

# Biologically Plausible Learning for NLP Using Spiking Neural Network

by

Tazrian Hossain

21201153

Adiba Amreen Alam

21241035

Nafisa Rahman

21201446

Maisa Kabir

21201018

Nafis Al Shams

21301372

A thesis submitted to the Department of Computer Science and Engineering  
in partial fulfillment of the requirements for the degree of  
B.Sc. in Computer Science

Department of Computer Science and Engineering  
Brac University  
June 2025

© 2025. Brac University  
All rights reserved.

# Declaration

It is hereby declared that

1. The thesis submitted is my/our own original work while completing degree at Brac University.
2. The thesis does not contain material previously published or written by a third party, except where this is appropriately cited through full and accurate referencing.
3. The thesis does not contain material which has been accepted, or submitted, for any other degree or diploma at a university or other institution.
4. We have acknowledged all main sources of help.

**Student's Full Name & Signature:**

---

Tazrian Hossain  
21201153

---

Adiba Amreen Alam  
21241035

---

Nafisa Rahman  
21201446

---

Maisa Kabir  
21201018

---

Nafis Al Shams  
21301372

# Approval

The thesis/project titled “Biologically Plausible Learning for NLP Using Spiking Neural Network” submitted by

1. Tazrian Hossain(21201153)
2. Adiba Amreen Alam(21241035)
3. Nafisa Rahman(21201446)
4. Maisa Kabir(21201018)
5. Nafis Al Shams(21301372)

Of Spring, 2025 has been accepted as satisfactory in partial fulfillment of the requirement for the degree of B.Sc. in Computer Science on June 18, 2025.

## Examining Committee:

Supervisor:

---

Dr. Md. Golam Rabiul Alam  
Professor  
Department of Computer Science and Engineering  
Brac University

Thesis Coordinator:

---

Dr. Md. Golam Rabiul Alam  
Professor  
Department of Computer Science and Engineering  
Brac University

Head of Department:  
(Chair)

---

Dr. Sadia Hamid Kazi, PhD  
Chairperson and Associate Professor  
Department of Computer Science and Engineering  
Brac University

# Abstract

Commonly referred to as the third generation of neural networks, Spiking Neural Networks (SNNs) have attracted plenty of research interest in the last decade mainly due to its energy efficient and biologically realistic approach. Although areas like computer vision and signal processing have benefited significantly from SNNs, it seems NLP is still uncharted territory in neuromorphic devices. Our research seeks to establish the capability of Spike-Timing-Dependent Plasticity (R-STDP) in SNNs to conduct sentiment analysis. R-STDP provides a reward based learning mechanism that adjusts the synaptic weights according to the spike timing and feedback such as classification accuracy. This duplicates dopamine-controlled learning in the human brain. We also employed the Forward-Forward algorithm which replaces traditional backpropagation with local, layer-wise learning based on positive and negative sample contrast, allowing for modular and decentralized training without the need for backward error signals which further enhances biological plausibility. In addition, we employ an optimized rate coding method to convert textual data into spike trains that can then be easily processed by SNN architectures. We show that by applying this model on a benchmark sentiment analysis and affective computing dataset, SNNs, using learning rules such as R-STDP, can harness energy efficiency and the event-based nature of neuromorphic platforms to achieve sentiment classification accuracy (48%, 73%) comparable to conventional approaches. To understand the findings of our work, we compare our model to the existing deep learning models. The results obtained are of particular interest in order to assess the performance of spiking models for low-power NLP tasks, and to tailor SNNs into further machine learning pipelines.

**Keywords:** Spiking Neural Networks; SNN; Synaptic Plasticity; STDP; Dopamine; R-STDP; Natural Language Processing; NLP; Sentiment Classification; Affective Computing; Forward-Forward

