

Column Networks for Collective Classification

Trang Pham, Truyen Tran, Dinh Phung, Svetha Venkatesh

Deakin University, Australia

{phtra, truyen.tran, dinh.phung, svetha.venkatesh}@deakin.edu.au

Abstract

Relational learning deals with data that are characterized by relational structures. An important task is collective classification, which is to jointly classify networked objects. While it holds a great promise to produce a better accuracy than non-collective classifiers, collective classification is computationally challenging and has not leveraged on the recent breakthroughs of deep learning. We present Column Network (CLN), a novel deep learning model for collective classification in multi-relational domains. CLN has many desirable theoretical properties: (i) it encodes multi-relations between any two instances; (ii) it is deep and compact, allowing complex functions to be approximated *at the network level* with a small set of free parameters; (iii) local and relational features are learned simultaneously; (iv) long-range, higher-order dependencies between instances are supported naturally; and (v) crucially, learning and inference are efficient with linear complexity in the size of the network and the number of relations. We evaluate CLN on multiple real-world applications: (a) delay prediction in software projects, (b) PubMed Diabetes publication classification and (c) film genre classification. In all of these applications, CLN demonstrates a higher accuracy than state-of-the-art rivals.

1 Introduction

Relational data are characterized by relational structures between objects or data instances. For example, research publications are linked by citations, web pages are connected by hyperlinks and movies are related through same directors or same actors. Using relations may improve performance in classification as relations between entities may be indicative of relations between classes. A canonical task in learning from this data type is *collective classification* in which networked data instances are classified simultaneously rather than independently to exploit the dependencies in the data (Macskassy and Provost 2007; Neville and Jensen 2007; Richardson and Domingos 2006; Sen et al. 2008). Collective classification is, however, highly challenging. Exact collective inference under general dependencies is intractable. For tractable learning, we often resort to surrogate loss functions such as (structured) pseudo-likelihood (Sutton and McCallum 2007), approximate gradient (Hinton 2002), or iterative schemes, stacked

learning (Choetkiertikul et al. 2015; Kou and Cohen 2007; Macskassy and Provost 2007; Neville and Jensen 2000).

Existing models designed for collective classification are mostly shallow and do not emphasize learning of local and relational features. Deep neural networks, on the other hand, offer automatic feature learning, which is arguably the key behind recent record-breaking successes in vision, speech, games and NLP (LeCun, Bengio, and Hinton 2015; Mnih et al. 2015). With known challenges in relational learning, *can we design a deep neural network that is efficient and accurate for collective classification?* There has been recent work that combines deep learning with structured prediction but the main learning and inference problems for general multi-relational settings remain open (Belanger and McCallum 2016; Do, Arti, and others 2010; Tompson et al. 2014; Yu, Wang, and Deng 2010; Zheng et al. 2015).

In this paper, we present *Column Network* (CLN), an efficient deep learning model for multi-relational data, with emphasis on collective classification. The design of CLN is partly inspired by the columnar organization of neocortex (Mountcastle 1997), in which cortical neurons are organized in vertical, layered mini-columns, each of which is responsible for a small receptive field. Communications between mini-columns are enabled through short-range horizontal connections. In CLN, each mini-column is a feedforward net that takes an input vector – which plays the role of a receptive field – and produces an output class. Each mini-column net not only learns from its own data but also exchanges features with neighbor mini-columns along the pathway from the input to output. Despite the short-range exchanges, the interaction range between mini-columns increases with depth, thus enabling long-range dependencies between data objects.

To be able to learn with hundreds of layers, we leverage the recently introduced highway nets (Srivastava, Greff, and Schmidhuber 2015) as models for mini-columns. With this design choice, CLN becomes a *network of interacting highway nets*. But unlike the original highway nets, CLN’s hidden layers share the same set of parameters, allowing the depth to grow without introducing new parameters (Liao and Poggio 2016; Pham et al. 2016). Functionally, if feedforward nets and highway nets are functional approximators for an input vector, CLN can be thought as an approximator of a grand function that takes a complex network of vectors as input and returns multiple outputs. CLN has many desirable theoretical

properties: (i) it encodes multi-relations between any two instances; (ii) it is deep and compact, allowing complex functions to be approximated *at the network level* with a small set of free parameters; (iii) local and relational features are learned simultaneously; (iv) long-range, higher-order dependencies between instances are supported naturally; and (v) crucially, learning and inference are efficient, linear in the size of network and the number of relations.

We evaluate CLN on real-world applications: (a) delay prediction in software projects, (b) PubMed Diabetes publication classification and (c) film genre classification. In all applications, CLN demonstrates a higher accuracy than state-of-the-art rivals.

