
Transformer Hawkes Process

Simiao Zuo¹ Haoming Jiang¹ Zichong Li² Tuo Zhao^{1,3} Hongyuan Zha^{4,5,6}

Abstract

Modern data acquisition routinely produce massive amounts of event sequence data in various domains, such as social media, healthcare, and financial markets. These data often exhibit complicated short-term and long-term temporal dependencies. However, most of the existing recurrent neural network based point process models fail to capture such dependencies, and yield unreliable prediction performance. To address this issue, we propose a Transformer Hawkes Process (THP) model, which leverages the self-attention mechanism to capture long-term dependencies and meanwhile enjoys computational efficiency. Numerical experiments on various datasets show that THP outperforms existing models in terms of both likelihood and event prediction accuracy by a notable margin. Moreover, THP is quite general and can incorporate additional structural knowledge. We provide a concrete example, where THP achieves improved prediction performance for learning multiple point processes when incorporating their relational information.

1. Introduction

Event sequence data are naturally observed in our daily life. Through social media such as Twitter and Facebook, we share our experiences and respond to other users' information (Yang et al., 2011). In these websites, each user has a sequence of events such as tweets and interactions. Hun-

dreds of millions of users generate large amounts of tweets, which are essentially sequences of events at different time stamps. Besides social media, event data also exist in domains like financial transactions (Bacry et al., 2015) and personalized healthcare (Wang et al., 2018). For example, in electronic medical records, tests and diagnoses of each patient can be treated as a sequence of events. Unlike other sequential data such as time series, event sequences tend to be asynchronous (Ross et al., 1996), which means time intervals between events are just as important as the order of them to describe their dynamics. Also, depending on specific application requirements, event data show sophisticated dependencies on their history.

Point process is a powerful tool for modeling sequences of discrete events in continuous time, and the technique has been widely applied. Hawkes process (Hawkes, 1971; Isham & Westcott, 1979) and Poisson point process are traditionally used as examples of point processes. However, the simplified assumptions of the complicated dynamics of point processes limit the models' practicality. As an example, Hawkes process states that all past events should have positive influences on the occurrence of current event. However, a user on Twitter may initiate tweets on different topics, and these events should be considered as unrelated instead of mutually-excited.

To alleviate the over-simplifications, likelihood-free methods (Xiao et al., 2017a; Li et al., 2018) and non-parametric models like kernel methods and splines (Vere-Jones et al., 1990) have been proposed, but the increasing complexity and quantity of collected data crave for more powerful models. With the development of neural networks, in particular deep neural networks, focuses have been placed on incorporating these flexible models into classical point processes. Because of the sequential nature of event streams, existing methods rely heavily on Recurrent Neural Networks (RNNs). Neural networks are known for their ability to capture complicated high-level features, in particular, RNNs have the representation power to model the dynamics of event sequence data. In previous works, either vanilla RNN (Du et al., 2016) or its variants (Mei & Eisner, 2017; Xiao et al., 2017b) have been used and significant progress in terms of likelihood and event prediction have been achieved.

However, there are two significant drawbacks with RNN-based models. First, recurrent neural networks, even those

¹Industrial and Systems Engineering, Georgia Institute of Technology, Atlanta, USA; ²School of the Gifted Young, University of Science and Technology of China, Hefei, China; ³Computational Science and Engineering, Georgia Institute of Technology, Atlanta, USA; ⁴Institute for Data and Decision Analytics, the Chinese University of Hong Kong, Shenzhen, Shenzhen, China; ⁵Shenzhen Institute of Artificial Intelligence and Robotics for Society, Shenzhen, China; ⁶Currently on leave from College of Computing, Georgia Institute of Technology. Correspondence to: Simiao Zuo <simiaozuo@gatech.edu>, Tuo Zhao <tourzhao@gatech.edu>, Hongyuan Zha <zahay@cuhk.edu.cn>.

equipped with forget gates, such as Long Short-Term Memory (Hochreiter & Schmidhuber, 1997) and Gated Recurrent Units (Chung et al., 2014), are unlikely to capture long-term dependencies. In financial transactions, short-term effects such as policy changes are important for modeling buy-sell behaviors of stocks. On the other hand, because of the delays in asset returns, stock transactions and prices often exhibit long-term dependencies on their history. As another example, in medical domains, at times we are interested in examining short-term dependencies on symptoms such as fever and cough for acute diseases like pneumonia. But for certain types of chronic diseases such as diabetes, long-term dependencies on disease diagnoses and medications are more critical. Desirable models should be able to capture these long-term dependencies. Yet with recurrent structures, interactions between two events located far in the temporal domain are always weak (Hochreiter et al., 2001), even though in reality they may be highly correlated. The reason is that the probability of keeping information in a state that is far away from the current state decreases exponentially with distance.

The second drawback is trainability of recurrent neural networks. Training deep RNNs (including LSTMs) is notoriously difficult because of gradient explosion and gradient vanishing (Pascanu et al., 2013). In practice, single-layer and two-layer RNNs are mostly used, and they may not successfully model sophisticated dependencies among data (Bengio et al., 1994). Additionally, inputs are fed into the recurrent models sequentially, which means future states must be processed after the current state, rendering it impossible to process all the events in parallel. This limits RNNs' ability to scale to large problems.

Recently, convolutional neural network variants that are tailored for analyzing sequential data (Oord et al., 2016; Gehring et al., 2017; Yin et al., 2017) have been proposed to better capture long-term effects. However, these models enforce many unnecessary dependencies. This particular downside plus the increased computational burdens deem these models insufficient.

To address the above concerns, we propose a Transformer Hawkes Process (THP) model that is able to capture both short-term and long-term dependencies whilst enjoying computational efficiency. Even though the Transformer (Vaswani et al., 2017) is widely adopted in natural language processing, it has rarely been used in other applications. We remark that such an architecture is not readily applicable to event sequences that are defined in a continuous-time domain. To the best of our knowledge, our proposed THP is the first of this type in point process literature.

Building blocks of THP are the self attention modules (Bahdanau et al., 2014). These modules directly model dependencies among events by assigning attention scores. A large

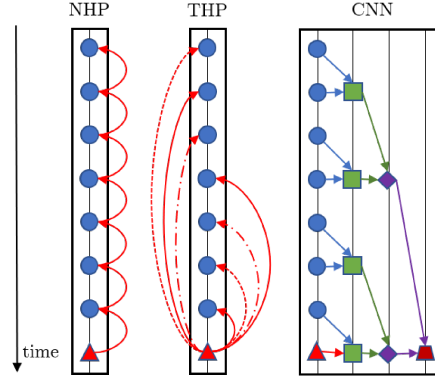


Figure 1. Illustration of dependency computation between the last event (the red triangle) and its history (the blue circles). RNN-based NHP can only model dependencies through recursion. THP directly and adaptively models event's dependencies on its history. Convolution-based models enforce static dependency patterns.

score between two events implies a strong dependency, and a small score implies a weak one. In this way, the modules are able to adaptively select events that are at any temporal distance from the current event. Therefore, THP has the ability to capture both short-term and long-term dependencies. Figure 1 demonstrates dependency computation of different models.

The non-recurrent structure of THP facilitates efficient training of multi-layer models. Transformer-based architectures can be as deep as dozens of layers (Devlin et al., 2018; Radford et al., 2019), where deeper layers capture higher order dependencies. The ability to capture such dependencies creates models that are more powerful than RNNs, which are often shallow. Also, THP allows full parallelism when calculating dependencies across all events, i.e., the computation between any two event pairs is independent with each other. This yields a model presenting strong efficiency.