# Table of Contents

|                                                                   |           |
|-------------------------------------------------------------------|-----------|
| Declaration                                                       | i         |
| Approval                                                          | ii        |
| Abstract                                                          | iii       |
| Table of Contents                                                 | iv        |
| List of Figures                                                   | vi        |
| List of Tables                                                    | 1         |
| <b>1 Introduction</b>                                             | <b>2</b>  |
| 1.1 Research Problem . . . . .                                    | 3         |
| 1.2 Research Objective . . . . .                                  | 4         |
| <b>2 Literature Review</b>                                        | <b>5</b>  |
| 2.1 Existing Works . . . . .                                      | 5         |
| 2.1.1 Biologically Inspired SNNs . . . . .                        | 5         |
| 2.1.2 SNNs for NLP tasks . . . . .                                | 6         |
| 2.2 Background . . . . .                                          | 6         |
| 2.2.1 Spike Train encoding . . . . .                              | 6         |
| 2.2.2 Spike-Timing-Dependent Plasticity (STDP) and its Variants . | 8         |
| 2.2.3 The Forward-Forward Algorithm . . . . .                     | 9         |
| 2.3 Key Findings . . . . .                                        | 10        |
| <b>3 Methodology</b>                                              | <b>11</b> |
| 3.1 Top Level Overview . . . . .                                  | 11        |
| 3.2 Datasets . . . . .                                            | 12        |
| 3.3 Data Preprocessing . . . . .                                  | 13        |
| 3.4 Proposed SNN model Architecture (SFF-RSTDp) . . . . .         | 14        |
| 3.4.1 Model Structure . . . . .                                   | 14        |
| 3.4.2 Parameters . . . . .                                        | 14        |
| 3.4.3 Decentralized structure . . . . .                           | 14        |
| 3.4.4 Forward-Forward Mechanism . . . . .                         | 15        |
| 3.4.5 Weight Update . . . . .                                     | 16        |

|          |                                                |           |
|----------|------------------------------------------------|-----------|
| <b>4</b> | <b>Experimental Results</b>                    | <b>19</b> |
| 4.1      | Training Module . . . . .                      | 20        |
| 4.2      | Analysis Metrics . . . . .                     | 20        |
| 4.3      | Analysis of proposed model’s results . . . . . | 21        |
| 4.3.1    | Initial Results . . . . .                      | 22        |
| 4.3.2    | CARER Emotion Dataset . . . . .                | 22        |
| 4.3.3    | IMDB dataset . . . . .                         | 23        |
| 4.4      | Comparison with other models . . . . .         | 25        |
| 4.4.1    | Ablation Study . . . . .                       | 30        |
| 4.4.2    | Discussion . . . . .                           | 31        |
| <b>5</b> | <b>Conclusion</b>                              | <b>33</b> |
|          | <b>Bibliography</b>                            | <b>37</b> |

# List of Figures

|      |                                                                                                                                                                                                                                                                                                   |    |
|------|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|----|
| 2.1  | How R-stdp affects learning . . . . .                                                                                                                                                                                                                                                             | 9  |
| 3.1  | Top Level Overview of the Proposed System . . . . .                                                                                                                                                                                                                                               | 11 |
| 3.2  | SFF-RSTD P Model (Note that the network doesn't output a class directly as it computes goodness per class label and chooses the class with maximum goodness) . . . . .                                                                                                                            | 15 |
| 3.3  | Decentralized Layer-wise Training Strategy . . . . .                                                                                                                                                                                                                                              | 16 |
| 3.4  | Different options for overlaying a label on a sample of a 10-class dataset for positive and negative sample generation [45] . . . . .                                                                                                                                                             | 16 |
| 3.5  | Neuron-level learning loop: Each neuron receives spike input from both positive and negative views. It computes its activation goodness (based on membrane potential), then adjusts synaptic weights using R-STDP and surrogate loss to favor correct labels and suppress incorrect ones. . . . . | 18 |
| 4.1  | Accuracy, MRR, Precision and Recall data for CARER dataset . . .                                                                                                                                                                                                                                  | 22 |
| 4.2  | SFF-RSTD P Confusion Matrix for CARER Dataset . . . . .                                                                                                                                                                                                                                           | 23 |
| 4.3  | Accuracy, MRR, Precision and Recall Data for IMDB dataset . . . .                                                                                                                                                                                                                                 | 24 |
| 4.4  | SFF-RSTD P Confusion Matrix for IMDB Dataset . . . . .                                                                                                                                                                                                                                            | 25 |
| 4.5  | Accuracy data for CARER . . . . .                                                                                                                                                                                                                                                                 | 26 |
| 4.6  | MRR for CARER . . . . .                                                                                                                                                                                                                                                                           | 26 |
| 4.7  | Precision and Recall Score for CARER . . . . .                                                                                                                                                                                                                                                    | 27 |
| 4.8  | Accuracy data for IMDB . . . . .                                                                                                                                                                                                                                                                  | 28 |
| 4.9  | MRR for IMDB . . . . .                                                                                                                                                                                                                                                                            | 29 |
| 4.10 | Precision and Recall Score for IMDB . . . . .                                                                                                                                                                                                                                                     | 29 |

