

Context-Aware Learning to Rank with Self-Attention

Anonymous Author(s)

ABSTRACT

In learning to rank, one is interested in optimising the global ordering of a list of items according to their utility for users. Popular approaches learn a scoring function that scores items individually (i.e. without the context of other items in the list) by optimising a pointwise, pairwise or listwise loss. The list is then sorted in the descending order of the scores. Possible interactions between items present in the same list are taken into account in the training phase at the loss level. However, during inference, items are scored individually, and possible interactions between them are not considered. In this paper, we propose a context-aware neural network model that learns item scores by applying a self-attention mechanism. The relevance of a given item is thus determined in the context of all other items present in the list, both in training and in inference. Finally, we empirically demonstrate significant performance gains of self-attention based neural architecture over Multi-Layer Perceptron baselines. This effect is consistent across popular pointwise, pairwise and listwise losses on datasets with both implicit and explicit relevance feedback.

CCS CONCEPTS

• **Information systems** → **Learning to rank**; *Information retrieval*;

KEYWORDS

learning to rank, self-attention, context-aware ranking

ACM Reference Format:

Anonymous Author(s). 2019. Context-Aware Learning to Rank with Self-Attention. In *Proceedings of The Web Conference*. ACM, New York, NY, USA, 7 pages. <https://doi.org/10.1145/nnnnnnn.nnnnnnn>

1 INTRODUCTION

Learning to rank (LTR) is an important area of machine learning research, lying at the core of many information retrieval (IR) systems. It arises in numerous industrial applications like search engines, recommender systems, question-answering systems, and others.

A typical machine learning solution to the LTR problem involves learning a scoring function, which assigns real-valued scores to each item of a given list, based on a dataset of item features and human-curated or implicit (e.g. clickthrough logs) relevance labels. Items are then sorted in the descending order of scores [19]. Performance of the trained scoring function is usually evaluated using an

IR metric like Mean Reciprocal Rank (MRR) [29], Normalised Discounted Cumulative Gain (NDCG) [16] or Mean Average Precision (MAP) [4].

In contrast to other classic machine learning problems like classification or regression, the main goal of a ranking algorithm is to determine relative preference among a group of items. Scoring items individually is a proxy of the actual learning to rank task. Users' preference for a given item on a list depends on other items present in the same list: an otherwise preferable item might become less relevant in the presence of other, more relevant items. Common learning to rank algorithms attempt to model such inter-item dependencies at the loss level. That is, items in a list are still scored individually, but the effect of their interactions on evaluation metrics is accounted for in the loss function, which usually takes a form of a pairwise (RankNet [6], LambdaLoss [30]) or a listwise (ListNet [9], ListMLE [31]) objective. For example, in LambdaMART [8] the gradient of the pairwise loss is rescaled by the change in NDCG of the list which would occur if a pair of items was swapped. Pointwise objectives, on the other hand, do not take such dependencies into account.

In this work, we propose a learnable, context-aware, self-attention [27] based scoring function, which allows for modelling of inter-item dependencies not only at the loss level but also in the computation of items' scores. Self-attention is a mechanism first introduced in the context of natural language processing. Unlike RNNs [14], it does not process the input items sequentially but allows the model to attend to different parts of the input regardless of their distance from the currently processed item. We adapt the Transformer [27], a popular self-attention based neural machine translation architecture, to the ranking task. We demonstrate that the obtained ranking model significantly improves performance over Multi-Layer Perceptron (MLP) baselines across a range of pointwise, pairwise and listwise ranking losses. Evaluation is conducted on MSLR-WEB30K [24], the benchmark LTR dataset with multi-level relevance judgments, as well as on clickthrough data coming from Allegro.pl, a large-scale e-commerce search engine.

We provide an open-source Pytorch [22] implementation of our self-attentive context-aware ranker available at `url_removed`.

The rest of the paper is organised as follows. In Section 2 we review related work. In Section 3 we formulate the problem solved in this work. In Section 4 we describe our self-attentive ranking model. Experimental results and their discussion are presented in Section 5. In Section 6 we conduct an ablation study of various hyperparameters of our model. Finally, a summary of our work is given in Section 7.

2 RELATED WORK

Learning to rank has been extensively studied and there is a plethora of resources available on classic pointwise, pairwise and listwise approaches. We refer the reader to [19] for the overview of the most popular methods.

