

From Source to Target: Leveraging Transfer Learning for Predictive Process Monitoring in Organizations

Sven Weinzierl, Sandra Zilker, Annina Liessmann, Martin Käppel, Weixin Wang, Martin Matzner

Business & Information Systems Engineering (2026)

Appendix (available online via <http://link.springer.com>)

Appendix A. Prediction performance for complete model compared to model sub-variants

We conducted an additional evaluation in the form of an ablation analysis to assess the importance of two main components of our proposed model on prediction performance. These components include the second LSTM layer, which transforms the model into an encoder, and the dedicated parameter initialization, which is performed before model training. When the dedicated parameter initialization is deactivated, all model weights and biases are initialized uniformly. Both model components are designed to improve learning from the event data of the source context and facilitate the transfer of process knowledge to the target context.

Table 12 shows that the complete model of our proposed TL-based technique clearly outperforms its sub-variants in terms of prediction performance when it is trained on L_S and tested on L_T for the inter-organizational use case. This confirms the importance of using both components in the prediction model for an effective process knowledge transfer from source to target contexts.

Table 12: Prediction performance of complete model compared to its sub-variants for the inter-organizational use case (average and standard deviation over five runs)

Scenario	Precision	Recall	F1-score	MCC	AUC _{ROC}
COMPLETE MODEL					
Model with two LSTM layers and dedicated parameter initialization					
Train and test on L_S	0.735 ($\pm$.003)	0.694 ($\pm$.003)	0.708 ($\pm$.002)	0.293 ($\pm$.007)	0.730 ($\pm$.005)
Train and test on L_T	0.769 ($\pm$.012)	0.769 ($\pm$.037)	0.766 ($\pm$.024)	0.330 ($\pm$.034)	0.775 ($\pm$.018)
Train on L_S , test on L_T	0.739 ($\pm$.020)	0.755 ($\pm$.050)	0.724 ($\pm$.021)	0.206 ($\pm$.029)	0.711 ($\pm$.020)
SUB-VARIANTS OF MODEL					
Model with two LSTM layers and without dedicated parameter initialization					
Train and test on L_S	0.737 ($\pm$.005)	0.691 ($\pm$.009)	0.706 ($\pm$.006)	0.296 ($\pm$.010)	0.730 ($\pm$.005)
Train and test on L_T	0.770 ($\pm$.011)	0.783 ($\pm$.015)	0.775 ($\pm$.011)	0.336 ($\pm$.027)	0.767 ($\pm$.015)
Train on L_S , test on L_T	0.712 ($\pm$.037)	0.701 ($\pm$.048)	0.702 ($\pm$.035)	0.167 ($\pm$.099)	0.666 ($\pm$.063)
Model with one LSTM layer and with dedicated parameter initialization					
Train and test on L_S	0.738 ($\pm$.004)	0.696 ($\pm$.009)	0.710 ($\pm$.007)	0.301 ($\pm$.007)	0.734 ($\pm$.004)
Train and test on L_T	0.756 ($\pm$.021)	0.743 ($\pm$.058)	0.745 ($\pm$.043)	0.292 ($\pm$.062)	0.728 ($\pm$.030)
Train on L_S , test on L_T	0.719 ($\pm$.023)	0.694 ($\pm$.046)	0.693 ($\pm$.009)	0.171 ($\pm$.058)	0.701 ($\pm$.022)
Model with one LSTM layer and without dedicated parameter initialization					
Train and test on L_S	0.739 ($\pm$.005)	0.686 ($\pm$.008)	0.702 ($\pm$.006)	0.299 ($\pm$.009)	0.731 ($\pm$.006)
Train and test on L_T	0.764 ($\pm$.016)	0.774 ($\pm$.025)	0.767 ($\pm$.019)	0.318 ($\pm$.045)	0.741 ($\pm$.032)
Train on L_S , test on L_T	0.702 ($\pm$.026)	0.688 ($\pm$.040)	0.687 ($\pm$.017)	0.133 ($\pm$.068)	0.670 ($\pm$.044)

Note. Best results for training on L_S and testing on L_T are marked in bold

Appendix B. Prediction performance for different prefix encoding approaches

We conducted an additional evaluation to assess the effect of index-based prefix encoding used in our proposed TL-based PPM technique on the prediction performance compared to alternative prefix encoding approaches. The two alternative prefix encoding approaches used in this evaluation were last-3-state encoding and aggregation encoding. Both are common prefix encoding approaches in outcome-oriented PPM (Teinemaa et al. 2019). The first approach, last-3-state encoding, encodes each prefix only by the last three events, or fewer, if the complete trace includes fewer than three events. The second approach, aggregation encoding, encodes each prefix in an aggregated form by calculating the mean value for each attribute across all events in the prefix.