2 Preliminaries

Notation convention: We use capital letters for matrices and bold lowercase letters for vectors. The sigmoid function of a scalar x is defined as $\sigma(x) = [1 + \exp(-x)]^{-1}$, $x \in \mathbb{R}$. A function g of a vector $\mathbf{x}$ is defined as $g(\mathbf{x}) = (g(x_1), \dots, g(x_n))$. The operator $*$ is used to denote element-wise multiplication. We use superscript t (e.g. $\mathbf{h}^t$) to denote layers or computational steps in neural networks, and subscript i for the i^{th} element in a set (e.g. $\mathbf{h}_i^t$ is the hidden activation at layer t of entity i in a graph).

2.1 Collective Classification in Multi-relational Setting

We describe the collective classification setting under multiple relations. Given a graph of entities $\mathbf{G} = \{\mathbf{E}, \mathbf{R}, \mathbf{X}, \mathbf{Y}\}$ where $\mathbf{E} = \{e_1, \dots, e_N\}$ are N entities that connect through relations in $\mathbf{R}$. Each tuple $\{e_j, e_i, r\} \in \mathbf{R}$ describes a relation of type r ($r = 1 \dots R$, where R is the number of relation types in $\mathbf{G}$) from entity e_j to entity e_i . Two entities can connect through multiple relations. A relation can be unidirectional or bidirectional. For example, movie A and movie B may be linked by a unidirectional relation `sequel` (A, B) and two bidirectional relations: `same-actor` (A, B) and `same-director` (A, B).

Entities and relations can be represented in an entity graph where a node represents an entity and an edge exists between two nodes if they have at least one relation. Furthermore, e_j is a neighbor of e_i if there is a link from e_j to e_i . Let $\mathcal{N}(i)$ be the set of all neighbors of e_i and $\mathcal{N}_r(i)$ be the set of neighbors related to e_i through relation r . This immediately implies $\mathcal{N}(i) = \cup_{r \in R} \mathcal{N}_r(i)$.

$\mathbf{X} = \{\mathbf{x}_1, \dots, \mathbf{x}_N\}$ is the set of local features, where $\mathbf{x}_i$ is feature vector of entity e_i ; and $\mathbf{Y} = \{y_1, \dots, y_N\}$ with each $y_i \in \{1, \dots, L\}$ is the label of e_i . y_i can either be observed or latent. Given a set of known label entities $\mathbf{E}_{obs}$, a collective classification algorithm simultaneously infers unknown labels of entities in the set $\mathbf{E}_{hid} = \mathbf{E} \setminus \mathbf{E}_{obs}$. In our probabilistic setting, we assume the classifier produces estimate of the joint conditional distribution $P(\mathbf{Y} | \mathbf{G})$.

It is challenging to learn and infer about $P(\mathbf{Y} | \mathbf{G})$. A popular strategy is to employ approximate but efficient iterative methods (Macskassy and Provost 2007). In the next subsection, we describe a highly effective strategy known as *stacked learning*, which partly inspires our work.

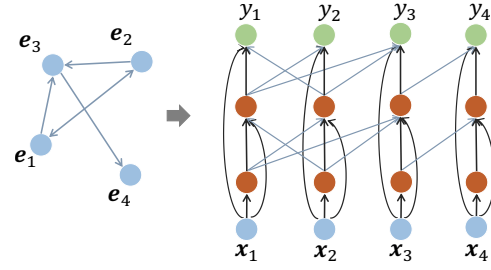


Figure 1: Collective classification with Stacked Learning (SL). (Left): A graph with 4 entities connected by unidirectional and bidirectional links, (Right): SL model for the graph with three steps where $\mathbf{x}_i$ is the feature vector of entity e_i . The bidirectional link between e_1 and e_2 is modeled as two unidirectional links from e_1 to e_2 and vice versa.

2.2 Stacked Learning

Stacked learning (Fig. 1) is a multi-step learning procedure for collective classification (Choetkiertikul et al. 2015; Kou and Cohen 2007; Yu, Wang, and Deng 2010). At step $t - 1$, a classifier is used to predict class probabilities for entity e_j , i.e., $\mathbf{p}_j^{t-1} = [P^{t-1}(y_j = 1), \dots, P^{t-1}(y_j = L)]$. These intermediate outputs are then used as relational features for neighbor classifiers in the next step. In (Choetkiertikul et al. 2015), each relation produces one set of *contextual features*, where all features of the same relation are averaged:

$$\mathbf{c}_{ir}^t = \frac{1}{|\mathcal{N}_r(i)|} \sum_{j \in \mathcal{N}_r(i)} \mathbf{p}_j^{t-1} \quad (1)$$

where $\mathbf{c}_{ir}^t$ is the relational feature vector for relation r at step t . The output at step t is predicted as follows

$$P^t(y_i) = f^t(\mathbf{x}_i, \mathbf{p}_i^{t-1}, [\mathbf{c}_{i1}^t, \mathbf{c}_{i2}^t, \dots, \mathbf{c}_{iR}^t]) \quad (2)$$

where f^t is the classifier at step t . When $t = 1$, the model uses local features of entities for classification, i.e., $\mathbf{c}_{ir}^1 = \mathbf{0}$ and $\mathbf{p}_i^0 = \mathbf{0}$. At each step, classifiers are trained sequentially with known-label entities.