Permission to make digital or hard copies of part or all of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for third-party components of this work must be honored. For all other uses, contact the owner/author(s).
The Web Conference, April, 2020, Taipei, Taiwan
© 2019 Copyright held by the owner/author(s).
ACM ISBN 978-x-xxxx-xxxx-x/YY/MM.
<https://doi.org/10.1145/nnnnnnn.nnnnnnn>

What the majority of LTR methods have in common is that their scoring functions score items individually. Inter-item dependencies are (if at all) taken into account at the loss level only. Previous attempts at modelling context of other items in a list in the scoring function include:

- a pairwise scoring function [12] and Groupwise Scoring Function (GSF) [2], which incorporates the former work as its special case. However, the proposed GSF method simply concatenates feature vectors of multiple items and passes them through an MLP. To desensitize the model to the order of concatenated items, Monte-Carlo sampling is used, which yields an unscalable algorithm,
- a seq2slate model [5] uses an RNN combined with a variant of Pointer Networks [28] in an encoder-decoder type architecture to both encode items in a context-aware fashion and then produce the optimal list by selecting items one-by-one. Authors evaluate their approach only on clickthrough data (both real and simulated from WEB30K). A similar, simpler approach known as Deep Listwise Context Model (DLCM) was proposed in [1]: an RNN is used to encode a set of items for re-ranking, followed by a single decoding step with attention,
- in [15], authors attempt to capture inter-item dependencies by adding so-called *delta* features which represent how different given item is from items surrounding it in the list. It can be seen as a simplified version of a local self-attention mechanism. Authors evaluate their approach on proprietary search logs only,
- authors of [17] formulate the problem of re-ranking of a list of items as that of a whole-list generation. They introduce ListCVAE, a variant of Conditional Variational Auto-Encoder [25] which learns the joint distribution of items in a list conditioned on users' relevance feedback and uses it to directly generate a ranked list of items. Authors claim NDCG unfairly favours greedy ranking methods and thus do not use that metric in their evaluation,
- similarly to our approach, Pei et al. [23] use the self-attention mechanism to model inter-item dependencies. Their approach, however, was not evaluated on a standard WEB30K dataset and the only loss functions considered was ListNet.

Our proposed solution to the problem of context-aware ranking makes use of the self-attention mechanism. It was first introduced as intra-attention in [11] and received more attention after the introduction of the Transformer architecture [27]. Our model can be seen as a special case of the encoder part of the Transformer.

We compare the proposed approach with those of the aforementioned methods which provided an evaluation on WEB30K in terms of NDCG@5. These include GSF of [2] and DLCM of [1]. We outperform both competing methods.

3 PROBLEM FORMULATION

In this section, we formulate problem at hand in learning to rank setting. Let X be the training set. It consists of pairs $(\mathbf{x}, \mathbf{y})$ of a list $\mathbf{x}$ of d_f -dimensional real-valued vectors x_i together with a list $\mathbf{y}$ of their relevance labels y_i (multi-level or binary). Note that lists $\mathbf{x}$ in the training set may be of varying length. The goal is to find

a scoring function f which maximises an IR metric of choice (e.g. NDCG) on the test set. Since IR metrics are rank based (thus, non-differentiable), the scoring function f is trained to minimise the average of a surrogate loss l over the training data.

$$\mathcal{L}(f) = \frac{1}{|X|} \sum_{(\mathbf{x}, \mathbf{y}) \in X} l((\mathbf{x}, \mathbf{y}), f),$$

while controlling for overfitting (e.g. by using dropout [26] in the neural network based scoring function f or adding L_1 or L_2 penalty term [20] to the loss function l). Thus, two crucial choices one needs to make when proposing a learning to rank algorithm are that of a scoring function f and loss function l . As discussed earlier, typically, f scores elements $x_i \in \mathbf{x}$ individually to produce scores $f(x_i)$, which are then input to loss function l together with ground truth labels y_i . In subsequent sections, we describe our construction of context-aware scoring function f which is able to model interactions between items x_i in a list $\mathbf{x}$. Our model is generic enough to be applicable with any of standard pointwise, pairwise or listwise loss. We thus experiment with a variety of popular ranking losses l .

4 SELF-ATTENTIVE RANKER

In this section, we describe the architecture of our self-attention based ranking model. We modify the Transformer architecture to work in the ranking setting and obtain a scoring function which, when scoring a single item, takes into account all other items present in the same list.

4.1 Self-Attention Mechanism

The key component of our model is the self-attention mechanism introduced in [27]. The attention mechanism can be described as taking the query vector and pairs of key and value vectors as input and producing a vector output. The output of the attention mechanism for a given query is a weighted sum of the value vectors, where weights represent how relevant to the query is the key of the corresponding value vector. Self-attention is a variant of attention in which query, key and value vectors are all the same - in our case, they are vector representations of items in the list. The goal of the self-attention mechanism is to compute a new, higher-level representation for each item in a list, by taking a weighted sum over all items in a list according to weights representing the relevance of these items to the query item.