We applied the index-based encoding, the last-3-state encoding, and the aggregation encoding to our proposed TL-based PPM technique for the inter-organizational use case. We also applied these three encoding approaches to two traditional PPM techniques, using LSTM and XGBoost models. We selected these two traditional PPM techniques because the first uses the same LSTM model as in our proposed technique but not in an TL-setting, and the second uses an XGBoost model, which performed relatively well in the first evaluation of our inter-case organizational use case (see Section 7.2.1).

Table 13 shows that our proposed technique with the index-based prefix encoding clearly outperforms all baseline approaches in terms of prediction performance on the test event data of the target context. This suggests that index-based encoding is an effective prefix encoding approach for our proposed TL-based PPM technique.

Table 13: Prediction performance of our proposed TL-based PPM technique and traditional PPM techniques with different prefix encoding approaches for the inter-organizational use case (average and standard deviation over five runs)

ML approach	Precision	Recall	F1-score	MCC	AUC _{ROC}
LSTM MODEL WITH ACTIVITY EMBEDDING AND RELATIVE CROSS-DOMAIN TIMESTAMP ENCODING					
Index-based encoding					
Train and test on L_S	0.735 ($\pm$.003)	0.694 ($\pm$.003)	0.708 ($\pm$.002)	0.293 ($\pm$.007)	0.730 ($\pm$.005)
Train and test on L_T	0.769 ($\pm$.012)	0.769 ($\pm$.037)	0.766 ($\pm$.024)	0.330 ($\pm$.034)	0.775 ($\pm$.018)
Train on L_S , test on L_T	0.739 ($\pm$.020)	0.755 ($\pm$.050)	0.724 ($\pm$.021)	0.206 ($\pm$.029)	0.711 ($\pm$.020)
Last-3-state encoding					
Train and test on L_S	0.745 ($\pm$.003)	0.680 ($\pm$.009)	0.698 ($\pm$.007)	0.310 ($\pm$.003)	0.745 ($\pm$.001)
Train and test on L_T	0.778 ($\pm$.020)	0.792 ($\pm$.026)	0.779 ($\pm$.017)	0.351 ($\pm$.044)	0.779 ($\pm$.013)
Train on L_S , test on L_T	0.725 ($\pm$.014)	0.694 ($\pm$.054)	0.699 ($\pm$.028)	0.195 ($\pm$.039)	0.709 ($\pm$.025)
Aggregation encoding					
Train and test on L_S	0.724 ($\pm$.016)	0.697 ($\pm$.015)	0.704 ($\pm$.003)	0.265 ($\pm$.039)	0.708 ($\pm$.025)
Train and test on L_T	0.731 ($\pm$.004)	0.741 ($\pm$.018)	0.734 ($\pm$.008)	0.225 ($\pm$.015)	0.706 ($\pm$.013)
Train on L_S , test on L_T	0.718 ($\pm$.023)	0.756 ($\pm$.046)	0.685 ($\pm$.007)	0.074 ($\pm$.015)	0.542 ($\pm$.027)
LSTM MODEL WITH ONE-HOT ACTIVITY ENCODING AND MIN-MAX TIME FEATURE ENCODING					
Index-based encoding					
Train and test on L_S	0.737 ($\pm$.002)	0.677 ($\pm$.004)	0.695 ($\pm$.003)	0.292 ($\pm$.005)	0.721 ($\pm$.005)
Train and test on L_T	0.676 ($\pm$.047)	0.775 ($\pm$.000)	0.678 ($\pm$.001)	0.011 ($\pm$.016)	0.567 ($\pm$.002)
Train on L_S , test on L_T	0.050 ($\pm$.000)	0.224 ($\pm$.000)	0.082 ($\pm$.000)	-0.029 ($\pm$.000)	0.430 ($\pm$.009)