3 Column Networks

In this section we present our main contribution, the Column Network (CLN).

3.1 Architecture

Inspired by the columnar organization in neocortex (Mountcastle 1997), the CLN has one mini-column per entity (or data instance), which is akin to a sensory receptive field. Each column is a feedforward net that passes information from a lower layer to a higher layer of its own, and higher layers of neighbors (see Fig. 2 for a CLN that models the graph in Fig. 1(Left)). The nature of the inter-column communication is dictated by the relations between the two entities.

Through multiple layers, long-range dependencies are established (see Sec. 3.4 for more in-depth discussion). This somewhat resembles the strategy used in stacked learning as described in Sec. 2.2. *The main difference is that in CLN the intermediate steps do not output class probabilities but*

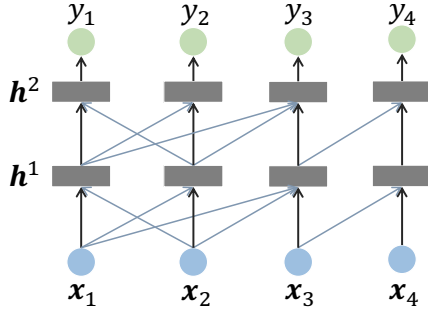


Figure 2: CLN for the graph in Fig. 1(Left) with 2 hidden layers (h^1 and h^2).

learn higher abstraction of instance features and relational features. As such, our model is *end-to-end* in the sense that receptive signals are passed from the bottom to the top, and abstract features are inferred along the way. Likewise, the training signals are passed from the top to the bottom.

Denote by $\mathbf{x}_i \in \mathbb{R}^M$ and $\mathbf{h}_i^t \in \mathbb{R}^{K^t}$ the input feature vector and the hidden activation at layer t of entity e_i , respectively. If there is a connection from entity e_j to e_i , $\mathbf{h}_j^{t-1}$ serves as an input for $\mathbf{h}_i^t$. Generally, $\mathbf{h}_i^t$ is a non-linear function of $\mathbf{h}_i^{t-1}$ and previous hidden states of its neighbors:

$$\mathbf{h}_i^t = g\left(\mathbf{h}_i^{t-1}, \mathbf{h}_{j_1}^{t-1}, \dots, \mathbf{h}_{j_{|\mathcal{N}(i)|}}^{t-1}\right)$$

where $j \in \mathcal{N}(i)$ and $\mathbf{h}_i^0$ is the input vector $\mathbf{x}_i$.

We borrow the idea of stacked learning (Sec. 2.2) to handle multiple relations in CLN. The context of relation r ($r = 1, \dots, R$) at layer t in Eq. (1) is replaced by

$$\mathbf{c}_{ir}^t = \frac{1}{|\mathcal{N}_r(i)|} \sum_{j \in \mathcal{N}_r(i)} \mathbf{h}_j^{t-1} \quad (3)$$

Furthermore, different from stacked learning, the context in CLN are abstracted features, i.e., we replace Eq. (2) by

$$\mathbf{h}_i^t = g\left(\mathbf{b}^t + W^t \mathbf{h}_i^{t-1} + \frac{1}{z} \sum_{r=1}^R V_r^t \mathbf{c}_{ir}^t\right) \quad (4)$$

where $W^t \in \mathbb{R}^{K^t \times K^{t-1}}$ and $V_r^t \in \mathbb{R}^{K^t \times K^{t-1}}$ are weight matrices and $\mathbf{b}^t$ is a bias vector for some activation function g ; z is a pre-defined constant which is used to prevent the sum of parameterized contexts from growing too large for complex relations.

At the top layer T , for example, the label probability for entity i is given as:

$$P(y_i = l) = \text{softmax}\left(\mathbf{b}_l + W_l \mathbf{h}_i^T\right)$$

Remark: There are several similarities between CLN and existing neural network operations. Eq. (3) implements mean-pooling, the operation often seen in CNN. The main difference with the standard CNN is that the mean pooling does not reduce the graph size. This suggests other forms of pooling such as max-pooling or sum-pooling. Asymmetric pooling

can also be implemented based on the concept of attention, that is, Eq. (3) can be replaced by:

$$\mathbf{c}_{ir}^t = \sum_{j \in \mathcal{N}_r(i)} \alpha_j \mathbf{h}_j^{t-1}$$

subject to $\sum_{j \in \mathcal{N}_r(i)} \alpha_j = 1$ and $\alpha_j \geq 0$.

Eq. (4) implements a convolution. For example, standard 3x3 convolutional kernels in images implement 8 relations: *left, right, above, below, above-left, above-right, below-left, below-right*. Supposed that the relations are shared between nodes, the CLN achieves *translation invariance*, similar to that in CNN.