Our proposed model is quite general, and can incorporate additional structural knowledge to learn more complicated event sequence data, such as multiple point processes over a graph. In social networks, each user has her own sequence of events, like tweets and comments. Sequences among users can be related, for example, a tweet from a user may trigger retweets from her followers. We can use graphs to model these follower-followee relationships (Zhou et al., 2013; Farajtabar et al., 2017), where each vertex corresponds to a specific user and each edge represents connections between the two associated users. We propose an extension to THP that integrates these relational graphs (Borgatti et al., 2009; Linderman & Adams, 2014) into the self-attention module via a similarity metric among users. Such a metric can be learned by our proposed graph regularization.

We experiment THP on five datasets to evaluate both validation likelihood and event prediction accuracy. Our THP

model exhibits superior performance to RNN-based models in all these experiments. We further test our structured-THP on two additional datasets, where the model achieves improved prediction performance for learning multiple point processes when incorporating their relational information.

2. Background

We briefly review Hawkes Process (Hawkes, 1971), Neural Hawkes Process (Mei & Eisner, 2017), and Transformer (Vaswani et al., 2017) in this section.

Hawkes Process is a doubly stochastic point process, whose intensity function is defined as

$$\lambda(t) = \mu + \sum_{j:t_j < t} \psi(t - t_j). \quad (1)$$

Here μ is the base intensity and $\psi(\cdot)$ is a pre-specified decaying function, i.e., exponential function and power-law function. Intuitively, Eq. 1 means that each of the past events has a positive contribution to occurrence of the current event, and this influence decreases through time. However, a major limitation of this formulation is the simplification that history events can never inhibit occurrence of future events, which is unrealistic in complex real-life scenarios.

Neural Hawkes Process generalizes the classical Hawkes process by parameterizing its intensity function with recurrent neural networks. Specifically,

$$\lambda(t) = \sum_{k=1}^K \lambda_k(t) = \sum_{k=1}^K f_k(\mathbf{w}_k^\top \mathbf{h}(t)), \quad t \in (0, T],$$

$$f_k(x) = \beta_k \log \left(1 + \exp \left(\frac{x}{\beta_k} \right) \right),$$

$$\mathbb{P}[k_t = k] = \frac{\lambda_k(t)}{\lambda(t)},$$

where $\lambda(t)$ is the intensity function, K is the number of event types, and $\mathbf{h}(t)$ s are the hidden states of the event sequence, obtained by a continuous-time LSTM (CLSTM) module. CLSTM is an interpolated version of the standard LSTM, and it allows us to generate outputs in a continuous-time domain. Also, $f_k(\cdot)$ is the softplus function with parameter β_k that guarantees a positive intensity. One downside of the neural Hawkes process is that intrinsic weaknesses of RNNs are still inherited, namely the model is unable to capture long-term dependencies and is difficult to train.

Transformer is an attention-based model that has been broadly applied in tasks such as machine translation (Devlin et al., 2018) and language modeling (Radford et al., 2019). Despite its success in natural language processing, it has rarely been used in other areas. We remark that the Transformer architecture is not directly applicable to model point processes. In particular, time intervals between any two events can be arbitrary in event streams, while in natural languages, words are observed on regularly spaced time

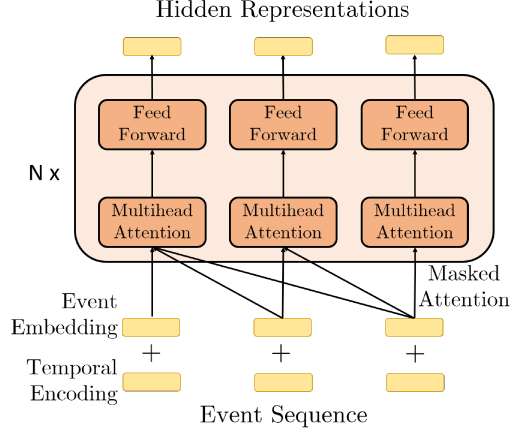


Figure 2. Architecture of the Transformer Hawkes Process. Each event sequence S is fed through embedding layers and N multi-head self-attention modules. Outputs of the THP are hidden representations of events in S , with history information encoded.

intervals. Therefore, we need to generalize the architecture to a continuous-time domain.

3. Model

We introduce our proposed Transformer Hawkes Process. Suppose we are given an event sequence $S = \{(t_j, k_j)\}_{j=1}^L$ of L events, where each event has type $k_j \in \{1, 2, \dots, K\}$, with a total number of K types. Then each pair (t_j, k_j) corresponds to an event of type k_j occurs at time t_j .

3.1. Transformer Hawkes Process

The key ingredient of our proposed THP model is the self-attention module. Different from RNNs, the attention mechanism discards recurrent structures. However, our model still needs to be aware of the temporal information of inputs, i.e., time stamps. Therefore, analogous to the original positional encoding method (Vaswani et al., 2017), we propose to use a temporal encoding procedure, defined by

$$[\mathbf{z}(t_j)]_i = \begin{cases} \cos(t_j/10000^{\frac{i-1}{M}}), & \text{if } i \text{ is odd,} \\ \sin(t_j/10000^{\frac{i}{M}}), & \text{if } i \text{ is even.} \end{cases} \quad (2)$$

Eq. 2 uses trigonometric functions to define a temporal encoding for each time stamp, i.e., for each t_j , we deterministically computes $\mathbf{z}(t_j) \in \mathbb{R}^M$, where M is the dimension of encoding. Other temporal encoding methods can also be applied, such as the relative position representation model (Shaw et al., 2018), where two temporal encoding matrices are learned instead of pre-defined.

Besides temporal encoding, we train an embedding matrix $\mathbf{U} \in \mathbb{R}^{M \times K}$ for the event types, where the k -th column of $\mathbf{U}$ is a M -dimensional embedding for event type k . For any event of type k_j , let $\mathbf{k}_j$ be its one-hot encoding (a K -dimensional vector with all 0s except for the k_j -th index, which has value 1), then its embedding is $\mathbf{U}\mathbf{k}_j$. Notice

that for any event and its corresponding time stamp (t_j, k_j) , the temporal encoding $\mathbf{z}(t_j)$ and the event embedding $\mathbf{U}\mathbf{k}_j$ both reside in $\mathbb{R}^M$. Embedding of the event sequence $\mathcal{S} = \{(t_j, k_j)\}_{j=1}^L$ is then specified by

$$\mathbf{X} = (\mathbf{U}\mathbf{Y} + \mathbf{Z})^\top, \quad (3)$$

where $\mathbf{Y} = [\mathbf{k}_1, \mathbf{k}_2, \dots, \mathbf{k}_L] \in \mathbb{R}^{K \times L}$ is the collection of event type embedding, and $\mathbf{Z} = [\mathbf{z}(t_1), \mathbf{z}(t_2), \dots, \mathbf{z}(t_L)] \in \mathbb{R}^{M \times L}$ is the concatenation of event time encodings. Notice that $\mathbf{X} \in \mathbb{R}^{L \times M}$ and each row of $\mathbf{X}$ corresponds to the embedding of a specific event in the sequence.