# List of Tables

|     |                                                                                 |    |
|-----|---------------------------------------------------------------------------------|----|
| 3.1 | Dataset Split Distribution . . . . .                                            | 13 |
| 4.1 | SFF-RSTD P Performance Metrics . . . . .                                        | 22 |
| 4.2 | CARER Emotion Dataset Results . . . . .                                         | 25 |
| 4.3 | Number of Operations for Each Model (CARER Dataset) . . . . .                   | 27 |
| 4.4 | IMDB Dataset Results . . . . .                                                  | 28 |
| 4.5 | Number of Operations for Each Model (IMDB Dataset) . . . . .                    | 30 |
| 4.6 | Ablation Study Results on IMDB Dataset . . . . .                                | 30 |
| 4.7 | Evaluation Metrics and Operation Count for CARER and IMDB<br>Datasets . . . . . | 31 |

# Chapter 1

## Introduction

Natural Language Processing (NLP) is a field of Artificial Intelligence that is focused on the interaction between computers and human beings with the use of natural language. It consists of algorithms and models which allow a machine to learn, understand, and respond to humans. NLP has a wide area of applications using different kinds of neural networks like Artificial Neural Networks (ANN), Convolutional Neural Networks (CNN) and Spiking Neural Networks (SNN). ANNs and CNNs rely on backpropagation - a learning method that, despite its effectiveness, is widely considered biologically implausible due to its reliance on global error signals, symmetric weight transport, and sequential updates, none of which are observed in the human brain [45].

SNNs, on the other hand, closely mimic the way biological brains function. SNN uses binary signals and inherently processes time-based data, making them ideal for neuromorphic computing, robotics or sensory data processing [30]. The two primary fields of SNN research at the moment are computer vision and information processing where it has achieved a performance comparable to real-valued equivalent neural networks (although with a slight performance reduction). Furthermore, the energy consumption of SNN is far lower than that of their real-valued equivalents [38]. This is why it is intriguing to examine popular NLP tasks in this context of neuromorphic computing since many research subjects, including NLP, are currently understudied [34] in the setting of spiking neural networks.

In this research, we explore the use of SNNs for NLP tasks by first employing rate coding [43] to convert word embeddings into spike trains, making linguistic data compatible with spiking models. We then utilize Reward-Modulated Spike-Timing-Dependent Plasticity (R-STDP), a biologically inspired learning rule that adjusts synaptic weights based on both spike timing [18] and external reward signals [16]—similar to dopaminergic modulation in biological brains. However, to further increase the biological plausibility and efficiency of learning, we integrate Forward-Forward (FF) learning [33] into our architecture. FF training replaces backward error propagation with two forward passes: one with positive (correct) data and one with negative (incorrect) data. Each layer learns locally to increase its activation for positive inputs and decrease it for negatives, aligning closely with how distributed and parallel learning occurs in the brain.

By combining R-STDP with the FF algorithm, our proposed hybrid framework not only avoids biologically implausible mechanisms but also enables efficient, reward-sensitive and modular learning for NLP tasks, pushing the boundary of neuromorphic computing into a largely unexplored yet promising domain.

## 1.1 Research Problem

NLP has been an appealing field of research for an extended period of time now. Consequently, the advancements in this field are highly notable. However, existing research in this field mostly focused on and prioritized the accuracy of the models and accepted the high energy consumption as a necessary trade-off. Thus, the resource-consuming and power-hungry process desperately needs optimization in regards to its energy consumption [23].