3.2 Highway Network as Mini-Column

We now specify the detail of a mini-column, which we implement by extending a recently introduced feedforward net called Highway Network (Srivastava, Greff, and Schmidhuber 2015). Recall that traditional feedforward nets have a major difficulty of learning with high number of layers. This is due to the nested non-linear structure that prevents the ease of passing information and gradient along the computational path. Highway nets solve this problem by partially opening the gate that lets previous states to propagate through layers, as follows:

$$\mathbf{h}^t = \alpha_1 * \tilde{\mathbf{h}}^t + \alpha_2 * \mathbf{h}^{t-1} \quad (5)$$

where $\tilde{\mathbf{h}}^t$ is a nonlinear candidate function of $\mathbf{h}^{t-1}$ and where $\alpha_1, \alpha_2 \in (0, 1)$ are learnable gates. Since the gates are never shut down completely, data signals and error gradients can propagate very far in a deep net.

For modeling relations, the candidate function $\tilde{\mathbf{h}}^t$ in Eq. (5) is computed using Eq. (4). Likewise, the gates are modeled as:

$$\alpha_1 = \sigma\left(\mathbf{b}_\alpha^t + W_\alpha^t \mathbf{h}_i^{t-1} + \frac{1}{z} \sum_{r=1}^R V_{\alpha r}^t \mathbf{c}_{ir}^t\right) \quad (6)$$

and $\alpha_2 = 1 - \alpha_1$ as for compactness (Srivastava, Greff, and Schmidhuber 2015). Other gating options exists, for example, the p -norm gates where $\alpha_1^p + \alpha_2^p = 1$ for $p > 0$ (Pham et al. 2016).

3.3 Parameter Sharing for Compactness

For feedforward nets, the number of parameters grow with number of hidden layers. In CLN, the number is multiplied by the number of relations (see Eq. (4)). In highway network implementation of mini-columns, a set of parameters for the gates is used thus doubling the number of parameters (see Eq. (6)). For a deep CLN with many relations, the number of parameters may grow faster than the size of training data, leading to overfitting and a high demand of memory. To address this challenge, we borrow the idea of parameter sharing in Recurrent Neural Network (RNN), that is, layers have identical parameters. There has been empirical evidence supporting this strategy in non-relational data (Liao and Poggio 2016; Pham et al. 2016).

With parameter sharing, the depth of the CLN can grow without increasing in model size. This may lead to good performance on small and medium datasets. Sec. 4 provides empirical evidences.

3.4 Capturing Long-range Dependencies

An important property of our proposed deep CLN is the ability to capture long-range dependencies despite only local state exchange as shown in Eqs. (4,6). To see how, let us consider the example in Fig. 2, where x_1 is modeled in h_3^1 and h_3^1 is modeled in h_4^2 , therefore although e_1 does not directly connect to e_4 but information of e_1 is still embedded in h_4^2 through h_3^1 . More generally, after k hidden layers, a hidden activation of an entity can contains information of its expanded neighbors of radius k . When the number of layers is large, the representation of an entity at the top layer contains not only its local features and its directed neighbors, but also the information of the entire graph. With highway networks, all of these levels of representations are accumulated through layers and used to predict output labels.

3.5 Training with mini-batch

As described in Sec. 3.1, h_i^t is a function of h_i^{t-1} and the previous layer of its neighbors. h_i^t therefore can contains information of the entire graph if the network is deep enough. This requires full-batch training which is expensive and not scalable. We propose a very simple yet efficient approximation method that allows mini-batch training. For each mini-batch, the neighbor activations are temporarily frozen to scalars, i.e., gradients are not propagated through this “blanket”. After the parameter update, the activations are recomputed as usual. Experiments showed that the procedure did converge and its performance is comparative with the full-batch training method.

4 Experiments and Results

In this section, we report three real-world applications of CLN on networked data: *software delay estimate*, *PubMed paper classification* and *film genre classification*.

4.1 Baselines

For comparison, we employed a comprehensive suit of baseline methods which include: (a) those designed for collective classification, and (b) deep neural nets for non-collective classification. For the former, we used NetKit¹, an open source toolkit for classification in networked data (Macskassy and Provost 2007). NetKit offers a classification framework consisting of 3 components: a *local classifier*, a *relational classifier* and a *collective inference* method. In our experiments, the local classifier is the Logistic Regression (LR) for all settings; relational classifiers are (i) weighted-vote Relational Neighbor (wvRN), (ii) logistic regression link-based classifier with normalized values (nbD), and (iii) logistic regression link-based classifier with absolute count values (nbC). Collective inference methods include Relaxation Labeling (RL) and Iterative Classification (IC). In total, there are 6 pairs

of “relational classifier – collective inference”: wvRN-RL, wvRN-IC, nbD-RL, nbD-IC, nbC-RL and nbC-IC. For each dataset, results of two best settings will be reported.