After the initial encoding and embedding layers, we pass $\mathbf{X}$ through the self-attention module. Specifically we compute attention output $\mathbf{S}$ by

$$\mathbf{S} = \text{Softmax}\left(\frac{\mathbf{Q}\mathbf{K}^\top}{\sqrt{M_K}}\right)\mathbf{V}, \quad (4)$$

$$\mathbf{Q} = \mathbf{X}\mathbf{W}^Q, \quad \mathbf{K} = \mathbf{X}\mathbf{W}^K, \quad \mathbf{V} = \mathbf{X}\mathbf{W}^V.$$

Here $\mathbf{Q}$, $\mathbf{K}$, and $\mathbf{V}$ are the query, key, and value matrices obtained by different transformations of $\mathbf{X}$, and $\mathbf{W}^Q, \mathbf{W}^K \in \mathbb{R}^{M \times M_K}$, $\mathbf{W}^V \in \mathbb{R}^{M \times M_V}$ are weights for the linear transformations, respectively. In practice using multi-head self-attention to increase model flexibility is more beneficial for data fitting. To facilitate this, different attention outputs $\mathbf{S}_1, \mathbf{S}_2, \dots, \mathbf{S}_H$ are computed using different sets of weights $\{W_j^Q, W_j^K, W_j^V\}_{j=1}^H$. The final attention output for the event sequence is then

$$\mathbf{S} = [\mathbf{S}_1, \mathbf{S}_2, \dots, \mathbf{S}_H]\mathbf{W}^O,$$

where $\mathbf{W}^O \in \mathbb{R}^{HM_V \times M}$ is an aggregation matrix.

We highlight that the self-attention module is able to directly select events whose occurrence time is at any distance from the current time. The j -th column of the attention weights $\text{Softmax}(\mathbf{Q}\mathbf{K}^\top/\sqrt{M_K})$ signifies event t_j 's extent of dependency on its history. In contrast, RNN-based models encode history information sequentially via hidden representations of the events, i.e., the state of t_j depends on that of t_{j-1} , which in turn depends on t_{j-2} , etc. Should any of these encodings be weak, i.e., the RNN fails to learn sufficient relevant information for event t_k , hidden representations of any event t_j where $j \geq k$ will be inferior.

The attention output $\mathbf{S}$ is then fed through a position-wise feed-forward neural network, generating hidden representations $\mathbf{h}(t)$ of the input event sequence:

$$\mathbf{H} = \text{ReLU}(\mathbf{S}\mathbf{W}_1^{\text{FC}} + \mathbf{b}_1)\mathbf{W}_2^{\text{FC}} + \mathbf{b}_2, \quad (5)$$

$$\mathbf{h}(t_j) = \mathbf{H}(j, :).$$

Here $\mathbf{W}_1^{\text{FC}} \in \mathbb{R}^{M \times M_H}$, $\mathbf{W}_2^{\text{FC}} \in \mathbb{R}^{M_H \times M}$, $\mathbf{b}_1 \in \mathbb{R}^{M_H}$, and $\mathbf{b}_2 \in \mathbb{R}^M$ are parameters of the neural network, and $\mathbf{W}_2^{\text{FC}}$ has identical columns. The resulting matrix $\mathbf{H} \in \mathbb{R}^{L \times M}$ contains hidden representations of all the events in the input sequence, where each row corresponds to a particular event.

To avoid ‘‘peeking into the future’’, our attention algorithm is equipped with masks. That is, when computing the attention output $\mathbf{S}_j$ (the j -th column of $\mathbf{S}$), we mask all the future positions, i.e., we set $\mathbf{Q}_{j+1}, \mathbf{Q}_{j+2}, \dots, \mathbf{Q}_L$ to 0. This will avoid the softmax function from assigning dependency to events in the future.

In practice we stack multiple self-attention modules together, and inputs are passed through each of these modules sequentially. In this way our model is able to capture high level dependencies. We remark that stacking RNN/LSTM is not plausible because gradient explosion and gradient vanishing will render the stacked model difficult to train. Figure 2 illustrates the architecture of THP.

3.2. Continuous Time Conditional Intensity

Dynamics of temporal point processes are described by a continuous conditional intensity function. Eq. 5 only generates hidden representations for discrete time stamps, and the associated intensity is also discrete. Therefore an interpolated continuous time intensity function is in need.

Let $\lambda(t|\mathcal{H}_t)$ be the conditional intensity function for our model, where $\mathcal{H}_t = \{(t_j, k_j)\}$ for all j such that $t_j < t$ is the history up to time t . We define different intensity functions for different event types, i.e., for every $k \in \{1, 2, \dots, K\}$, define $\lambda_k(t|\mathcal{H}_t)$ as the conditional intensity function for events of type k . The conditional intensity function for the entire event sequence is defined by

$$\lambda(t|\mathcal{H}_t) = \sum_{k=1}^K \lambda_k(t|\mathcal{H}_t),$$

where each of the type-specific intensity takes the form

$$\lambda_k(t|\mathcal{H}_t) = f_k\left(\underbrace{\alpha_k \frac{t - t_j}{t_j}}_{\text{current}} + \underbrace{\mathbf{w}_k^\top \mathbf{h}(t)}_{\text{history}} + \underbrace{b_k}_{\text{base}}\right). \quad (6)$$

In Eq. 6, time is defined on interval $t \in [t_j, t_{j+1})$, and $f_k(x) = \beta_k \log(1 + \exp(x/\beta_k))$ is the softplus function with ‘‘softness’’ parameter β_k . The reason for choosing this particular function is two-fold: first, the softplus function ensures that the intensity is positive; second, ‘‘softness’’ of the softplus function guarantees stable computation and avoids dramatic changes in intensity.

Now we explain each term in Eq. 6 in detail:

◊ The ‘‘current’’ influence is an interpolation between two observed time stamps t_j and t_{j+1} , and α_k modulates importance of the interpolation. When $t = t_j$, i.e., a new observation comes in, this influence is 0. When $t \rightarrow t_{j+1}$, the conditional intensity function is no longer continuous. As a matter of fact, Eq. 6 is continuous everywhere except for the observed events $\{(t_j, k_j)\}$. However, these ‘‘jumps’’ in intensity is a non-factor when computing likelihood.

◊ The “*history*” term contains two parts: a vector $\mathbf{w}_k$ that transforms the hidden states of the THP model into a scalar, and the hidden states $\mathbf{h}(t)$ (Sec. 3.1) themselves that encode past events up to time t .

◊ The “*base*” intensity represents probability of occurrence of events without considering history information.

