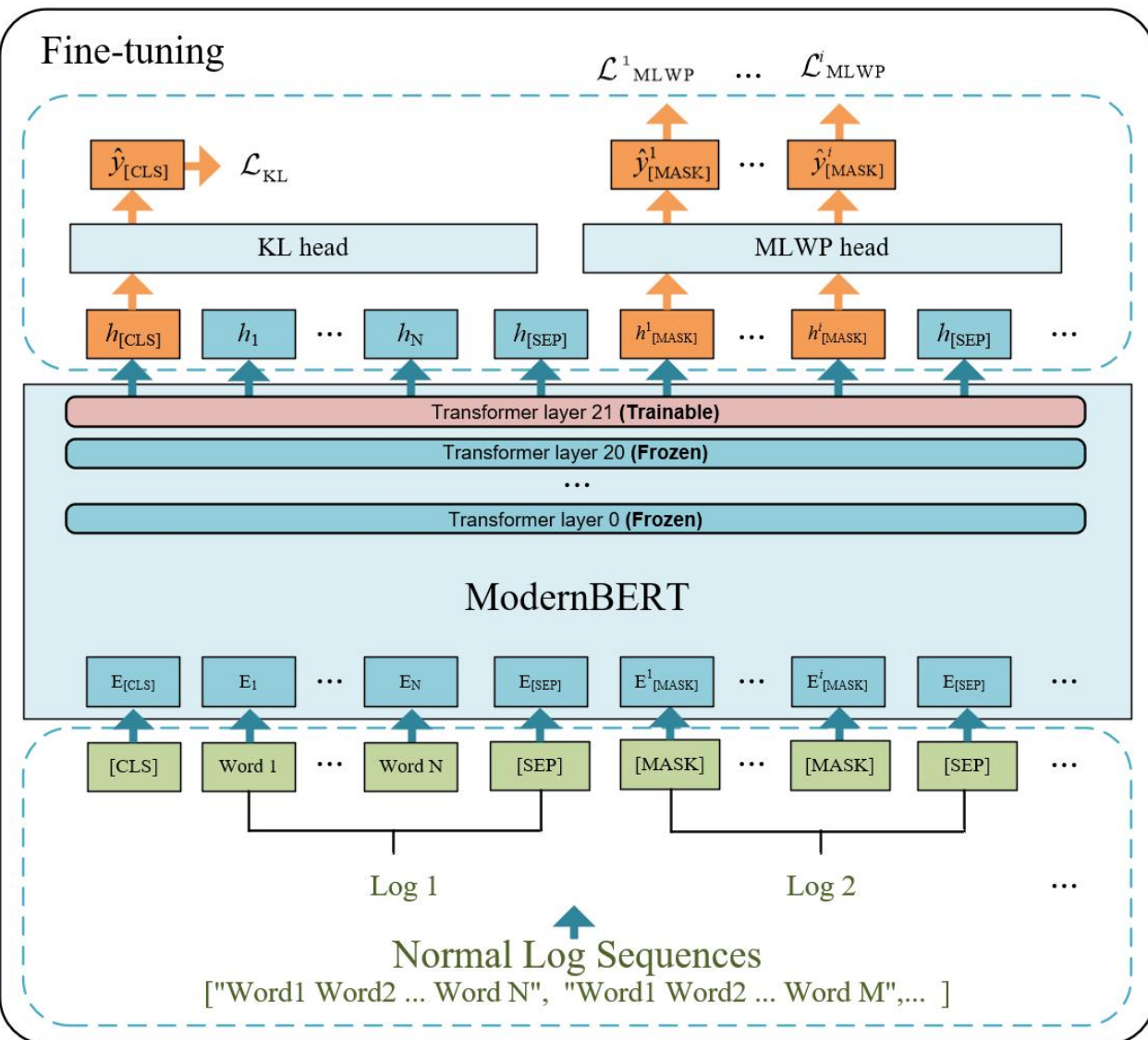
Real Distribution

Backpropagation:

$$\mathcal{L} = \mathcal{L}_{\text{MLWP}} + \alpha \mathcal{L}_{\text{KL}}$$



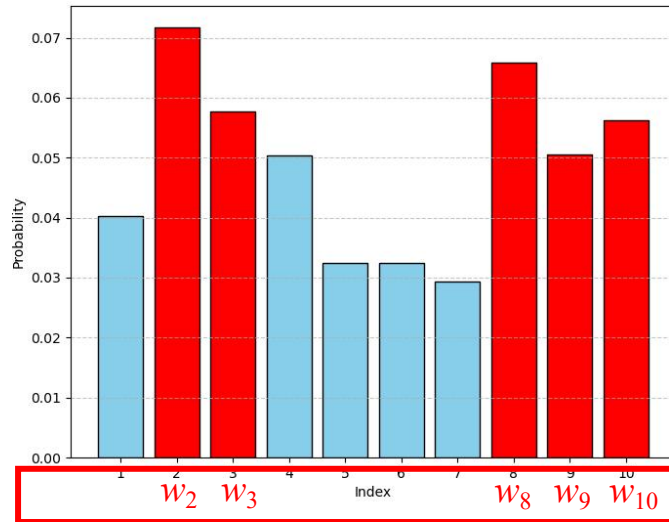
Our Architecture 2: ModernLogBERT - MLWP (Masked Log Word Prediction)



ModernLogBERT - MLWP Anomaly Detection

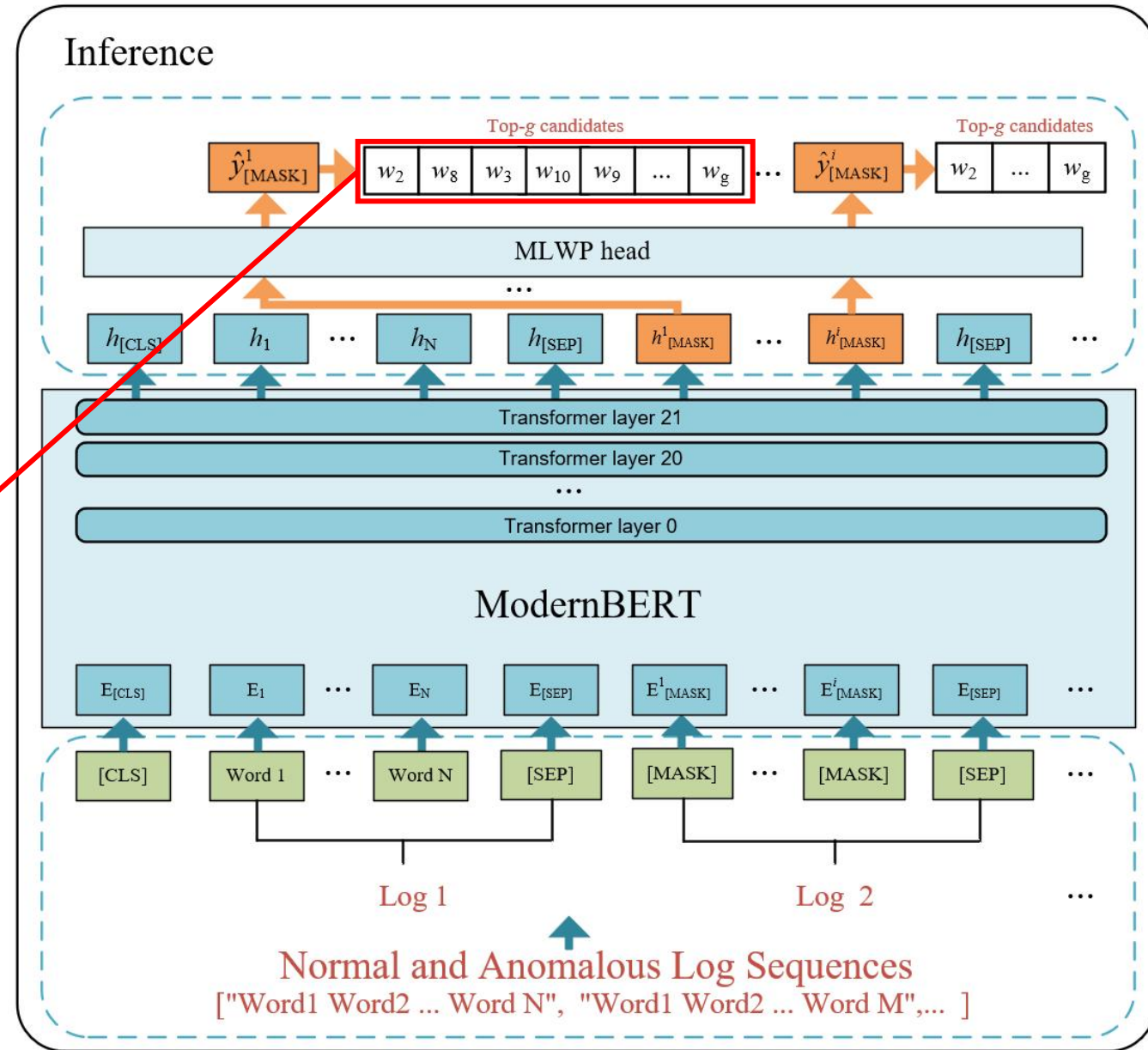
$$\hat{y}_{\text{MASK}_i}^j = \text{Softmax}(W h_{\text{MASK}_i}^j + b)$$

Probability for each token:



Threshold g and r :

- g : Determine whether the true **word** is in the top- g highest likelihoods candidate set.
- r : If there are more than r **words** not in the candidate set, the sequence will be abnormal



Dataset

HDFS_v1 Dataset (<https://github.com/logpai/loghub>).