We also implemented the state-of-the-art collective classifiers following (Choetkiertikul et al. 2015; Kou and Cohen 2007; Yu, Wang, and Deng 2010): stacked learning with logistic regression (SL-LR) and with random forests (SL-RF).

For deep neural nets, following the latest results in (Liao and Poggio 2016; Pham et al. 2016), we implemented highway network with shared parameters among layers (HWN-noRel). This is essentially a special case of CLN without relational connections.

4.2 Experiment Settings

We report three variants of the CLN: a basic version that uses standard Feedforward Neural Network as mini-column (CLN-FNN) and two versions of CLN-HWN that use highway nets with shared parameters (CLN-HWN-full for full-batch mode and CLN-HWN-mini for mini-batch mode, as described in Sections 3.2, 3.3 and 3.5). All neural nets use ReLU in the hidden layers.

Dropout is applied before and after the recurrent layers of CLN-HWNs and at every hidden layers of CLN-FNN. Each dataset is divided into 3 separated sets: training, validation and test sets. For hyper-parameter tuning, we search for (i) number of hidden layers: 2, 6, 10, ..., 30, (ii) hidden dimensions, and (iii) optimizers: Adam or RMSprop. CLN-FNN has 2 hidden layers and the same hidden dimension with CLN-HWN so that the two models have equal number of parameters. The best training setting is chosen by the validation set and the results of the test set are reported. The result of each setting is reported by the mean result of 5 runs. Code for our model can be found on Github²

4.3 Software Delay Prediction

This task is to predict potential delay for an *issue*, which is an unit of task in an iterative software development life-cycle (Choetkiertikul et al. 2015). The prediction point is when issue planning has completed. Due to the dependencies between issues, the prediction of delay for an issue must take into account all related issues. We use the largest dataset reported in (Choetkiertikul et al. 2015), the JBoss, which contains 8,206 issues. Each issue is a vector of 15 features and connects to other issues through 12 relations (unidirectional such as `blocked-by` or bidirectional such as `same-developer`). The task is to predict whether a software issue is at risk of getting delays (i.e., binary classification).

Fig. 3 visualizes CLN-HWN-full performance with different numbers of layers ranging from 2 to 30 and hidden dimensions from 5, 10 to 20. The F1-score peaks at 10 hidden layers and dimension size of 10.

Table 1 reports the F1-scores of all methods. The two best classifiers in NetKit are wvRN-IC and wvRN-RL. The non-collective HWN-noRel works surprisingly well – almost reaching the performance of the best collective SL-RF with 2 points short. This demonstrates that deep neural nets are

¹<http://netkit-srl.sourceforge.net/>

²https://github.com/trangptm/Column_networks

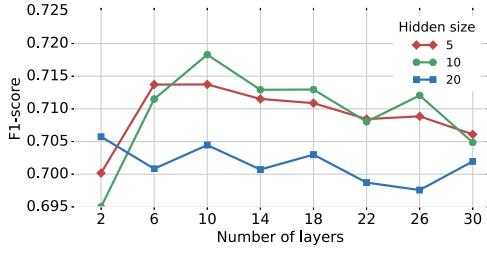


Figure 3: Performance (F1-score) of CLN-HWN on Software delay prediction task with different numbers of layers and hidden sizes

<i>Non-neural</i>	<i>F1</i>	<i>Neural net</i>	<i>F1</i>
wvRN-IC	54.7	HWN-noRel	66.8
wvRN-RL	55.8	CLN-FNN	70.5
SL-LR	65.3	CLN-HWN-full	71.9
SL-RF(*)	68.8	CLN-HWN-mini	71.2

Table 1: Software delay prediction performance. (*) Result reported in (Choetkiertikul et al. 2015).

highly competitive in this domain, and to the best of our knowledge, this fact has not been established. CLN-HWN-full beats the best collective-method, the SL-RF by 3.1 points. We lost 0.7% in mini-batch training mode but the gain of training speed was substantial - roughly 6x.

4.4 PubMed Publication Classification

We used the Pubmed Diabetes dataset consisting of 19,717 scientific publications and 44,338 citation links among them³. Each publication is described by a TF/IDF weighted word vector from a dictionary which consists of 500 unique words. We conducted experiments of classifying each publication into one of three classes: Diabetes Melitus - Experimental, Diabetes Melitus type 1, and Diabetes Mellitus type 2.

Visualization of hidden layers We randomly picked 2 samples of each class and visualized their ReLU units activations through 10 layers of the CLN-HWN (Fig. 4). Interestingly, the activation strength seems to grow with higher layers, suggesting that learn features are more discriminative as they are getting closer to the outcomes. For each class a number of hidden units is turned off in every layer. Figures of samples in the same class have similar patterns while figures of samples from different classes are very different.