With our proposed conditional intensity function, next time stamp prediction and next event type prediction is given by¹

$$\begin{aligned} p(t|\mathcal{H}_t) &= \lambda(t|\mathcal{H}_t) \exp\left(-\int_{t_j}^t \lambda(\tau|\mathcal{H}_\tau) d\tau\right), \\ \hat{t}_{j+1} &= \int_{t_j}^{\infty} t \cdot p(t|\mathcal{H}_t) dt, \\ \hat{k}_{j+1} &= \operatorname{argmax}_k \frac{\lambda_k(t_{j+1}|\mathcal{H}_{j+1})}{\lambda(t_{j+1}|\mathcal{H}_{j+1})}. \end{aligned} \quad (7)$$

3.3. Training

For any sequence $\mathcal{S}$ over an observation interval $[t_1, t_L]$, given its conditional intensity function $\lambda(t|\mathcal{H}_t)$, the log-likelihood is

$$\ell(\mathcal{S}) = \underbrace{\sum_{j=1}^L \log \lambda(t_j|\mathcal{H}_j)}_{\text{event log-likelihood}} - \underbrace{\int_{t_1}^{t_L} \lambda(t|\mathcal{H}_t) dt}_{\text{non-event log-likelihood}}. \quad (8)$$

Model parameters are learned by maximizing the log-likelihood across all sequences. Concretely, suppose we have N sequences $\mathcal{S}_1, \mathcal{S}_2, \dots, \mathcal{S}_N$, then the goal is to find parameters that solve

$$\max \sum_{i=1}^N \ell(\mathcal{S}_i),$$

where $\ell(\mathcal{S}_i)$ is the log-likelihood of event sequence $\mathcal{S}_i$. This optimization problem can be efficiently solved by stochastic gradient type algorithms like ADAM (Kingma & Ba, 2014). Additionally, techniques that help stabilizing training such as layer normalization (Ba et al., 2016) and residual connection (He et al., 2016) are also applied.

In Eq. 8, one challenge is to compute $\Lambda = \int_{t_1}^{t_L} \lambda(t|\mathcal{H}_t) dt$, the non-event log-likelihood. Because of the softplus function, there is no closed-form computation for this integral, and a proper approximation is needed.

The first approach to approximate the non-event log-likelihood is by using Monte Carlo integration (Robert & Casella, 2013):

$$\begin{aligned} \hat{\Lambda}_{\text{MC}} &= \sum_{j=2}^L (t_j - t_{j-1}) \left(\frac{1}{N} \sum_{i=1}^N \lambda(u_i) \right), \\ \nabla \hat{\Lambda}_{\text{MC}} &= \sum_{j=2}^L (t_j - t_{j-1}) \left(\frac{1}{N} \sum_{i=1}^N \nabla \lambda(u_i) \right). \end{aligned} \quad (9)$$

¹Without causing any confusion, denote $\mathcal{H}_{t_j}$ as $\mathcal{H}_j$.

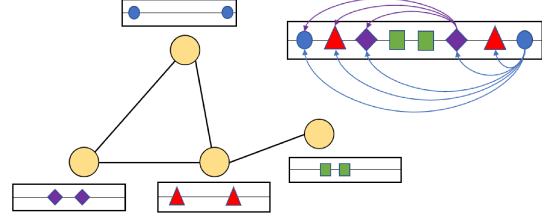


Figure 3. Illustration of event sequences on a graph. Sequences on vertices are aligned temporally to form a long sequence, and relational information among events are shown in arrows. Notice that only the structural information of the last event (the blue circle) and the third to the last event (the purple diamond) are shown. Like before, events cannot attend to future.

Here $u_i \sim \text{Unif}(t_{j-1}, t_j)$ is sampled from a uniform distribution with support $[t_{j-1}, t_j]$. Notice that $\lambda(u_i)$ and $\nabla \lambda(u_i)$ can be calculated by feed-forward and back-propagation through the model, respectively. Moreover, Eq. 9 yields an unbiased estimation to the integral, i.e., $\mathbb{E}[\hat{\Lambda}_{\text{MC}}] = \Lambda$.

The second approach is to apply numerical integration methods, which are faster because of the elimination of sampling. For example, the trapezoidal rule (Stoer & Bulirsch, 2013) states that

$$\hat{\Lambda}_{\text{NU}} = \sum_{j=2}^L \frac{t_j - t_{j-1}}{2} \left(\lambda(t_j|\mathcal{H}_j) + \lambda(t_{j-1}|\mathcal{H}_{j-1}) \right) \quad (10)$$

qualifies as an approximation to Λ . Other higher order methods such as the Simpson’s rule (Stoer & Bulirsch, 2013) can also be applied. Even though approximations build upon numerical integration algorithms are biased, in practice they are affordable. This is because the conditional intensity (Eq. 6) uses softplus as its activation function, which is highly smooth and ensures bias introduced by linear interpolations (Eq. 10) between consecutive events are small.

4. Structured Transformer Hawkes Process

THP is quite general and can incorporate additional structural knowledge. We consider multiple point processes, where any two of them can be related. Such relationships are often described by a graph $\mathcal{G} = (\mathcal{V}, \mathcal{E})$, where $\mathcal{V}$ is the vertex set, and each vertex is associated with a point process. Also, $\mathcal{E}$ is the edge set, where each edge signifies relational information between the corresponding two vertices. Figure 3 illustrates event sequences on a graph.

The graph encodes relationships among vertices, and further indicates potential interactions. We propose to model all the point processes with a single THP, and the heterogeneity of the vertices’ point processes is handled by a vertex embedding approach.

Suppose we have an event sequence $\mathcal{S} = \{(t_j, k_j, v_j)\}_{j=1}^L$, where t_j and k_j are time stamps and event types as before. Further, $v_j \in \{1, 2, \dots, |\mathcal{V}|\}$ is an indicator to which ver-

text the event belongs. In addition to the event embedding and the temporal encoding (Eq. 3), we introduce a vertex embedding matrix $\mathbf{E} \in \mathbb{R}^{M \times |\mathcal{V}|}$, where the j -th column of $\mathbf{E}$ denotes the M -dimensional embedding for vertex j . Let $\mathbf{v}_j$ be the one-hot encoding of v_j , then embedding of $\mathcal{S}$ is specified by

$$\mathbf{X} = (\mathbf{U}\mathbf{Y} + \mathbf{E}\mathbf{V} + \mathbf{Z})^\top,$$

where $\mathbf{V} = [\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_L] \in \mathbb{R}^{|\mathcal{V}| \times L}$ is the concatenation of vertex embedding, and other terms are defined in Eq. 3.

The graph attention output is defined by

$$\mathbf{S} = \text{Softmax} \left(\frac{\mathbf{Q}\mathbf{K}^\top}{\sqrt{M_K}} + \mathbf{A} \right) \mathbf{V}_{\text{value}}, \quad (11)$$

$$\mathbf{A} = (\mathbf{E}\mathbf{V})^\top \mathbf{\Omega} (\mathbf{E}\mathbf{V}),$$

where $\mathbf{Q}$, $\mathbf{K}$, and $\mathbf{V}_{\text{value}}$ are the same² as in Eq. 4. Matrix $\mathbf{A} \in \mathbb{R}^{L \times L}$ is a vertex similarity matrix, where each entry $\mathbf{A}_{ij}$ signifies the similarity between two vertices v_i and v_j , and $\mathbf{\Omega} \in \mathbb{R}^{M \times M}$ is a metric to be learned. To extend the graph self-attention module to a multi-head setting, we use different metric matrices $\{\mathbf{\Omega}_j\}_{j=1}^H$ for different heads.

We remark that unlike RNN-based shallow models, in structured-THP, multiple multi-head self-attention modules can be stacked (Figure 2) to learn high level representations, a feature that enables learning of complicated similarities among vertices. Moreover, the vertex similarity matrix enables modeling of even more complicated structured data, such as sequences on dynamically evolving graphs.