This is where SNNs fit into the picture. As SNNs are event-driven in nature, during most of its life-cycle, the neurons remain dormant until activated by the spiking mechanism. This seemed like a suitable approach to the problem of energy optimization. However, in the case of SNN, the trade-off between accuracy and energy consumption still maintains a gap and without bridging this gap, researchers cannot obtain the same performance they have come to expect from ANNs and conventional models [26]. Therefore, researchers shifted their focus to optimize the accuracy of the SNN models without significant energy consumption trade-off. Different researchers have come up with different ways to optimize the SNN models. Some researchers proposed better binarization and encoding methods [43]. On the other hand, researchers also focused on the neural processing of the models itself. Since SNN exclusively uses discrete inputs and generates discrete spikes, conventional training methods like backpropagation are poorly suited for SNNs. Several solutions to this have been proposed and they have variable results. There are mainly 2 different research paths that have been widely used. One is using a surrogate gradient to back propagate and reduce errors [21]. Another path is using Spike Time Dependent Plasticity (STDP), an unsupervised process that strengthens and weakens neural connections [15]. In order to avoid the additional computations of a surrogate gradient, some researchers have focused on optimizing the STDP itself. One of those ways to improve on STDP is R-STDP (Reward Modulated) which has already been used in several researches [29],[17],[32].

Despite their potential, three underexplored opportunities hinder progress. Firstly, most R-STDP implementations use basic rate-coding or ANN-to-SNN conversion techniques [39], neglecting advanced spike-encoding methods that could enhance both accuracy and energy efficiency. Moreover, backpropagation’s incompatibility with neuromorphic systems [6] contrasts with emerging alternatives like the forward-forward algorithm, which mimics cortical layer-wise learning. However, this method has not been systematically adapted for SNNs in NLP. Also no studies combine R-STDP with optimized rate-coding and biologically plausible learning rules (e.g., forward-forward) to create end-to-end efficient SNN architectures for language tasks.

## 1.2 Research Objective

This research aims to advance energy-efficient and biologically plausible NLP by exploring an under-investigated combination of Spiking Neural Networks, Reward-Modulated STDP, and backpropagation-free learning. The main objectives of the research is given below

- To optimize energy consumption of NLP tasks (Sentiment Analysis and Affective Computing) using SNNs. Existing NLP tasks are processed by Transformer based Artificial Neural Networks and LLM models which take a significant number of operations to complete and are inefficient in terms of energy consumption.
- To enhance the learning capability of SNNs for NLP tasks. Existing SNN models compromise performance in order to conserve energy. In the present state, SNN models can't effectively compete with other models in terms of performance.
- To provide biological plausibility to sentiment classification and affective computation tasks. SNN models, by nature, are more close to a biological brain than other existing models.
- To analyze and compare the performance to other models and combinations with different rate-coding techniques. SNN models require spike-trains as inputs and the process of generating these spike-trains can have a notable effect on the model's performance. Comparison of models and their rate-coding techniques is a core objective of our research.
- To contribute practical insights and recommendations for future research on building biologically inspired, energy-aware models for language processing.

# Chapter 2

## Literature Review

### 2.1 Existing Works

#### 2.1.1 Biologically Inspired SNNs

Spiking Neural Networks (SNNs) offer a more biological and realistic approach as an alternative to traditional artificial neural networks due to the nature of using spike-based communication, membrane potential modeling, and temporal dynamics. SNN uses various neuron models, among which Leaky Integrate-and-Fire (LIF) helps to naturally generate spikes (discrete events) [8] which are more relevant in the field of SNN. LIF has a lot of advantages such as biological plausibility, compatibility with neuromorphic hardware, scalability in large networks, and analytical tractability. However, there are some challenges while implementing LIF such as the non-differentiable activation function. To overcome such problems, different approaches can be used. One of the approaches is surrogate gradient descent [22]. Surrogate gradients provide an approximate way to compute gradients during the backpropagation process in SNN. Instead of using the non-differentiable spiking function, a smooth surrogate function is introduced, which allows computation of gradients while the original function remains the same. This approach enables gradient-based optimization, preserves temporal information while also maintaining the spiking dynamics, and also allows a much more efficient path for training deep SNN. The SNN networks align with different cognitive processes that are observed in the human brain such as excitatory and inhibitory signaling, sparse activation, and local plasticity [42]. Models such as excitatory-inhibitory cooperative iterative learning (EICIL) have emulated neural behaviors by maintaining a biologically grounded ratio of 80:20 between the excitatory and inhibitory neurons [40]. On the other hand, other biologically inspired works include dendritic spine-based synaptic pruning strategies [31] and adaptive threshold mechanisms like Bio-inspired Dynamic Energy-Temporal Threshold (BDETT), which stabilizes the firing rate in a dynamic manner that changes the environments [32].