Last-3-state encoding					
Train and test on L_S	0.745 ($\pm$.002)	0.675 ($\pm$.005)	0.694 ($\pm$.004)	0.307 ($\pm$.002)	0.741 ($\pm$.002)
Train and test on L_T	0.624 ($\pm$.050)	0.776 ($\pm$.000)	0.678 ($\pm$.000)	0.004 ($\pm$.009)	0.556 ($\pm$.004)
Train on L_S , test on L_T	0.143 ($\pm$.208)	0.225 ($\pm$.001)	0.084 ($\pm$.003)	-0.007 ($\pm$.015)	0.442 ($\pm$.014)
Aggregation encoding					
Train and test on L_S	0.725 ($\pm$.002)	0.702 ($\pm$.005)	0.711 ($\pm$.003)	0.271 ($\pm$.004)	0.713 ($\pm$.003)
Train and test on L_T	0.602 ($\pm$.000)	0.776 ($\pm$.000)	0.678 ($\pm$.000)	0.000 ($\pm$.000)	0.528 ($\pm$.003)
Train on L_S , test on L_T	0.050 ($\pm$.000)	0.224 ($\pm$.000)	0.082 ($\pm$.000)	-0.029 ($\pm$.000)	0.450 ($\pm$.027)
XGBOOST WITH ONE-HOT ACTIVITY ENCODING					
Index-based encoding					
Train and test on L_S	0.714 ($\pm$.002)	0.743 ($\pm$.002)	0.720 ($\pm$.001)	0.234 ($\pm$.004)	0.738 ($\pm$.002)
Train and test on L_T	0.793 ($\pm$.000)	0.811 ($\pm$.000)	0.793 ($\pm$.000)	0.386 ($\pm$.000)	0.759 ($\pm$.000)
Train on L_S , test on L_T	0.828 ($\pm$.003)	0.257 ($\pm$.044)	0.143 ($\pm$.080)	0.086 ($\pm$.059)	0.604 ($\pm$.004)
Last-3-state encoding					
Train and test on L_S	0.729 ($\pm$.000)	0.750 ($\pm$.000)	0.735 ($\pm$.000)	0.279 ($\pm$.001)	0.759 ($\pm$.000)
Train and test on L_T	0.782 ($\pm$.002)	0.801 ($\pm$.002)	0.785 ($\pm$.002)	0.361 ($\pm$.006)	0.763 ($\pm$.002)
Train on L_S , test on L_T	0.707 ($\pm$.003)	0.541 ($\pm$.002)	0.578 ($\pm$.002)	0.128 ($\pm$.005)	0.616 ($\pm$.002)
Aggregation encoding					
Train and test on L_S	0.723 ($\pm$.000)	0.743 ($\pm$.001)	0.730 ($\pm$.000)	0.265 ($\pm$.001)	0.749 ($\pm$.001)
Train and test on L_T	0.736 ($\pm$.002)	0.759 ($\pm$.004)	0.744 ($\pm$.002)	0.236 ($\pm$.004)	0.692 ($\pm$.003)
Train on L_S , test on L_T	0.831 ($\pm$.000)	0.378 ($\pm$.011)	0.351 ($\pm$.016)	0.225 ($\pm$.009)	0.614 ($\pm$.002)

Note. Best results for training on L_S and testing on L_T are marked in bold

Appendix C. Prediction performance for encoding activity and timestamp information of prefixes separately compared to encoding both jointly

We conducted an additional evaluation to test whether traditional, text-based contextualized embedding models can effectively encode activity and timestamp information about prefixes in a joint manner. For this evaluation, we constructed a baseline technique that jointly encodes the activity and timestamp information about prefixes using one of the four pre-trained contextualized embedding models, which we introduced in the setup of our evaluation. Specifically, the information about a prefix of the length N is fed into the embedding models as “ $a_1; t_1; \dots; a_N; t_N$ ”, where a is an activity and t is a timestamp in the format %d-%m-%Y %H:%M:%S.