With the incorporation of relational information, we need to modify the conditional intensity function accordingly. As an extension to Eq. 6, where each type of events has its own intensity, we define a different intensity function for each event type and each vertex. Specifically,

$$\lambda(t|\mathcal{H}_t) = \sum_{k=1}^K \sum_{v=1}^{|\mathcal{V}|} \lambda_{k,v}(t|\mathcal{H}_t), \quad t \in [t_j, t_{j+1}),$$

$$\lambda_{k,v}(t|\mathcal{H}_t) = f_{k,v} \left(\alpha_{k,v} \frac{t - t_j}{t_j} + \mathbf{w}_{k,v}^\top \mathbf{h}(t) + b_{k,v} \right).$$

Model parameters are learned by maximizing the log-likelihood (Eq. 8) across all sequences. Concretely, suppose we have N sequences $\mathcal{S}_1, \mathcal{S}_2, \dots, \mathcal{S}_N$, then parameters are obtained by solving

$$\max \sum_{i=1}^N \ell(\mathcal{S}_i) + \mu L_{\text{graph}}(\mathbf{V}, \mathbf{\Omega}),$$

where μ is a hyper-parameter and

$$L_{\text{graph}}(\mathbf{V}, \mathbf{\Omega}) = \sum_{k=1}^{|\mathcal{V}|} \sum_{j=1}^k -\log(1 + \exp(\mathbf{V}_j \mathbf{\Omega} \mathbf{V}_k))$$

$$+ \mathbf{1}\{(v_j, v_k) \in \mathcal{E}\} (\mathbf{V}_j \mathbf{\Omega} \mathbf{V}_k).$$

²We use $\mathbf{V}_{\text{value}}$ to denote the value matrix instead of $\mathbf{V}$, which denotes the vertex embedding.

Table 1. Datasets statistics. From left to right columns: name of the dataset, number of event types, number of events in the dataset, and average length per sequence.

Dataset	K	# events	Avg. length
Retweets	3	2, 173, 533	109
MemeTrack	5000	123, 639	3
Financial	2	414, 800	2074
MIMIC-II	75	2, 419	4
StackOverflow	22	480, 413	72
911-Calls	3	290, 293	403
Earthquake	2	256, 932	500

Here $L_{\text{graph}}(\mathbf{V}, \mathbf{\Omega})$ is a regularization term that encourages $\mathbf{V}_j \mathbf{\Omega} \mathbf{V}_k$ to be large when there exists an edge between v_j and v_k . Which means if two vertices are connected in graph $\mathcal{G}$, then the regularizer will promote attention between them, and vice versa.

Notice that in the simplest case, $\mathbf{A}$ in Eq. 11 can be some transformation of the adjacency matrix, i.e., $\mathbf{A}_{ij} = 1$ if $(v_i, v_j) \in \mathcal{E}$, and 0 otherwise. However, we believe that this constraint is too strict, i.e., some connected vertices may not be similar. Therefore, we treat the graph as a guide and introduce a regularization term that encourages $\mathbf{A}$ to be similar to the adjacency matrix, but not enforce it. In this way, our model is more flexible.

5. Experiments

We compare THP against existing models: Recurrent Marked Temporal Point Process (RMTPP) (Du et al., 2016), Neural Hawkes Process (NHP) (Mei & Eisner, 2017), Time Series Event Sequence (TSES) (Xiao et al., 2017b), and Self-attentive Hawkes Processes (SAHP) (Zhang et al., 2019)³. We evaluate the models by per-event log-likelihood (in nats) and event prediction accuracy on held-out test sets. Details about training are deferred to the appendix.

5.1. Datasets

We adopt several datasets to evaluate the models. Table 1 summarizes statistics of the datasets.

Retweets (Zhao et al., 2015): The Retweets dataset contains sequences of tweets, where each sequence contains an origin tweet (i.e., some user initiates a tweet), and some follow-up tweets. We record the time and the user tag of each tweet. Further, users are grouped into three categories based on the number of their followers: “small”, “medium”, and “large”.

MemeTrack (Leskovec & Krevl, 2014): This dataset contains mentions of 42 thousand different memes spanning ten months. We collect data on over 1.5 million documents (blogs, web articles, etc.) from over 5000 websites. Each sequence in this dataset is the life-cycle of a particular meme,

³This is a concurrent work that also employs the Transformer architecture, and we only include results reported in their paper.

where each event (usage of meme) in the sequence is associated with a time stamp and a website id.

Financial Transactions (Du et al., 2016): This financial dataset contains transaction records of a stock in one day. We record the time (in milliseconds) and the action that was taken in each transaction. The dataset is a single long sequence with only two types of events: “buy” and “sell”. The event sequence is further partitioned by time stamps.

Electrical Medical Records (Johnson et al., 2016): MIMIC-II medical dataset collects patients’ visit to a hospital’s ICU in a seven-year period. We treat the visits of each patient as a separate sequence, where each event in the sequence contains a time stamp and a diagnosis.

StackOverflow (Leskovec & Krevl, 2014): StackOverflow is a question-answering website. The website rewards users with badges to promote engagement in the community, and the same badge can be rewarded multiple times to the same user. We collect data in a two-year period, and we treat each user’s reward history as a sequence. Each event in the sequence signifies receipt of a particular medal.

*911-Calls*⁴: The 911-Calls dataset contains emergency phone call records. Calling time, location of the caller, and nature of the emergency are logged for each record. We consider three types of emergencies: EMS, fire, and traffic. We treat location of callers (given by zipcodes) as vertices on a relational information graph. Zipcodes are ranked based on the number of recorded calls, and only the top 75 zipcodes are kept. An undirected edge exists between two vertices if their zipcodes are within 10 of each other.

*Earthquake*⁵: This dataset contains time and location of earthquakes in China in an eight-year period. We partition the records into two categories: “small” and “large”. A relational information graph is built based on geographical locations of the earthquakes, i.e., each province is a vertex and earthquakes are sequences on the vertices. Two vertices are connected if their associated provinces are neighbors.

5.2. Likelihood Comparison

We fit THP and NHP on Retweets and MemeTrack. From Figure 4, we can see that THP outperforms NHP during the entire training process by a large margin on both of the datasets. The reason is because of the complicated nature of social media data, and RNN-based models such as NHP are not powerful enough to model the dynamics.

In the Retweets dataset, we often observe time gaps between two consecutive retweets become larger, and this dynamic can be successfully modeled by temporal encoding. Also,

⁴The dataset is available on www.kaggle.com/mchirico/montcoalert.

⁵The dataset is provided by China Earthquake Data Center. (<http://data.earthquake.cn>)

Table 2. Log-likelihood comparison. Here RT is the Retweets dataset, MT is the MemeTrack dataset, FIN is the Financial Transactions dataset, and SO is the StackOverflow dataset.

Model	RT	MT	FIN	MIMIC-II	SO
RMTTP	-5.99	-6.04	-3.89	-1.35	-2.60
NHP	-5.60	-6.23	-3.60	-1.38	-2.55
SAHP	-4.56	—	—	-0.52	-1.86
THP	-2.04	0.68	-1.11	0.820	0.042

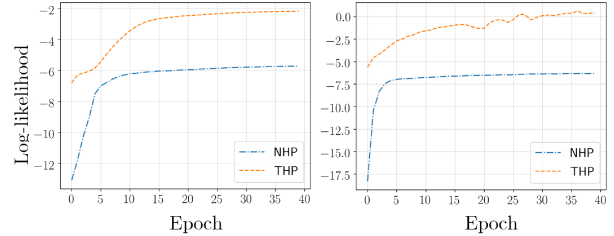


Figure 4. Training curves of NHP and THP fitted on Retweets (left figure) and MemeTrack (right figure).

unlike RNN-based models, our model is able to capture long-term dependencies that exist in long sequences. In the MemeTrack dataset, we have extremely short sequences, i.e., average sequence length is 3. Even though the data only exhibit short-term dependencies, we still need to model latent properties of memes such as topics and targeted users. We build deep THP models to capture these high-level features, and we remark that constructing deep NHP is not plausible because of the difficulty in training.