Consequently, SNN's biological significance makes it a perfect place to use brain activities that are inspired by biology in a variety of computational domains. In addition to the methods discussed above, the biologically inspired SNN rule that will mainly concern our research is that of synaptic plasticity.

### 2.1.2 SNNs for NLP tasks

The traditional NLP relies on dense continuous embeddings, which are not compatible with the discrete and temporal inputs that are required by SNNs. [43] proposed a deterministic spike encoding scheme and their model achieved 32x inference and 60x training efficiency compared to ANNs. Another model, SpikeGPT [44], replaced the multi-head-self-attention mechanism with a sequential attention process. It achieved competitive accuracy while reducing operations by 32x on neuromorphic hardware. It became one of the largest backpropagation-trained SNN models to date.

However, it needs to be noted that both SNNLP and SpikeGPT focused on energy efficiency rather than strict biological plausibility. Therefore, exploring the use of biologically inspired methods such as R-STDP and forward-forward learning for NLP tasks still remains an intriguing field of study.

## 2.2 Background

### 2.2.1 Spike Train encoding

In natural Language Processing, the initial salient part of the entire process is to encode the input text into a continuous n-dimensional feature vector termed as ‘embedding’ [4] which maps each token in a multidimensional vector space. Extensive research has been done already in this field to encode input data into embedding vectors. Many word embedding techniques have already been proposed. For example, FastText [9], Word2Vec [3], GloVe [5], etc. In addition, research has not been only limited to Word Embedding. NLP has been a field of extensive research for years and recent research works are capable of encoding larger language aspects like sentences [10], [14] and phrases [9] while retaining the syntactic information of the original input.

In conventional NLP methods, embedding itself is enough. However, due to SNN’s discrete nature of input, these embedded feature vectors cannot be processed using spiking architecture as it requires a spike train input. To solve this problem, the embedded vectors are binarized and the vectors are converted into vectors that only consist of 0s and 1s. There are 2 widely used methods to obtain these binarized vectors. They are:

**A) Binarization Methods:** There are several proposed methods for binarization like the ‘Near-lossless’ [25] binarization method. However, most of these methods use an autoencoder to convert embedded vectors into vectors of 0/1. Which, by itself, is not that energy consuming considering it only needs to be done once. However, when it comes to Sentences, binarization requires use of Transformer [11] models. In such a case, the energy consumption of a model using both autoencoders and Transformer models is not ideal for a paper like this since the entire point of using the Spiking Neural Structure was to optimize energy consumption.

**B) Rate-Coding:** In this method, floating-point numbers are easily converted to a series of spikes based on a Poisson process, where, for a number of time-steps  $k$ :

$$Y_i = \begin{cases} 1 & \text{if } X_i(t) < p_i \\ 0 & \text{if } X_i(t) \geq p_i \end{cases} \quad (2.1)$$

**Given  $\mathbf{x} \sim \text{Binomial}(\mathbf{k}, \mathbf{p})$**

This process is repeated for all items in the input vector  $\mathbf{v}$ , generating a spike train for each of the input vector’s feature dimensions [43].

For the purpose of this paper, we looked at several different rate-coding methods until [43] proposed its own rate-coding technique that was appropriate for our energy optimization task. Some rate-coding techniques as well as [43]’s proposed “SNN-rate” has been briefly discussed below:

- ANN - This is a typical ANN using regular floating-point embeddings
- SNN-rate-rand - This is a Poisson rate-coded embedding being used in a downstream SNN
- SNN-bin - This is a binarized embedding being used in a downstream SNN
- SNN-rate [43]: Poisson rate-coded embedding method that generates deterministic spike trains.

**SNN-rate:** As this is the encoding system we will use for this paper, we need more details on how it works. This method aims to inherently modify the stochastic Poisson rate-coding method (referred to as “SNN-rate-rand”) which generates spike trains that approximate the float-point input values stochastically and instead generates deterministic spike trains for the given embedding. The paper [43] generates deterministic encoding as follows:

$$V_i(t) = V_i(t - 1) + x_i(t) \quad (2.2)$$

$$Y_i(t) = \begin{cases} 1 & \text{if } V_i(t) \geq \alpha_i \\ 0 & \text{otherwise} \end{cases} \quad (2.3)$$

Where  $x_i$  is the floating-point value being encoded,  $V_i(t)$  is the membrane potential of the neuron at time-step  $t$ , and  $\alpha_i$  is the membrane threshold of the neuron.