- Source: Hadoop Distributed File System logs, generated from MapReduce tasks on Amazon EC2.
- Size: 11,175,629 log entries, 1.47 GB.
- Structure: Semi-structured logs
- Annotations: Expert-labeled

Raw log

```
081109 203518 143 INFO dfs.DataNode$DataXceiver: Receiving block blk_-
1608999687919862906 src: /10.250.19.102:54106 dest:
/10.250.19.102:50010

081109 203518 35 INFO dfs.FSNamesystem: BLOCK*
NameSystem.allocateBlock:
/mnt/hadoop/mapred/system/job_200811092030_0001/job.jar. blk_-
1608999687919862906

081109 203519 143 INFO dfs.DataNode$DataXceiver: Receiving block blk_-
1608999687919862906 src: /10.250.10.6:40524 dest: /10.250.10.6:50010

081109 203519 145 INFO dfs.DataNode$DataXceiver: Receiving block blk_-
1608999687919862906 src: /10.250.14.224:42420 dest:
/10.250.14.224:50010

.....
```

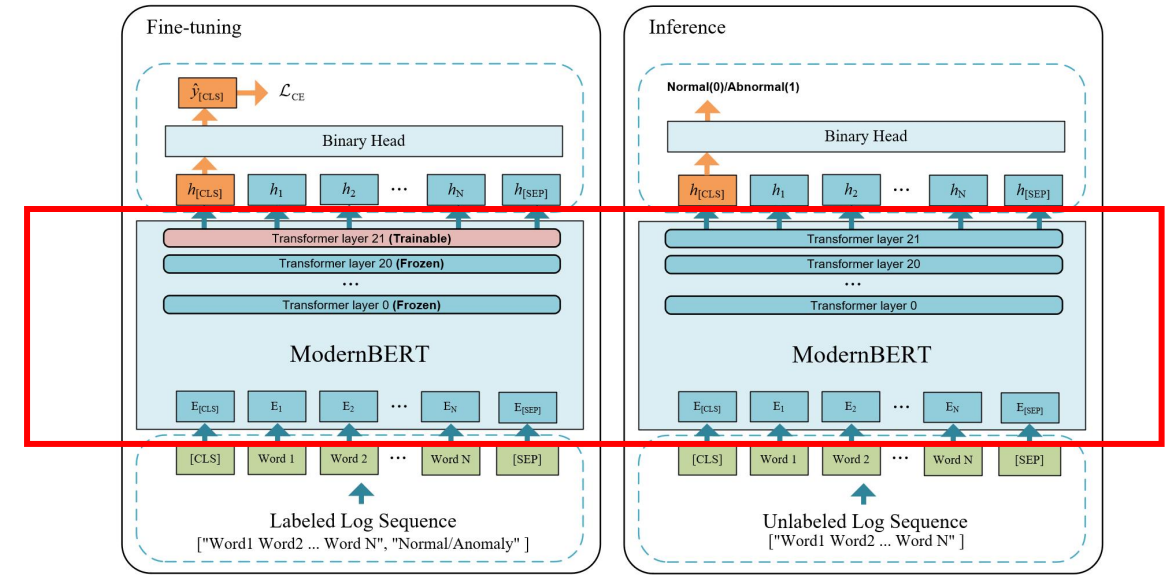
Template

EventId (key)	EventTemplate
E1	[*]Adding an already existing block[*]
E2	[*]Verification succeeded for[*]
...	...
E5	[*]Receiving block[*]src:[*]dest:[*]
...	...
E28	[*]BLOCK* NameSystem[*]addStoredBlock:
E29	[*]PendingReplicationMonitor timed out block[*]