Table 2 summarizes results on other datasets. Note that TSES is likelihood-free. Our THP model fits the data well and outperforms all the baselines in all the experiments.

Figure 5 visualizes attention patterns of THP. We can see that each attention head employs a different pattern to capture dependencies. Moreover, while attention heads in the first layer tend to focus on individual events, the attention patterns in the last layer are more uniformly distributed. This is because features in deeper layers are already transformed by attention heads in shallow layers.

5.3. Event Prediction Comparison

For point processes, event prediction is just as important as data fitting. Eq. 7 enables us to predict future events. In practice, however, adding additional prediction layers on top of the THP model yields better performance. Specifically, given the hidden representation $\mathbf{h}(t_j)$ for event (t_j, k_j) , the next event type and time predictions are as follows.

◊ The next event type prediction is

$$\begin{aligned}\hat{\mathbf{p}}_{j+1} &= \text{Softmax}(\mathbf{W}^{\text{type}} \mathbf{h}(t_j)), \\ \hat{k}_{j+1} &= \underset{k}{\operatorname{argmax}} \hat{\mathbf{p}}_{j+1}(k),\end{aligned}$$

where $\mathbf{W}^{\text{type}} \in \mathbb{R}^{K \times M}$ is the parameter of the event type predictor, and $\hat{\mathbf{p}}_j(k)$ is the k -th element of $\hat{\mathbf{p}}_j \in \mathbb{R}^K$.

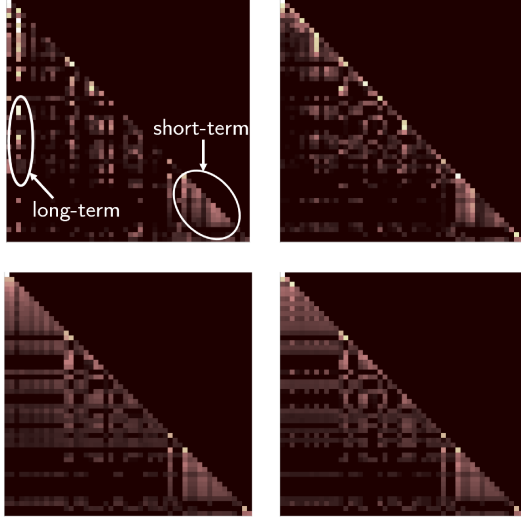


Figure 5. Visualization of attention patterns of different attention heads in different layers. Pixel (i, j) in each figure signifies the attention weight of event t_j attending to event t_i . Attention heads in the upper two figures are from the first layer, while they are from the last layer in the lower two figures.

◇ The next event time prediction is

$$\hat{t}_{j+1} = \mathbf{W}^{\text{time}} \mathbf{h}(t_j),$$

where $\mathbf{W}^{\text{time}} \in \mathbb{R}^{1 \times M}$ is the predictor parameter.

To learn the predictor parameters, the loss function is equipped with a cross-entropy term for event type prediction and a mean square error term for event time prediction. Concretely, for an event sequence $\mathcal{S} = \{(t_j, k_j)\}_{j=1}^L$, let $\mathbf{k}_1, \mathbf{k}_2, \dots, \mathbf{k}_L$ be the ground-truth one-hot encodings for the event types, we define

$$L_{\text{type}}(\mathcal{S}) = \sum_{j=2}^L -\mathbf{k}_j^\top \log(\hat{\mathbf{p}}_j),$$

$$L_{\text{time}}(\mathcal{S}) = \sum_{j=2}^L (t_j - \hat{t}_j)^2,$$

notice that we do not predict the first event. Then, given event sequences $\{\mathcal{S}_i\}_{i=1}^N$, we seek to solve

$$\min \sum_{i=1}^N -\ell(\mathcal{S}_i) + L_{\text{type}}(\mathcal{S}_i) + L_{\text{time}}(\mathcal{S}_i),$$

where $\ell(\mathcal{S}_i)$ is the log-likelihood (Eq. 8) of $\mathcal{S}_i$.

To evaluate model performance, we predict every held-out event (t_j, k_j) given its history $\mathcal{H}_j$, i.e., for a test sequence of length L , we make $L-1$ predictions. We evaluate event type prediction by accuracy and event time prediction by Root Mean Square Error (RMSE). Table 3 and Table 4 summarize experiment results. We can see that THP outperforms the baselines in all these tasks. The datasets we adopted vary significantly in average sequence length, i.e., the average length in Financial Transactions is 2074 while it is only 4 in MIMIC-II. In all the three datasets, THP improves upon RNN-based models by a notable margin. The results

Table 3. Event type prediction accuracy comparison.

Model	Financial	MIMIC-II	StackOverflow
RMTTP	61.95	81.2	45.9
NHP	62.20	83.2	46.3
TSES	62.17	83.0	46.2
THP	62.64	85.3	47.0

Table 4. Event time prediction RMSE comparison.

Model	Financial	MIMIC-II	StackOverflow
RMTTP	1.56	6.12	9.78
NHP	1.56	6.13	9.83
TSES	1.50	4.70	8.00
SAHP	—	3.89	5.57
THP	0.93	0.82	4.99

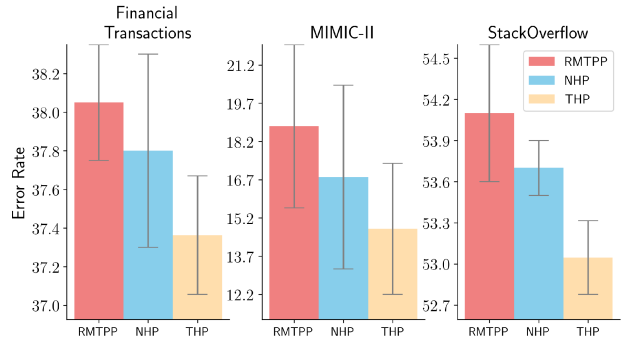


Figure 6. Prediction error rates of THP, NHP, and RMTTP. Based on a same train-dev-test splitting ratio, each dataset is sampled five times to produce different train, development and test sets. Error bars are generated according to these experiments.

demonstrate that THP is able to capture both short-term and long-term dependencies better than existing methods.

Figure 6 illustrates run-to-run variance of THP, NHP, and RMTTP. The error bars are wide because of how the data are split. Held-out test sets are constructed by randomly sampling some events from the entire dataset. That is, at times “important” events are sampled out and that will yield unsatisfactory model performance. Our results are better than all the baselines in all the individual experiments.

5.4. THP vs. Structured-THP

Now we demonstrate by incorporating relational information, THP achieves improved performance.

Baseline models are constructed as following: for each vertex on a relational graph $\mathcal{G}$, there exists a point process that consists of time and type of events. These event sequences are learned separately by both THP and NHP, i.e., we do not allow information sharing among vertices in these models.