This deterministic encoding, used in downstream SNNs, ensures higher fidelity spike trains, leading to improved model performance compared to traditional rate-coding methods. The overall goal is to maintain high accuracy while leveraging SNN’s energy efficiency benefits in comparison to conventional neural network models.

### 2.2.2 Spike-Timing-Dependent Plasticity (STDP) and its Variants

STDP is an algorithm where the time difference between pre- and post-synaptic neurons determine the strengthening or weakening of a synapse. Traditionally, the equation of STDP is as follows:

$$\Delta t = t_{\text{post}} - t_{\text{pre}}, \quad (2.4)$$

$$W(\Delta t) = \begin{cases} A^+ e^{-\Delta t/\tau^+}, & \text{if } \Delta t \geq 0, \\ -A^- e^{\Delta t/\tau^-}, & \text{if } \Delta t < 0, \end{cases} \quad (2.5)$$

$$\Delta w = \sum_{t_{\text{pre}}} \sum_{t_{\text{post}}} W(\Delta t). \quad (2.6)$$

Here,  $w$  is the synaptic weight.  $\Delta w$  is the change of the synaptic weight.  $t_{\text{pre}}$  and  $t_{\text{post}}$  stand for the timing of the firing spike from pre-neuron and post-neuron.  $A^+$  is the magnitude of the weight change when we have a positive correlation, i.e. the pre-synaptic neuron fires prior to the post-synaptic neuron and  $A^-$  is the magnitude when the correlation is negative.  $\tau^+$  and  $\tau^-$  are time constants defining the width of the positive and negative learning window. [12]

A simple image classification problem which used spike coding to represent images and STDP to update weights, found that their model offered 20 times faster computing speed while retaining accuracy when compared to traditional spiking neuron models [7]. One of the most common ways for training artificial neural networks (ANNs) is the backpropagation algorithm but it does not resemble how humans learn and requires a large amount of labeled data. On the other hand, despite reducing computational and memory complexity, SNNs trained solely with STDP tend to be suboptimal due to local optimization and neuronal coordination issues. Research exists on a number of adaptive mechanisms to improve STDP such as introduction of a dynamic firing threshold, removal of noisy inputs, lateral inhibition and STB-STDP which batches both samples and time-steps. Although the accuracy achieved by this model on the CIFAR10 dataset was relatively low compared to supervised learning algorithms, it was remarkable for a solely unsupervised method [35]. Another method, called sym-STDP, utilizes the concept of symmetric STDP where synaptic changes occur regardless of the order of the firing neuron. It is a biologically plausible, high-performance supervised learning framework that uses spike-based learning rules and achieved high accuracy on the MNIST dataset [28]. Therefore, biologically plausible mechanisms that improve STDP are valuable assets for both supervised and unsupervised learning.

In an effort to modify STDP, a reinforcement learning rule, R-STDP was introduced which was inspired by the role of neuromodulators such as Dopamine (DA) or Acetylcholine (Ach) in the modulation of STDP. Dopamine increases the plasticity of the eligible synapses. All eligible synapses are either working against obtaining the reward, toward it, or the activity is unrelated to the reward. If they systematically work against the behavior being rewarded, STDP will weaken these synapses. If they work toward the goal, the effect of STDP will be to strengthen them. Any uncorrelated activity will lead to a very small change in synaptic strength [2]. [20]

created a 3 layer DCSNN where each convolutional layer was followed by a pooling layer. For every decision made, it was compared with the original label. Accordingly, a reward or punishment signal was generated that would be passed back to the layers that implemented R-STDP. This study achieved a high accuracy on the MNIST dataset and also predicted that R-STDP can be useful when the dataset is populated with frequent distractors. Another study [29], utilized R-STDP in medical image classification to address the issue of class imbalance. Higher rewards were provided to the minority class whereas lower rewards were provided to the majority class which significantly improved the sensitivity of disease identification. Additionally, many hierarchical SNNs with R-STDP learning rules have been used to solve object recognition tasks and achieved notable results [24], [13]. Hence, while the dopamine mechanism has been extensively used across various computer vision tasks, SNN itself remains significantly underused in the domain of NLP. [27] For the task of sentiment classification, a simple reward signal can be computed which rewards correct classification and punishes wrong classification. The weight of the synapses can then be updated by R-STDP accordingly. 2.1 illustrates how R-STDP affects learning in spiking neural networks. Part (a) shows the timing of spikes from pre- and post-synaptic neurons. Part (b) shows the eligibility trace, which increases or decreases based on spike timing and decays over time. Part (c) shows external reward or punishment signals and part (d) shows changes in synaptic weight, which increase after a reward and decrease after a punishment. Overall, the image demonstrates how spike timing combined with reward signals drives synaptic plasticity and learning.