There are many ways in which one may compute the relevance of key vectors to query vectors. We use the variant of self-attention known as *Scaled Dot-Product Attention*. Suppose Q is a d_{model} -dimensional matrix representing all items (queries) in the list. Let K and V be the keys and values matrices, respectively. Then

$$\text{Attention}(Q, K, V) = \text{softmax}\left(\frac{QK^T}{\sqrt{d_{\text{model}}}}\right)V.$$

The scaling factor of $\frac{1}{\sqrt{d_{\text{model}}}}$ is added to avoid small gradients in the softmax operation for large values of d_{model} .

4.2 Multi-Headed Self-Attention

As described in [27], it is beneficial to perform the self-attention operation multiple times and concatenate the outputs. To avoid growing the size of the resulting output vector, matrices Q , K and V

are first linearly projected H times to d_q , d_k and d_v dimensional spaces, respectively. Usually, $d_q = d_k = d_v = d_{\text{model}}/H$. Each of H computations of linear projection of Q , K , V , followed by self-attention mechanism is referred to as a single attention head. Note that each head has its own learnable projection matrices. The outputs of each head are concatenated and once again linearly projected, usually to the vector space of the same dimension as that of input matrix Q . Similarly to the Transformer, our model also uses multiple attention heads.

4.3 Positional Encodings

Transformer architecture was designed to solve a neural machine translation (NMT) task. In NMT, the order of input tokens should be taken into account. Unlike RNNs, self-attention based encoder has no way of discerning the order of input tokens. Authors of the original Transformer paper proposed to solve the problem by using either fixed or learnable positional encodings.

The ranking problem can be viewed as either ordering a set of (unordered) items or as re-ranking, where the input list has already been sorted according to a weak ranking model. In the former case, the use of positional encodings is not needed. In the latter, they may boost model's performance. We experiment with both ranking and re-ranking settings and when positional encodings are used, we test the fixed encodings variant ¹. Details can be found in Section 5.

4.4 Model Architecture

We adapt the Transformer model to the ranking setting as follows. Items on a list are treated as tokens and item features as input token embeddings. We denote the length of an input list as l and the number of features as d_f . Each item is first passed through a shared fully connected layer of size d_{fc} . Next, hidden representations are passed through an encoder part of Transformer architecture with N encoder blocks, H heads and hidden dimension d_h . Recall that an encoder block in the Transformer consists of a multi-head attention layer with a skip-connection [13] to the input, followed by layer normalisation [3], time-distributed feed-forward layer and another skip connection followed by layer normalisation. Dropout is applied before performing summation in residual blocks. Finally, after N encoder blocks, a fully-connected layer shared across all items in the list is used to compute a score for each item. The model can be seen as an encoder part of the Transformer with extra linear projection on the input. By using self-attention in the encoder, we ensure that in the computation of a score of a given item, hidden representation of all other items were accounted for. Obtained scores, together with ground truth labels, can provide input to any ranking loss of choice. If the loss is a differentiable function of scores (and thus, of model's parameters), one can use SGD to optimise it. We thus obtain a general, context-aware model for scoring items on a list that can readily be used with any differentiable ranking loss.

5 EXPERIMENTS

5.1 Datasets

Learning to rank datasets come in two flavours: they can have either multi-level or binary relevance labels. Usually, multi-level relevance

labels are human-curated, whereas binary labels are derived from clickthrough logs and are considered implicit feedback. We evaluate our context-aware ranker on both types of data.

For the first type, we use the popular WEB30K dataset, which consists of more than 30,000 queries together with lists of associated search results. Every search result is encoded as a 136-dimensional real-valued vector and has associated with it a relevance label on the scale from 0 (irrelevant) to 4 (most relevant). We standardise the features before inputting them into a learning algorithm. The dataset comes partitioned into five folds with roughly the same number of queries per fold. We perform 5-fold cross-validation by training our models on three folds, validating on one and testing on the final fold. All results reported are averages across five folds together with the standard deviation of results. Since lists in the dataset are of unequal length, we pad or subsample to equal length for training, but use full length (i.e. pad to maximum length present in the dataset) for validation and testing.

For a dataset with binary labels, we use clickthrough logs of Allegro.pl, a large scale e-commerce search engine. The search engine already has a ranking model deployed, which is trained using XGBoost [10] with rank:pairwise loss. We thus treat learning on this dataset as a re-ranking problem and use fixed positional encodings in context-aware scoring functions. This lets the models leverage items' positions returned by the base ranker. The search logs consist of 1M lists, each of length at most 60. Nearly all lists (95%) have only one relevant item with label 1; remaining items were not clicked and are deemed irrelevant (label 0). Each item in a list is represented by a 45-dimensional, real-valued vector. We do not perform cross-validation on this set, but we use the usual train, validation and test splits of the data.