Classification accuracy The best setting for CLN-HWN is with 40 hidden dimensions and 10 recurrent layers. Results are measured in MicroF1-score and MacroF1-score (See Table 2). The non-relational highway net (HWN-noRel) outperforms two best baselines from NetKit. The two version of CLN-HWN perform best in both F1-score measures.

4.5 Film Genre Prediction

We used the MovieLens Latest Dataset (Harper and Konstan 2016) which consists of 33,000 movies. The task is to predict

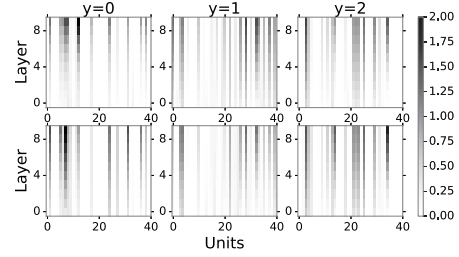


Figure 4: The dynamics of activations of 40 ReLU units through 10 hidden layers of 2x3 samples - each column is for a class.

<i>Method</i>	<i>MicroF1-score</i>	<i>MacroF1-score</i>
wvRN-IC	82.6	81.4
wvRN-RL	82.4	81.2
SL-LR	88.2	87.9
HWN-noRel	87.9	87.9
CLN-FNN	89.4	89.2
CLN-HWN-full	89.8	89.6
CLN-HWN-mini	89.8	89.6

Table 2: Pubmed Diabetes classification results measured by MicroF1-score and MacroF1-score.

genres for each movie given plot summary. Local features were extracted from movie plot summary downloaded from IMDB database⁴. After removing all movies without plot summary, the dataset remains 18,352 movies. Each movie is described by a Bag-of-Words vector of 1,000 most frequent words. Relations between movies are (same-actor, same-director). To create a rather balanced dataset, 20 genres are collapsed into 9 labels: (1) Drama, (2) Comedy, (3) Horror + Thriller, (4) Adventure + Action, (5) Mystery + Crime + Film-Noir, (6) Romance, (7) Western + War + Documentary, (8) Musical + Animation + Children, and (9) Fantasy + Sci-Fi. The frequencies of 9 labels are reported in Table 3.

<i>Label</i>	<i>0</i>	<i>1</i>	<i>2</i>	<i>3</i>	<i>4</i>
Freq(%)	46.3	32.5	24.0	19.1	15.9
<i>Label</i>	<i>5</i>	<i>6</i>	<i>7</i>	<i>8</i>	
Freq(%)	16.5	14.8	10.4	11.5	

Table 3: The frequencies of 9 collapsed labels on Movielens

On this dataset, CLN-HWNs work best with 30 hidden dimensions and 10 recurrent layers. Table 4 reports the F-scores. The two best settings with NetKit are nbC-IC and nbC-RL. CLN-FNN performs well on Micro-F1 but fails to improve MacroF1-score of prediction. CLN-HWN-mini outperforms CLN-HWN-full by 1.3 points on Macro-F1.

Fig. 5 shows why CLN-FNN performs badly on MacroF1 (MacroF1 is the average of all classes' F1-scores). While

³Download: <http://linqs.umiaccs.umd.edu/projects/projects/lbc/>

⁴<http://www.imdb.com>

Method	Micro-F1	Macro-F1
nbC-IC	46.6	38.0
nbC-RL	43.5	40.4
SL-LR	53.4	48.9
HWN-noRel	50.8	45.2
CLN-FNN	54.3	41.8
CLN-HWN-full	57.4	52.7
CLN-HWN-mini	57.5	54.1

Table 4: Movie Genre Classification Performance reported in MicroF1 score and MacroF1 score.

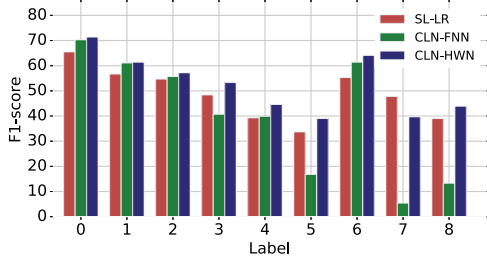


Figure 5: Genre prediction F1-score of SL-LR, CLN-FNN, CLN-HWN on each label. Best viewed in color.

CLN-FNN works well with balanced classes (in the first three classes, its performance is nearly as good as CLN-HWN), it fails to handle imbalanced classes (See Table 3 for label frequencies). For example, F1-score is only 5.4% for label 7 and 13.3% for label 8. In contrast, CLN-HWN performs well on all classes.

5 Related Work

This paper sits at the intersection of two recent independently developed areas: Statistical Relational Learning (SRL) and Deep Learning (DL). Started in the late 1990s, SRL has advanced significantly with noticeable works such as Probabilistic Relational Models (Getoor and Sahami 1999), Conditional Random Fields (Lafferty, McCallum, and Pereira 2001), Relational Markov Network (Taskar, Pieter, and Koller 2002) and Markov Logic Networks (Richardson and Domingos 2006). Collective classification is a canonical task in SRL, also known in various forms as structured prediction (Dietterich et al. 2008) and classification on networked data (Macskassy and Provost 2007).