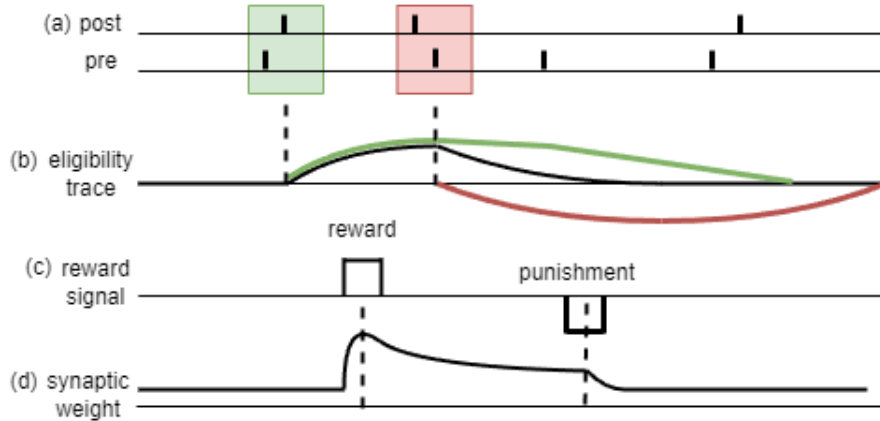


Figure 2.1: How R-stdp affects learning

### 2.2.3 The Forward-Forward Algorithm

The FF algorithm was introduced by [33], which is a biologically inspired training paradigm that eliminates the use of backpropagation. Instead of using traditional backward gradient propagation techniques, it employs two forward passes, one with real (positive) and the other with corrupted (negative) data. Each of the layers updates the weight based on a local system of “goodness” points, typically a function that uses the sum of squared activations.

$$G^{(l)} = \sum_i \left(a_i^{(l)}\right)^2 \quad (2.4)$$

Here,  $a_i^{(l)}$  denotes the activation of the  $i$ -th neuron in layer  $l$ . This metric serves as a local, unsupervised signal that quantifies how strongly the layer responds to the input.

$$\Delta W^{(l)} \propto G_{\text{pos}}^{(l)} - G_{\text{neg}}^{(l)} \quad (2.5)$$

Where  $G_{\text{pos}}^{(l)}$  and  $G_{\text{neg}}^{(l)}$  represent the goodness scores from the positive and negative inputs, respectively. This formulation allows each layer to locally assess whether its representation favors the real data over the corrupted counterpart, enabling learning without the need for backpropagation.

[36] extended FF algorithm to NLP, specifically for sentiment analysis on the IMDb dataset. They trained the FF networks without backpropagation using Word2Vec embeddings and embedded label representations. They introduced a scheme called pyramidal threshold optimization scheme, where deeper layers had higher loss thresholds, simulating the hierarchical abstraction of the human brain. This network then achieved a comparable performance to backpropagation-trained models with an improvement of up to 11% in accuracy under optimized threshold scheduling. Their study also proved that unbounded activations like ReLU were more compatible for the FF algorithm than sigmoid or tanh.

Based on the foundation built by [33], [46] adapted the FF Framework for the spiking neural architectures by using LIF neurons and spike count-based activations. They calculated the “goodness” from the squared spike count of each neuron across 2 forward passes. They also introduced surrogate gradient approximations and deterministic spike encoding, which allowed the model to train effectively without the use of traditional backpropagation. The FF-SNN was tested upon various datasets like MNIST, Fashion-MNIST, CIFAR-10, SHD and N-MNIST which demonstrated that the accuracy was competitive with conventional backpropagation methods. Its compatibility with neuromorphic hardware and local weight update rules makes the FF-SNNs more suitable for low-powered AI systems.