5.2 Loss Functions

To evaluate the performance of the proposed context-aware ranking model, we use several popular ranking losses. Pointwise losses used are RMSE of predicted scores and ordinal loss [21] (with minor modification to make it suitable for ranking). For pairwise losses, we use NDCGLoss 2++ (one of the losses of LambdaLoss framework) and its special cases, RankNet and LambdaRank [7]. Listwise losses used consist of ListNet and ListMLE.

Below, we briefly describe all of the losses used. For a more thorough treatment, please refer to the original papers. Throughout, X denotes the training set, $\mathbf{x}$ denotes an input list of items, $\mathbf{s} = f(\mathbf{x})$ is a vector of scores obtained via the ranking function f and $\mathbf{y}$ is the vector of ground truth relevancy labels.

5.2.1 Pointwise RMSE. The simplest baseline is a pointwise loss, in which no interaction between items is taken into account. We use RMSE loss:

$$l(\mathbf{s}, \mathbf{y}) = \sqrt{\sum_i (y_i - s_i)^2}$$

In practice, we used sigmoid activation function on the outputs of the scoring function f and rescaled them by multiplying by maximum relevance value (e.g. 4 for WEB30K).

¹We found learnable positional encodings to yield similar results.

Table 1: NDCG@5 test results on WEB30K

Loss	Self-attention	MLP
Ordinal loss	49.25±0.34	45.72±0.43
NDCGLoss 2++	48.59±0.28	46.04±0.46
LambdaRank	48.26±0.35	45.66±0.39
ListNet	48.21±0.30	44.69±0.43
RMSE	47.23±0.52	45.13±0.51
RankNet	46.90±0.38	44.42±0.46
ListMLE	46.23±0.56	43.87±0.46
	XGBoost	
rank:pairwise	46.8	

5.2.2 Ordinal Loss. We formulated ordinal loss as follows. Multi-level ground truth labels were converted to vectors as follows:

$$\begin{aligned}
0 &\mapsto [0, 0, 0, 0], \\
1 &\mapsto [1, 0, 0, 0], \\
2 &\mapsto [1, 1, 0, 0], \\
3 &\mapsto [1, 1, 1, 0], \\
4 &\mapsto [1, 1, 1, 1].
\end{aligned}$$

The self-attentive scoring function was modified to return four outputs and each output was passed through a sigmoid activation function. Thus, each neuron of the output predicts a single relevancy level, but by the reformulation of ground truth, their relative order is maintained, i.e. if, say, label 2 is predicted, label 1 should be predicted as well (although it is not strictly enforced and model is allowed to predict label 2 without predicting label 1). The final loss value is the mean of binary cross-entropy losses for each relevancy level. During inference, the outputs of all output neurons are summed to produce the final score of an item.

5.2.3 LambdaLoss, RankNet and LambdaRank. We used NDCG-Loss2++ of [30], formulated as follows:

$$l(\mathbf{s}, \mathbf{y}) = - \sum_{y_i > y_j} \log_2 \sum_{\pi} \left(\frac{1}{1 + e^{-\sigma(s_i - s_j)}} \right)^{(\rho_{ij} + \mu \delta_{ij}) |G_i - G_j|} H(\pi | \mathbf{s})$$

where

$$\begin{aligned}
G_i &= \frac{2^{y_i} - 1}{\max \text{DCG}}, \\
\rho_{ij} &= \left| \frac{1}{D_i} - \frac{1}{D_j} \right|, \\
\delta_{ij} &= \left| \frac{1}{D_{|i-j|}} - \frac{1}{D_{|i-j|} + 1} \right|, \\
D_i &= \log_2(1 + i)
\end{aligned}$$

and $H(\pi | \mathbf{s})$ is a hard assignment distribution of permutations, i.e.

$$H(\hat{\pi} | \mathbf{s}) = 1 \text{ and } H(\pi | \mathbf{s}) = 0 \text{ for all } \pi \neq \hat{\pi}$$

where $\hat{\pi}$ is the permutation in which all items are sorted by decreasing scores $\mathbf{s}$. Fixed parameter μ is set to 10.0.

By removing the exponent in $l(\mathbf{s}, \mathbf{y})$ formula we obtain the RankNet loss function, weighing each score pair identically. Similarly, we may obtain differently weighted RankNet variants by changing the formula in the exponent.

Table 2: Relative percentage NDCG@60 improvement on e-commerce search logs dataset

Loss	Self-attention	MLP
NDCGLoss 2++	3.00	1.51
LambdaRank	2.97	1.39
ListNet	2.93	1.24
RankNet	2.68	1.19
	XGBoost	
rank:pairwise	1.83	

To obtain a LambdaRank formula, replace the exponent with

$$\Delta \text{NDCG}(i, j) = |G_i - G_j| \rho_{ij}.$$

5.2.4 ListNet and ListMLE. ListNet loss [9] is given by the following formula:

$$l(\mathbf{s}, \mathbf{y}) = - \sum_j \text{softmax}(\mathbf{y})_j \times \log(\text{softmax}(\mathbf{s})_j)$$

In binary version, softmax of ground truth $\mathbf{y}$ is omitted for single-click lists and replaced with normalisation by the number of clicks for multiple-click lists.