Two key components of collective classifiers are relational classifier and collective inference (Macskassy and Provost 2007). *Relational classifier* makes use of predicted classes (or class probabilities) of entities from neighbors as features. Examples are wvRN (Macskassy and Provost 2007), logistic based (Domke 2013) or stacked graphical learning (Choetkierikul et al. 2015; Kou and Cohen 2007). *Collective inference* is the task of jointly inferring labels for entities. This is a subject of AI with abundance of solutions including message passing algorithms (Pearl 1988), variational mean-field (Opper and Saad 2001) and discrete optimization (Tran and Dinh Phung 2014). Among existing collective classifiers, the closest to ours is stacked graphical learning where collec-

tive inference is bypassed through stacking (Kou and Cohen 2007; Yu, Wang, and Deng 2010). The idea is based on learning a stack of models that take intermediate prediction of neighborhood into account.

The other area is Deep Learning (DL), where the current wave has offered compact and efficient ways to build multilayered networks for function approximation (via feed-forward networks) and program construction (via recurrent networks) (LeCun, Bengio, and Hinton 2015; Schmidhuber 2015). However, much less attention has been paid to general networked data (Monner, Reggia, and others 2013), although there has been work on pairing structured outputs with deep networks (Belanger and McCallum 2016; Do, Arti, and others 2010; Tompson et al. 2014; Yu, Wang, and Deng 2010; Zheng et al. 2015). Parameter sharing in feedforward networks was recently analyzed in (Liao and Poggio 2016; Pham et al. 2016). The sharing eventually transforms the networks in to recurrent neural networks (RNNs) with only one input at the first layer. The empirical findings were that the performance is good despite the compactness of the model. Among deep neural nets, the closest to our work is RNCC model (Monner, Reggia, and others 2013), which also aims at collective classification using RNNs. There are substantial differences, however. RNCC shuffles neighbors of an entities to a random sequence and uses *horizontal RNN* to integrate the sequence of neighbors. Ours emphasizes on *vertical depth*, where parameter sharing gives rise to the *vertical RNNs*. Ours is conceptually simpler – all nodes are trained simultaneously, not separately as in RNCC.

6 Discussion

This paper has proposed Column Network (CLN), a deep neural network with an emphasis on fast and accurate collective classification. CLN has linear complexity in data size and number of relations in both training and inference. Empirically, CLN demonstrates a competitive performance against rival collective classifiers on three real-world applications: (a) delay prediction in software projects, (b) PubMed Diabetes publication classification and (c) film genre classification.

As the name suggests, CLN is a network of narrow deep networks, where each layer is extended to incorporate as input the preceding neighbor layers. It somewhat resembles the columnar structure in neocortex (Mountcastle 1997), where each narrow deep network plays a role of a mini-column. We wish to emphasize that although we use highway networks in actual implementation due to its excellent performance (Pham et al. 2016; Srivastava, Greff, and Schmidhuber 2015; Tran, Phung, and Venkatesh 2016), *any feedforward networks can be potentially be used* in our architecture. When *parameter sharing* is used, the feedforward networks become recurrent networks, and CLN becomes a network of interacting RNNs. Indeed, the entire network can be collapsed into a giant feedforward network with $n - n$ input/output mappings. When relations are shared among all nodes across the network, CLN enables *translation invariance* across the network, similar to those in CNN. However, the CLN is not limited to a single network with shared relations. Alternatively, networks can be IID according to some distribution and this allows relations to be specific to nodes.

There are open rooms for future work. One extension is to learn the pooling operation using attention mechanisms. We have considered only homogeneous prediction tasks here, assuming instances are of the same type. However, the same framework can be easily extended to multiple instance types.

Acknowledgement

Dinh Phung is partially supported by the Australian Research Council under the Discovery Project DP150100031