To integrate $\mathcal{G}$ into THP, we consider two approaches. The first approach is by allowing full attention, i.e., information from one vertex can be shared with all the other vertices. The second approach is by using the neighborhood graph,

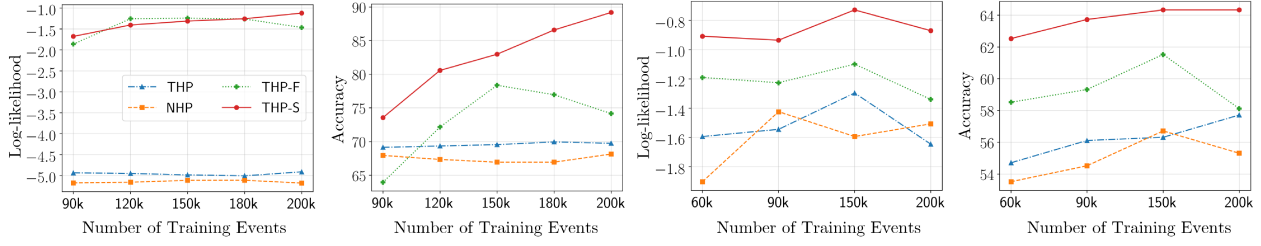


Figure 7. Log-likelihood and prediction accuracy of NHP, THP, THP with full attention (THP-F), and structured-THP (THP-S) fitted on the 911-Calls (left two figures) and the Earthquake (right two figures) datasets. Models are trained using different number of events.

which is constructed based on spatial proximity. In this approach, a specific vertex can only share information with its neighbors. We fit a structured-THP to both of the cases.

Figure 7 summarizes experimental results. We can see that THP is comparable or better than NHP in both validation likelihood and event prediction, which further demonstrates that THP can model complicated dynamics better than RNN-based models. Notice that THP-F, the structured-THP with full attention, yields a much better likelihood than the baseline models, which means relational information sharing can help the models in capturing latent dynamics. However, unlike likelihood, THP-F does not show consistent improvements in event prediction. This is because when the number of training events is small, the model cannot build a sufficient information-sharing heuristic. Also, the performance drop when the number of training events is large is due to the inhomogeneity of data. This demonstrates that the full attention scheme results in undesirable dependencies on which the attention heads focus. THP-S successfully resolves this issue by eliminating such dependencies from the attention heads’ span based on spatial closeness of vertices. In this way THP-S further improves upon THP-F, especially in event prediction tasks.

5.5. Ablation Study

We perform ablation study on Retweets and MemeTrack, and we evaluate models by validation log-likelihood. We inspect variants of THP by removing self-attention and temporal encoding mechanisms. Moreover, we test the effect of temporal encoding on NHP. Table 5 summarizes experimental results. As shown, both the self-attention module and the temporal encoding contribute to model performance.

We examine the models’ sensitivity to the number of parameters on the Retweets dataset. As shown in Table 6, our model is not sensitive to its number of parameters. Without the recurrent structure, Transformer-based models often have large number of parameters, but our THP model can outperform RNN-based models with fewer parameters. In all the experiments, using a small model (about 100-200k parameters) will suffice. In comparison, NHP has about 1000k and TSES has about 2000k parameters to achieve the

Table 5. Log-likelihood of variants of NHP and THP fitted on Retweets and MemeTrack. TE stands for temporal encoding (Eq. 2), and PE stands for positional encoding (Vaswani et al., 2017).

Model	Retweets	MemeTrack
NHP	−5.60	−6.23
NHP + TE	−2.50	−1.64
Atten	−5.29	−5.09
Atten + PE	−5.25	−4.70
Atten + TE	−2.03	0.68

Table 6. Sensitivity to the number of parameters and run-time comparison. Speedup is the speed of THP against NHP.

# parameters	Log-likelihood		Speedup
	THP	NHP	
100k	−2.090	−6.019	×1.985
200k	−2.072	−5.595	×2.564
500k	−2.058	−5.590	×2.224
1000k	−2.060	−5.614	×1.778

best performance, which are much larger than THP. We also include run-time comparison in Table 6. We conclude that THP is efficient in both model size and training speed.

6. Conclusion

In this paper we present Transformer Hawkes Process, a framework for analyzing event streams. Event sequence data are common in our daily life, and they exhibit sophisticated short-term and long-term dependencies. Our proposed model utilizes the self-attention mechanism to capture both of these dependencies, and meanwhile enjoys computational efficiency. Moreover, THP is quite general and can integrate structural knowledge into the model. This facilitates analyzing more complicated data, such as event sequences on graphs. Experiments on various real-world datasets demonstrate that THP achieves state-of-the-art performance in terms of both likelihood and event prediction accuracy.

Acknowledgement

The work done by Haoming Jiang and Tuo Zhao is partially supported by NSF III 1717916. The work done by Hongyuan Zha is supported by Shenzhen Institute of Artificial Intelligence and Robotics for Society, and Shenzhen Research Institute of Big Data.