ListMLE [31] is given by:

$$l(\mathbf{s}, \mathbf{y}) = - \log P(\mathbf{y} | \mathbf{s})$$

where

$$P(\mathbf{y} | \mathbf{s}) = \prod_i \frac{\exp(f(x_{y(i)}))}{\sum_{k=i}^n \exp(f(x_{y(k)}))}$$

and $y(i)$ is the index of object which is ranked at position i .

5.3 Experimental setup

We train both our context-aware ranking models and MLP models on both datasets, using all loss functions discussed in Section 5.2². We also train XGBoost models with rank:pairwise loss similar to the production model of the e-commerce search engine for both datasets. Hyperparameters of all models (number of encoder blocks, number of attention heads, dropout, etc.) are tuned on the validation set of Fold 1 for each loss separately. MLP models are constructed to have a similar number of parameters to context-aware ranking models. For optimisation of neural network models, we use Adam optimiser [18] with the learning rate tuned separately for each model. Details of hyperparameters used can be found in Appendix A. In Section 6 we provide an ablation study of the effect of various hyperparameters on the model's performance.

5.4 Results

On WEB30K, models' performance is evaluated using NDCG@5³, which is the usual metric reported for this dataset. Results are reported in Table 1. On e-commerce search logs, we report a relative percentage increase in NDCG@60 over production XGBoost model,

²For the clickthrough logs dataset, we used only the losses which can be applied to binary relevance labels.

³Expressed as a percentage.

presented in Table 2. We observe consistent and significant performance improvement of the proposed self-attention based model over MLP baseline across all types of loss functions considered. In particular, for ListNet we observe a 7.9% performance improvement over MLP baseline on WEB30K. Note also that the best performing MLP model is outperformed even by the worst-performing self-attention based model on both datasets. We thus observe that incorporating context-awareness into the model architecture has a more pronounced effect on the performance of the model than varying the underlying loss function. Surprisingly, ordinal loss outperforms more established and better-studied losses like ListNet, ListMLE or NDCGLoss 2++ on multi-level relevancy data. Another surprising finding is a good performance of models trained with RMSE loss, especially as compared to models trained to optimise RankNet and ListMLE. For comparison with the current state-of-the-art, we provide results on WEB30K reported in other works in Table 3. For models with multiple variants, we cite the best result reported in the original work. In all tables, boldface is the best value column-wise.

Table 3: Results reported in other papers

Method	WEB30K NDCG@5
GSF	41.50
GSF + LambdaMART (re-ranking)	44.90
DLCM	45.00
NDCGLoss 2++ (LightGBM)	51.21

Table 4: NDCG@5 on re-ranking task

Loss	With PE	Self-attention w/o PE
Ordinal loss	49.68	49.21
NDCGLoss 2++	49.24	48.41
RMSE	48.85	47.23
ListNet	48.77	48.34
LambdaRank	48.51	48.23
ListMLE	47.91	46.20
RankNet	47.58	46.91

5.5 Re-ranking

All experiments on WEB30K described above were conducted in the ranking setting - input lists of items were treated as unordered, thus positional encoding was not used. To verify the effect of positional encoding on the model’s performance, we conduct the following experiments on WEB30K. To avoid information leak, training data⁴ is divided into five folds and five XGBoost models are trained, each on four folds. Each model predicts scores for the remaining fold, and the entire dataset is sorted according to these scores.

Finally, we train the same models⁵ as earlier on the sorted dataset, but use fixed positional encoding. Results are presented in Table

⁴In experiments with positional encoding, we used only Fold 1 of the dataset.

⁵The only loss for which we had to modify model’s hyperparameters as compared to training on the ranking task was NDCGLoss 2++. We observed severe overfitting when training on re-ranking task and thus reduced the number of encoder blocks N from 4 to 2 and hidden dimension d_h from 512 to 256.

4. As expected, the models are able to learn positional information and demonstrates improved performance over the plain ranking setting.