The encoded prefixes are then mapped via an MLP model onto the values of the *in-time* prediction target. The MLP model has three fully connected blocks, each of which is followed by tanh activation, batch normalization, and dropout. In addition, the model incorporates residual connections after the second and third blocks using projection layers. The final output is passed through a sigmoid activation function. As with all other DL models used in our evaluation, the internal model parameters were optimized using the Adam optimizer (Kingma and Ba 2015).

Table 14 shows that our TL-based PPM technique achieves a clearly higher prediction performance when trained on L_S and tested on L_T compared to the baseline techniques using the four different pre-trained contextualized embedding models. This suggests that it is ineffective to use traditional, text-based, contextualized embedding models to jointly encode activity and timestamp information about prefixes in our TL setting.

Table 14: Prediction performance for our TL-based PPM technique that encodes activity and timestamp information about a trace separately compared to a baseline technique that encodes activity and timestamp information about a trace jointly. The baseline technique was tested using the four pre-trained contextualized embedding models of our setup for the inter-organizational use case (average and standard deviation over five runs)

Embedding model	Precision	Recall	F1-score	MCC	AUC _{ROC}
SEPARATE ENCODING OF ACTIVITY AND TIMESTAMP INFORMATION					
Proposed approach					
Train and test on L_S	0.735 (± 0.003)	0.694 (± 0.003)	0.708 (± 0.002)	0.293 (± 0.007)	0.730 (± 0.005)
Train and test on L_T	0.769 (± 0.012)	0.769 (± 0.037)	0.766 (± 0.024)	0.330 (± 0.034)	0.775 (± 0.018)
Train on L_S , test on L_T	0.739 (± 0.020)	0.755 (± 0.050)	0.724 (± 0.021)	0.206 (± 0.029)	0.711 (± 0.020)
JOINED ENCODING OF ACTIVITY AND TIMESTAMP INFORMATION					
bert-base-cased					
Train and test on L_S	0.725 (± 0.006)	0.591 (± 0.031)	0.616 (± 0.031)	0.233 (± 0.008)	0.680 (± 0.005)
Train and test on L_T	0.602 (± 0.000)	0.776 (± 0.000)	0.678 (± 0.000)	0.000 (± 0.000)	0.517 (± 0.020)
Train on L_S , test on L_T	0.688 (± 0.006)	0.466 (± 0.035)	0.499 (± 0.038)	0.075 (± 0.016)	0.566 (± 0.007)
bert-base-uncased					
Train and test on L_S	0.714 (± 0.008)	0.666 (± 0.010)	0.682 (± 0.006)	0.238 (± 0.016)	0.682 (± 0.006)
Train and test on L_T	0.602 (± 0.000)	0.776 (± 0.000)	0.678 (± 0.000)	0.000 (± 0.000)	0.526 (± 0.013)
Train on L_S , test on L_T	0.727 (± 0.025)	0.397 (± 0.069)	0.393 (± 0.101)	0.109 (± 0.015)	0.569 (± 0.004)
all-MiniLM-L12-v2					
Train and test on L_S	0.649 (± 0.086)	0.668 (± 0.078)	0.638 (± 0.029)	0.131 (± 0.120)	0.598 (± 0.096)
Train and test on L_T	0.620 (± 0.041)	0.750 (± 0.057)	0.676 (± 0.005)	0.023 (± 0.052)	0.537 (± 0.055)
Train on L_S , test on L_T	0.645 (± 0.040)	0.713 (± 0.076)	0.670 (± 0.028)	0.036 (± 0.035)	0.540 (± 0.044)
all-mpnet-base-v2					
Train and test on L_S	0.651 (± 0.087)	0.688 (± 0.054)	0.657 (± 0.021)	0.140 (± 0.131)	0.609 (± 0.102)
Train and test on L_T	0.656 (± 0.031)	0.763 (± 0.014)	0.684 (± 0.006)	0.021 (± 0.021)	0.594 (± 0.009)
Train on L_S , test on L_T	0.636 (± 0.033)	0.772 (± 0.003)	0.679 (± 0.002)	0.006 (± 0.012)	0.527 (± 0.036)

Note. Best results for training on L_S and testing on L_T are marked in bold