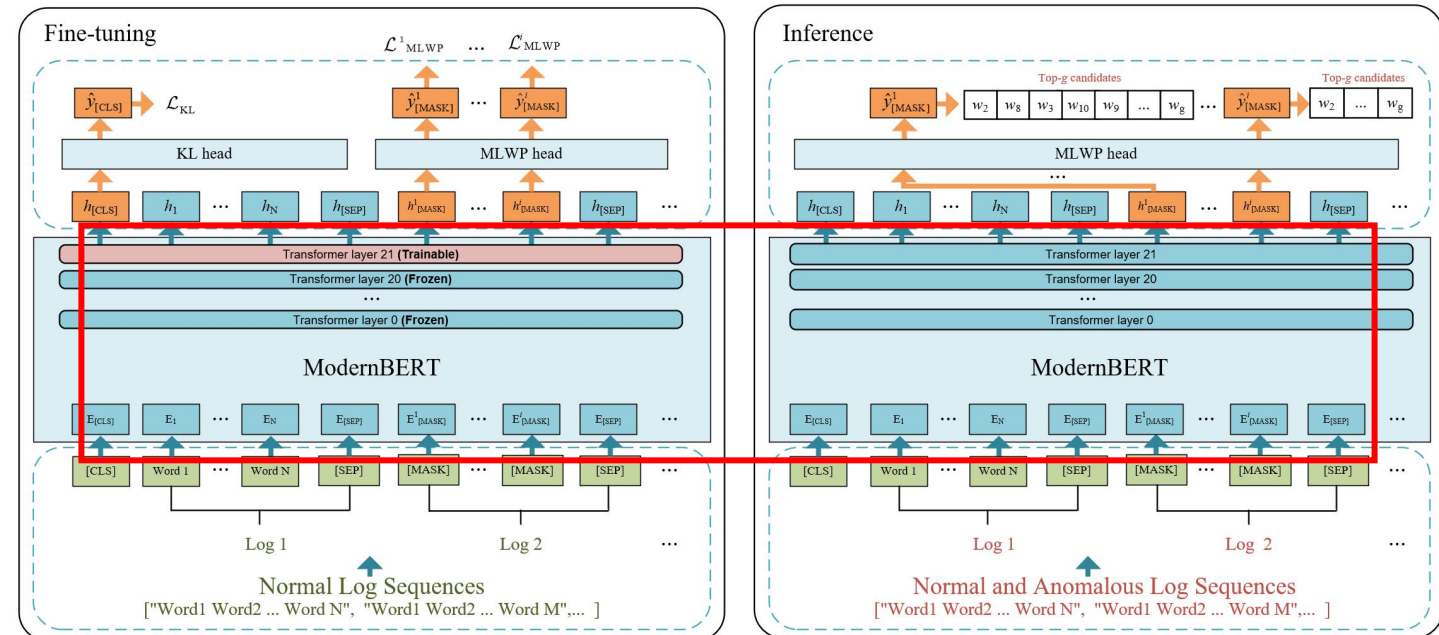
Our Architecture 1: ModernLogBERT - BC

Pretrained ModernBERT Models

- Experiment 1:
A Pretrained ModernBERT-base Model
- Experiment 2:
Pretrain a logtokenizer + ModernBERT-base Model



Our Architecture 2: ModernLogBERT - MLWP



Experiments — Parameter Settings

Table 1: Parameter Settings for Log Anomaly Detection Models

Model	Mask Ratio	Batch Size	Frozen Layers	Epochs	lr	α	g	r
LogBERT	0.15	32	-	200	2e-5	1.0	5	2
ModernLogBERT - BC	-	32	21	3	2e-5	-	-	-
ModernLogBERT - BC (Logtokenizer)	-	32	21	3	2e-5	-	-	-
ModernLogBERT - MLWP	0.3	32	21	3	2e-5	1.0	5	20
ModernLogBERT - MLWP (Logtokenizer)	0.3	32	21	3	2e-5	1.0	5	20

Table 2: Performance Comparison of Log Anomaly Detection Models

Model	Accuracy	Precision	Recall	Macro F1
LogBERT	87.02	-	78.10	82.32
ModernLogBERT - BC	99.43	81.11	97.77	88.66
ModernLogBERT - BC (Logtokenizer)	98.63	72.53	85.75	78.59
ModernLogBERT - MLWP	97.83	79.37	30.96	71.72
ModernLogBERT - MLWP (Logtokenizer)	99.10	85.71	85.47	92.57

Conclusion

- It is effective to use ModernBert to extract the textual semantics in the log information.
- Training a tokenizer for logs, combined with the method of pre-training +Finetuning ModernBERT, has great potential in the task of log anomaly detection

Future work

- Ablation experiment
- Hyperparameter optimization
- Determine the optimal number of frozen layers
- Verify the generalization performance on other datasets



LUND
UNIVERSITY