Table 5: Ablation study

Parameter	Value	Params	WEB30K NDCG@5
baseline		950K	49.01
H	1	950K	48.99
	4	950K	48.97
N	1	250K	47.33
	2	490K	48.73
d_h	64	430K	47.92
	128	509K	48.31
	256	650K	48.82
	1024	1540K	49.12
p_{drop}	0.0	950K	39.18
	0.1	950K	47.36
	0.2	950K	48.84
	0.3	950K	49.00
	0.5	950K	48.32
l	30	950K	47.48
	60	950K	48.14
	120	950K	48.72
	360	950K	48.89

6 ABLATION STUDY

To gauge the effect of various hyperparameters of self-attention based ranker on its performance, we performed the following ablation study. We trained the context-aware ranker with the ordinal loss on Fold 1 of WEB30K dataset and experimented with a different number N of encoder blocks, H attention heads, length l of longest list used in training, dropout rate p_{drop} and size d_h of hidden dimension. Results are summarised in Table 5. Baseline model (i.e. the best performing context-aware ranker trained with ordinal loss) had the following values of hyperparameters: $N = 4$, $H = 2$, $l = 240$, $p_{drop} = 0.4$ and $d_h = 512$. We observe that a high value of dropout is essential to prevent overfitting. Even though it is better to use multiple attention heads as opposed to a single attention head, using too many results in performance degradation. Notice that increasing hidden dimension yields better performance than one reported in Table 1, however, this comes at a price of a large increase in the number of parameters and thus longer training times. Finally, stacking multiple encoder blocks increases performance. However, we did not test the effect of stacking more than 4 encoder blocks due to GPU memory constraints.

7 CONCLUSION

In this work, we addressed the problem of constructing a context-aware scoring function for learning to rank. We adapted the self-attention based Transformer architecture from the neural machine translation literature to propose a new type of scoring function for LTR. We demonstrated considerable performance gains of proposed neural architecture over MLP baselines across different losses and types of data, both in ranking and re-ranking setting. These

Table 6: Details of hyperparameters used in self-attentive models

Loss	d_{fc}	N	H	d_h	lr	p_{drop}	l	params
WEB30K								
NDCGLoss 2++	128	4	4	512	1e-3	0.3	240	811K
LambdaRank	128	4	4	512	1e-3	0.3	240	811K
ListMLE	256	4	4	512	1e-3	0.3	240	2.14M
ListNet	128	4	4	512	1e-3	0.3	240	811K
Ordinal loss	144	4	2	512	1e-3	0.4	240	949K
RankNet	144	4	2	512	1e-3	0.3	240	949K
E-commerce								
NDCGLoss 2++	128	2	2	128	1e-3	0.0	60	206K
LambdaRank	128	2	2	128	1e-3	0.0	60	206K
ListNet	128	2	2	128	2e-3	0.0	60	206K
RankNet	128	2	2	128	2e-3	0.0	60	206K

Table 7: Details of hyperparameters used in MLP models

Loss	Hidden dimensions	lr	p_{drop}	l	params
WEB30K					
NDCGLoss 2++	[256, 512, 1024, 512, 256]	1e-3	0.3	240	1.35M
LambdaRank					
ListMLE					
ListNet					
Ordinal loss					
RankNet					
E-commerce					
NDCGLoss 2++	[256, 384, 256]	1e-3	0.0	60	210K
LambdaRank	[256, 384, 256]	1e-3	0.0	60	210K
ListNet	[256, 384, 256]	1e-4	0.0	60	210K
RankNet	[256, 384, 256]	1e-4	0.0	60	210K

experiments provide strong evidence that the gains are due to the ability of the model to score items simultaneously. As a result of our empirical study, we observed the strong performance of models trained to optimise ordinal loss function. Such models outperformed models trained with well-studied losses like LambdaLoss or LambdaMART, which were previously shown to provide tight bounds on IR metrics like NDCG. On the other hand, we observed the surprisingly poor performance of models trained to optimise RankNet and ListMLE losses. In future work, we plan to investigate the reasons for both good and poor performance of the aforementioned losses, in particular, the relation between ordinal loss and NDCG.

A EXPERIMENTS DETAILS

Above we provide hyperparameters used for all models reported in Table 1. Models trained on WEB30K were trained for 100 epochs with the learning rate decayed by 0.1 halfway through the training. On e-commerce search logs, we trained the models for 10 epochs and decayed the learning rate by 0.1 after 5-th epoch. The meaning of the columns in Table 6 is as follows: d_{fc} is the dimension of the linear projection done on the input data before passing it to the context-aware ranker, N is the number of encoder blocks, H is the number of attention heads, d_h is the hidden dimension used throughout computations in encoder blocks, lr is the learning rate,

p_{drop} is the dropout probability and l is the list length (lists of items were either padded or subsampled to that length). The last column shows the number of learnable parameters of the model.

In Table 7, *Hidden dimensions* column gives dimensions of subsequent layers of MLP models. The remaining columns have the same meaning as in the previous table.