References

- Belanger, D., and McCallum, A. 2016. Structured prediction energy networks. *ICML*.
- Choetkiertikul, M.; Dam, H. K.; Tran, T.; and Ghose, A. 2015. Predicting delays in software projects using networked classification. In *30th IEEE/ACM International Conference on Automated Software Engineering*.
- Dietterich, T. G.; Domingos, P.; Getoor, L.; Muggleton, S.; and Tadepalli, P. 2008. Structured machine learning: the next ten years. *Machine Learning* 73(1):3–23.
- Do, T.; Arti, T.; et al. 2010. Neural conditional random fields. In *International Conference on Artificial Intelligence and Statistics*, 177–184.
- Domke, J. 2013. Structured learning via logistic regression. In *Advances in Neural Information Processing Systems*, 647–655.
- Getoor, L., and Sahami, M. 1999. Using probabilistic relational models for collaborative filtering. In *Workshop on Web Usage Analysis and User Profiling (WEBKDD'99)*.
- Harper, F. M., and Konstan, J. A. 2016. The movielens datasets: History and context. *ACM Transactions on Interactive Intelligent Systems (TiiS)* 5(4):19.
- Hinton, G. 2002. Training products of experts by minimizing contrastive divergence. *Neural Computation* 14:1771–1800.
- Kou, Z., and Cohen, W. W. 2007. Stacked Graphical Models for Efficient Inference in Markov Random Fields. In *SDM*, 533–538. SIAM.
- Lafferty, J.; McCallum, A.; and Pereira, F. 2001. Conditional random fields: Probabilistic models for segmenting and labeling sequence data. In *Proceedings of the International Conference on Machine learning (ICML)*, 282–289.
- LeCun, Y.; Bengio, Y.; and Hinton, G. 2015. Deep learning. *Nature* 521(7553):436–444.
- Liao, Q., and Poggio, T. 2016. Bridging the gaps between residual learning, recurrent neural networks and visual cortex. *arXiv preprint arXiv:1604.03640*.
- Macskassy, S., and Provost, F. 2007. Classification in networked data: A toolkit and a univariate case study. *The Journal of Machine Learning Research* 8:935–983.
- Mnih, V.; Kavukcuoglu, K.; Silver, D.; Rusu, A. A.; Veness, J.; Bellemare, M. G.; Graves, A.; Riedmiller, M.; Fidjeland, A. K.; Ostrovski, G.; et al. 2015. Human-level control through deep reinforcement learning. *Nature* 518(7540):529–533.
- Monner, D. D.; Reggia, J.; et al. 2013. Recurrent neural collective classification. *Neural Networks and Learning Systems, IEEE Transactions on* 24(12):1932–1943.
- Mountcastle, V. B. 1997. The columnar organization of the neocortex. *Brain* 120(4):701–722.
- Neville, J., and Jensen, D. 2000. Iterative classification in relational data. In *Proc. AAAI-2000 Workshop on Learning Statistical Models from Relational Data*, 13–20.
- Neville, J., and Jensen, D. 2007. Relational dependency networks. *Journal of Machine Learning Research* 8(Mar):653–692.
- Oppor, M., and Saad, D. 2001. *Advanced mean field methods: Theory and practice*. Massachusetts Institute of Technology Press (MIT Press).
- Pearl, J. 1988. *Probabilistic Reasoning in Intelligent Systems: Networks of Plausible Inference*. San Francisco, CA: Morgan Kaufmann.
- Pham, T.; Tran, T.; Phung, D.; and Venkatesh, S. 2016. Faster training of very deep networks via p-norm gates. *ICPR'16*.
- Richardson, M., and Domingos, P. 2006. Markov logic networks. *Machine Learning* 62:107–136.
- Schmidhuber, J. 2015. Deep learning in neural networks: An overview. *Neural Networks* 61:85–117.
- Sen, P.; Namata, G.; Bilgic, M.; Getoor, L.; Galligher, B.; and Eliassi-Rad, T. 2008. Collective classification in network data. *AI magazine* 29(3):93.
- Srivastava, R. K.; Greff, K.; and Schmidhuber, J. 2015. Training very deep networks. In *Advances in neural information processing systems*, 2377–2385.
- Sutton, C., and McCallum, A. 2007. Piecewise pseudolikelihood for efficient CRF training. In *Proceedings of the International Conference on Machine Learning (ICML)*, 863–870.
- Taskar, B.; Pieter, A.; and Koller, D. 2002. Discriminative probabilistic models for relational data. In *Proceedings of the 18th Conference on Uncertainty in Artificial Intelligence (UAI)*, 485–49. Morgan Kaufmann.
- Tompson, J. J.; Jain, A.; LeCun, Y.; and Bregler, C. 2014. Joint training of a convolutional network and a graphical model for human pose estimation. In *Advances in neural information processing systems*, 1799–1807.
- Tran, T., and Dinh Phung, S. V. 2014. Tree-based Iterated Local Search for Markov Random Fields with Applications in Image Analysis. *Journal of Heuristics* DOI:10.1007/s10732-014-9270-1.
- Tran, T.; Phung, D.; and Venkatesh, S. 2016. Neural choice by elimination via highway networks. *Trends and Applications in Knowledge Discovery and Data Mining, Lecture Notes in Computer Science* 9794:15–25.
- Yu, D.; Wang, S.; and Deng, L. 2010. Sequential labeling using deep-structured conditional random fields. *Selected Topics in Signal Processing, IEEE Journal of* 4(6):965–973.
- Zheng, S.; Jayasumana, S.; Romera-Paredes, B.; Vineet, V.; Su, Z.; Du, D.; Huang, C.; and Torr, P. H. 2015. Conditional random fields as recurrent neural networks. In *Proceedings of the IEEE International Conference on Computer Vision*, 1529–1537.