References

- Ba, J. L., Kiros, J. R., and Hinton, G. E. Layer normalization. *arXiv preprint arXiv:1607.06450*, 2016.
- Bacry, E., Mastromatteo, I., and Muzy, J.-F. Hawkes processes in finance. *Market Microstructure and Liquidity*, 1(01):1550005, 2015.
- Bahdanau, D., Cho, K., and Bengio, Y. Neural machine translation by jointly learning to align and translate. *arXiv preprint arXiv:1409.0473*, 2014.
- Bengio, Y., Simard, P., and Frasconi, P. Learning long-term dependencies with gradient descent is difficult. *IEEE transactions on neural networks*, 5(2):157–166, 1994.
- Borgatti, S. P., Mehra, A., Brass, D. J., and Labianca, G. Network analysis in the social sciences. *science*, 323(5916):892–895, 2009.
- Chung, J., Gulcehre, C., Cho, K., and Bengio, Y. Empirical evaluation of gated recurrent neural networks on sequence modeling. *arXiv preprint arXiv:1412.3555*, 2014.
- Devlin, J., Chang, M.-W., Lee, K., and Toutanova, K. Bert: Pre-training of deep bidirectional transformers for language understanding. *arXiv preprint arXiv:1810.04805*, 2018.
- Du, N., Dai, H., Trivedi, R., Upadhyay, U., Gomez-Rodriguez, M., and Song, L. Recurrent marked temporal point processes: Embedding event history to vector. In *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, pp. 1555–1564. ACM, 2016.
- Farajtabar, M., Wang, Y., Gomez-Rodriguez, M., Li, S., Zha, H., and Song, L. Coevolve: A joint point process model for information diffusion and network evolution. *The Journal of Machine Learning Research*, 18(1):1305–1353, 2017.
- Gehring, J., Auli, M., Grangier, D., Yarats, D., and Dauphin, Y. N. Convolutional sequence to sequence learning. In *Proceedings of the 34th International Conference on Machine Learning-Volume 70*, pp. 1243–1252. JMLR. org, 2017.
- Hawkes, A. G. Spectra of some self-exciting and mutually exciting point processes. *Biometrika*, 58(1):83–90, 1971.
- He, K., Zhang, X., Ren, S., and Sun, J. Deep residual learning for image recognition. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pp. 770–778, 2016.
- Hochreiter, S. and Schmidhuber, J. Long short-term memory. *Neural computation*, 9(8):1735–1780, 1997.
- Hochreiter, S., Bengio, Y., Frasconi, P., Schmidhuber, J., et al. Gradient flow in recurrent nets: the difficulty of learning long-term dependencies, 2001.
- Isham, V. and Westcott, M. A self-correcting point process. *Stochastic Processes and Their Applications*, 8(3):335–347, 1979.
- Johnson, A. E., Pollard, T. J., Shen, L., Li-wei, H. L., Feng, M., Ghassemi, M., Moody, B., Szolovits, P., Celi, L. A., and Mark, R. G. Mimic-iii, a freely accessible critical care database. *Scientific data*, 3:160035, 2016.
- Kingma, D. P. and Ba, J. Adam: A method for stochastic optimization. *arXiv preprint arXiv:1412.6980*, 2014.
- Leskovec, J. and Krevl, A. Snap datasets: Stanford large network dataset collection, 2014.
- Li, S., Xiao, S., Zhu, S., Du, N., Xie, Y., and Song, L. Learning temporal point processes via reinforcement learning. In *Advances in neural information processing systems*, pp. 10781–10791, 2018.
- Linderman, S. and Adams, R. Discovering latent network structure in point process data. In *International Conference on Machine Learning*, pp. 1413–1421, 2014.
- Mei, H. and Eisner, J. M. The neural hawkes process: A neurally self-modulating multivariate point process. In *Advances in Neural Information Processing Systems*, pp. 6754–6764, 2017.
- Oord, A. v. d., Dieleman, S., Zen, H., Simonyan, K., Vinyals, O., Graves, A., Kalchbrenner, N., Senior, A., and Kavukcuoglu, K. Wavenet: A generative model for raw audio. *arXiv preprint arXiv:1609.03499*, 2016.
- Pascanu, R., Mikolov, T., and Bengio, Y. On the difficulty of training recurrent neural networks. In *International conference on machine learning*, pp. 1310–1318, 2013.
- Radford, A., Wu, J., Child, R., Luan, D., Amodei, D., and Sutskever, I. Language models are unsupervised multitask learners. *OpenAI Blog*, 1(8), 2019.
- Robert, C. and Casella, G. *Monte Carlo statistical methods*. Springer Science & Business Media, 2013.
- Ross, S. M., Kelly, J. J., Sullivan, R. J., Perry, W. J., Mercer, D., Davis, R. M., Washburn, T. D., Sager, E. V., Boyce, J. B., and Bristow, V. L. *Stochastic processes*, volume 2. Wiley New York, 1996.
- Shaw, P., Uszkoreit, J., and Vaswani, A. Self-attention with relative position representations. *arXiv preprint arXiv:1803.02155*, 2018.

- Stoer, J. and Bulirsch, R. *Introduction to numerical analysis*, volume 12. Springer Science & Business Media, 2013.
- Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., Kaiser, Ł., and Polosukhin, I. Attention is all you need. In *Advances in neural information processing systems*, pp. 5998–6008, 2017.
- Vere-Jones, D., of Wellington. Institute of Statistics, V. U., and Research, O. *Statistical Methods for the Description and Display of Earthquake Catalogues*. Technical report (Victoria University of Wellington. Institute of Statistics and Operations Research). Victoria University of Wellington, 1990.
- Wang, L., Zhang, W., He, X., and Zha, H. Supervised reinforcement learning with recurrent neural network for dynamic treatment recommendation. In *Proceedings of the 24th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining*, pp. 2447–2456. ACM, 2018.
- Xiao, S., Farajtabar, M., Ye, X., Yan, J., Song, L., and Zha, H. Wasserstein learning of deep generative point process models. In *Advances in Neural Information Processing Systems*, pp. 3247–3257, 2017a.
- Xiao, S., Yan, J., Yang, X., Zha, H., and Chu, S. M. Modeling the intensity function of point process via recurrent neural networks. In *Thirty-First AAAI Conference on Artificial Intelligence*, 2017b.
- Yang, S.-H., Long, B., Smola, A., Sadagopan, N., Zheng, Z., and Zha, H. Like like alike: joint friendship and interest propagation in social networks. In *Proceedings of the 20th international conference on World wide web*, pp. 537–546, 2011.
- Yin, W., Kann, K., Yu, M., and Schütze, H. Comparative study of cnn and rnn for natural language processing. *arXiv preprint arXiv:1702.01923*, 2017.
- Zhang, Q., Lipani, A., Kirnap, O., and Yilmaz, E. Self-attentive hawkes processes. *arXiv preprint arXiv:1907.07561*, 2019.
- Zhao, Q., Erdogdu, M. A., He, H. Y., Rajaraman, A., and Leskovec, J. Seismic: A self-exciting point process model for predicting tweet popularity. In *Proceedings of the 21th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, pp. 1513–1522. ACM, 2015.
- Zhou, K., Zha, H., and Song, L. Learning social infectivity in sparse low-rank networks using multi-dimensional hawkes processes. In *Artificial Intelligence and Statistics*, pp. 641–649, 2013.

A. Training Details

In this section we provide details about training.

To facilitate comparison with previous works, all the datasets are used by Du et al. (2016) and Mei & Eisner (2017), except for 911-Calls and Earthquake. Details about data pre-processing and train-dev-test split, as well as downloadable links, can be found in the aforementioned papers. For the 911-Calls dataset, we exclude zipcodes (and the associated events) whose occurrence is scarce, i.e., we only keep zipcodes that have the top 75 frequent occurrences. The dataset contains 141 types of events, and we cluster them into three categories, namely EMS, fire, and traffic. We do not exclude any event in the Earthquake dataset. Earthquakes are partitioned into two categories, “small” and “large”, where small earthquakes are the ones whose Richter scale is equal to or lower than 1.0. We perform this partition because of the imbalance in data, i.e., most of the recorded earthquakes are on small magnitude. Models are trained on 911-Calls and Earthquake with different number of training events. In each experiment, we equally divide the events that are not in the training set in half to construct the development set and the test set.

There are three sets of hyper-parameters that we use throughout the experiments, and they are summarized in Table 7. Besides layer normalization and residual connection, we also employ the dropout technique to avoid overfitting problems. Table 8 contains the specific parameters that are applied for the training of each dataset. In the table, from left to right columns specify: name of the dataset, the set of applied hyper-parameters, batch size, learning rate, and solver for the approximation of integral (MC stands for Monte Carlo integration, and NU stands for numerical integration with the trapezoidal rule), respectively. In the 911-Calls and the Earthquakes datasets, we also employ the graph regularization method, and the corresponding hyper-parameter is set to be 0.01 for all of the experiments. We use a single NVIDIA RTX graphics card to run all the experiments.

Table 7. Sets of hyper-parameters used in training.

Parameters	# head	# layer	M
Set 1	3	3	64
Set 2	6	6	128
Set 3	4	4	512
Parameters	$M_K = M_V$	M_H	dropout
Set 1	16	256	0.1
Set 2	64	2048	0.1
Set 3	512	1024	0.1

Table 8. Hyper-parameters used for training each dataset.

Dataset	set	batch	lr	solver
Retweets	1	16	5×10^{-3}	MC
MemeTrack	1	128	1×10^{-3}	MC
Financial	2	1	1×10^{-4}	NU
MIMIC-II	1	1	1×10^{-4}	NU
StackOverflow	3	4	1×10^{-4}	NU
911-Calls	2	1	1×10^{-5}	MC
Earthquake	3	1	1×10^{-5}	MC