REFERENCES

- [1] Qingyao Ai, Keping Bi, Jiafeng Guo, and W. Bruce Croft. 2018. Learning a Deep Listwise Context Model for Ranking Refinement. *CoRR* abs/1804.05936 (2018). arXiv:1804.05936 <http://arxiv.org/abs/1804.05936>
- [2] Qingyao Ai, Xuanhui Wang, Sebastian Bruch, Nadav Golbandi, Mike Bendersky, and Marc Najork. 2019. Learning Groupwise Multivariate Scoring Functions Using Deep Neural Networks. In *Proceedings of the 5th ACM SIGIR International Conference on the Theory of Information Retrieval (ICTIR)*.
- [3] Jimmy Ba, Jamie Ryan Kiros, and Geoffrey E. Hinton. 2016. Layer Normalization. *ArXiv* abs/1607.06450 (2016).
- [4] Ricardo A. Baeza-Yates and Berthier Ribeiro-Neto. 1999. *Modern Information Retrieval*. Addison-Wesley Longman Publishing Co., Inc., Boston, MA, USA.
- [5] Irwan Bello, Sayali Kulkarni, Sagar Jain, Craig Boutilier, Ed Hui-hsin Chi, Elad Eban, Xiyang Luo, Alan Mackey, and Ofer Meshi. 2018. Seq2Slate: Re-ranking and Slate Optimization with RNNs. *CoRR* abs/1810.02019 (2018). arXiv:1810.02019 <http://arxiv.org/abs/1810.02019>
- [6] Chris Burges, Tal Shaked, Erin Renshaw, Ari Lazier, Matt Deeds, Nicole Hamilton, and Greg Hullender. 2005. Learning to Rank Using Gradient Descent. In *Proceedings of the 22Nd International Conference on Machine Learning (ICML '05)*. ACM, New York, NY, USA, 89–96. <https://doi.org/10.1145/1102351.1102363>

- [7] Christopher J. Burges, Robert Ragno, and Quoc V. Le. 2007. Learning to Rank with Nonsmooth Cost Functions. In *Advances in Neural Information Processing Systems 19*, B. Schölkopf, J. C. Platt, and T. Hoffman (Eds.). MIT Press, 193–200. <http://papers.nips.cc/paper/2971-learning-to-rank-with-nonsmooth-cost-functions.pdf>
- [8] Christopher J. C. Burges. 2010. *From RankNet to LambdaRank to LambdaMART: An Overview*. Technical Report. Microsoft Research. http://research.microsoft.com/en-us/um/people/cburges/tech_reports/MSR-TR-2010-82.pdf
- [9] Zhe Cao, Tao Qin, Tie-Yan Liu, Ming-Feng Tsai, and Hang Li. 2007. Learning to Rank: From Pairwise Approach to Listwise Approach. In *Proceedings of the 24th International Conference on Machine Learning (ICML '07)*. ACM, New York, NY, USA, 129–136. <https://doi.org/10.1145/1273496.1273513>
- [10] Tianqi Chen and Carlos Guestrin. 2016. XGBoost: A Scalable Tree Boosting System. In *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining (KDD '16)*. ACM, New York, NY, USA, 785–794. <https://doi.org/10.1145/2939672.2939785>
- [11] Jianpeng Cheng, Li Dong, and Mirella Lapata. 2016. Long Short-Term Memory Networks for Machine Reading. In *Proceedings of the 2016 Conference on Empirical Methods in Natural Language Processing*. Association for Computational Linguistics, Austin, Texas, 551–561. <https://doi.org/10.18653/v1/D16-1053>
- [12] Mostafa Dehghani, Hamed Zamani, Aliaksei Severyn, Jaap Kamps, and W. Bruce Croft. 2017. Neural Ranking Models with Weak Supervision. In *Proceedings of the 40th International ACM SIGIR Conference on Research and Development in Information Retrieval (SIGIR '17)*. ACM, New York, NY, USA, 65–74. <https://doi.org/10.1145/3077136.3080832>
- [13] Kaiming He, Xiangyu Zhang, Shaoqing Ren, and Jian Sun. 2015. Deep Residual Learning for Image Recognition. *2016 IEEE Conference on Computer Vision and Pattern Recognition (CVPR)* (2015), 770–778.
- [14] Sepp Hochreiter and Jürgen Schmidhuber. 1997. Long Short-Term Memory. *Neural Comput.* 9, 8 (Nov. 1997), 1735–1780. <https://doi.org/10.1162/neco.1997.9.8.1735>
- [15] Saratchandra Indrakanti, Svetlana Strunjas, Shubhangi Tandon, and Manojkumar Rangasamy Kannadasan. 2019. Exploring the Effect of an Item's Neighborhood on its Sellability in eCommerce. *arXiv:cs.IR/1908.03825*
- [16] Kalervo Järvelin and Jaana Kekäläinen. 2002. Cumulated Gain-based Evaluation of IR Techniques. *ACM Trans. Inf. Syst.* 20, 4 (Oct. 2002), 422–446. <https://doi.org/10.1145/582415.582418>
- [17] Ray Jiang, Sven Gowal, Yuqiu Qian, Timothy Mann, and Danilo J. Rezende. 2019. Beyond Greedy Ranking: Slate Optimization via List-CVAE. In *International Conference on Learning Representations*. <https://openreview.net/forum?id=r1xX42R5Fm>
- [18] Diederik P. Kingma and Jimmy Ba. 2014. Adam: A Method for Stochastic Optimization. <http://arxiv.org/abs/1412.6980> cite arxiv:1412.6980Comment: Published as a conference paper at the 3rd International Conference for Learning Representations, San Diego, 2015.
- [19] Tie-Yan Liu. 2009. Learning to Rank for Information Retrieval. *Found. Trends Inf. Retr.* 3, 3 (March 2009), 225–331. <https://doi.org/10.1561/15000000016>
- [20] Andrew Y. Ng. 2004. Feature Selection, L1 vs. L2 Regularization, and Rotational Invariance. In *Proceedings of the Twenty-first International Conference on Machine Learning (ICML '04)*. ACM, New York, NY, USA, 78–. <https://doi.org/10.1145/1015330.1015435>
- [21] Zhenxing Niu, Mo Zhou, Le Wang, and Xinbo Gao. 2016. Ordinal Regression with Multiple Output CNN for Age Estimation. 4920–4928. <https://doi.org/10.1109/CVPR.2016.532>
- [22] Adam Paszke, Sam Gross, Soumith Chintala, Gregory Chanan, Edward Yang, Zachary DeVito, Zeming Lin, Alban Desmaison, Luca Antiga, and Adam Lerer. 2017. Automatic Differentiation in PyTorch. In *NIPS Autodiff Workshop*.
- [23] Changhua Pei, Yi Zhang, Yongfeng Zhang, Fei Sun, Xiao Lin, Hanxiao Sun, Jian Wu, Peng Jiang, and Wenwu Ou. 2019. Personalized Re-ranking for Recommendation. *arXiv:cs.IR/1904.06813*
- [24] Tao Qin and T. M. Liu. 2013. Introducing LETOR 4.0 Datasets. *ArXiv abs/1306.2597* (2013).
- [25] Kihyuk Sohn, Honglak Lee, and Xinchen Yan. 2015. Learning Structured Output Representation using Deep Conditional Generative Models. In *Advances in Neural Information Processing Systems 28*, C. Cortes, N. D. Lawrence, D. D. Lee, M. Sugiyama, and R. Garnett (Eds.). Curran Associates, Inc., 3483–3491. <http://papers.nips.cc/paper/5775-learning-structured-output-representation-using-deep-conditional-generative-models.pdf>
- [26] Nitish Srivastava, Geoffrey Hinton, Alex Krizhevsky, Ilya Sutskever, and Ruslan Salakhutdinov. 2014. Dropout: A Simple Way to Prevent Neural Networks from Overfitting. *J. Mach. Learn. Res.* 15, 1 (Jan. 2014), 1929–1958. <http://dl.acm.org/citation.cfm?id=2627435.2670313>
- [27] Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N Gomez, Łukasz Kaiser, and Illia Polosukhin. 2017. Attention is All you Need. In *Advances in Neural Information Processing Systems 30*, I. Guyon, U. V. Luxburg, S. Bengio, H. Wallach, R. Fergus, S. Vishwanathan, and R. Garnett (Eds.). Curran Associates, Inc., 5998–6008. <http://papers.nips.cc/paper/7181-attention-is-all-you-need.pdf>
- [28] Oriol Vinyals, Meire Fortunato, and Navdeep Jaitly. 2015. Pointer Networks. In *Advances in Neural Information Processing Systems 28*, C. Cortes, N. D. Lawrence, D. D. Lee, M. Sugiyama, and R. Garnett (Eds.). Curran Associates, Inc., 2692–2700. <http://papers.nips.cc/paper/5866-pointer-networks.pdf>
- [29] Ellen M. Voorhees. 1999. The TREC-8 Question Answering Track Report. In *In Proceedings of TREC-8*. 77–82.
- [30] Xuanhui Wang, Cheng Li, Nadav Golbandi, Mike Bendersky, and Marc Najork. 2018. The LambdaLoss Framework for Ranking Metric Optimization. In *Proceedings of The 27th ACM International Conference on Information and Knowledge Management (CIKM '18)*. 1313–1322.
- [31] Fen Xia, Tie-Yan Liu, Jue Wang, Wensheng Zhang, and Hang Li. 2008. Listwise Approach to Learning to Rank: Theory and Algorithm. In *Proceedings of the 25th International Conference on Machine Learning (ICML '08)*. ACM, New York, NY, USA, 1192–1199. <https://doi.org/10.1145/1390156.1390306>