## 2.3 Key Findings

Both the methods, R-STDP and FF Algorithm offer a biologically plausible alternative to gradient backpropagation while targeting different facets of neural computation. R-STDP mainly focuses on the synaptic plasticity modulated using reward signals and excels in dynamic learning scenarios, like in vision. On the other hand, FF provides a structured, gradient-free framework that can be used for both ANN and SNN settings while achieving competitive accuracy with energy-efficient operations. When these two methods are combined with novel spike-train encoding strategies, they present promising paths for the development of brain-like neural architectures that are capable of learning from both images and texts without having to sacrifice computational efficiency.

# Chapter 3

## Methodology

### 3.1 Top Level Overview

In order to optimize sentiment analysis and affective computing and compare the results with existing data, we first have to implement the proposed model itself.

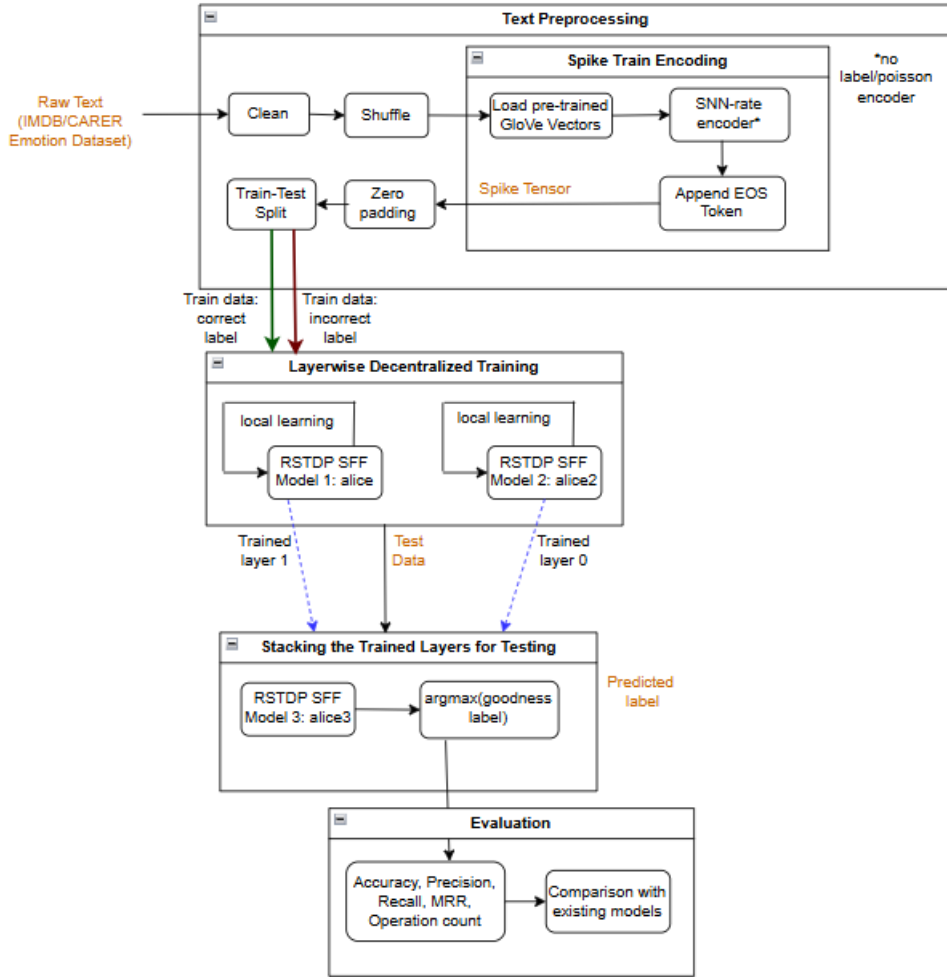


Figure 3.1: Top Level Overview of the Proposed System

Figure 3.1 illustrates the complete operational pipeline of a biologically inspired

SNN model designed for sentiment and emotion classification in NLP. The system replaces traditional Poisson and label encoders with a deterministic spike-based input encoding mechanism (SNN-rate), converting GloVe embeddings into spike trains. It uses decentralized layer-wise training using R-STDP and FF learning paradigm. Two model clones, *alice* and *alice2*, independently train the second and first hidden layers, respectively. Their learned weights are then merged into *alice3* for final inference using spike-based “goodness” scores. The model’s performance is assessed using accuracy, MRR, precision, recall, and energy efficiency, and benchmarked against conventional ANN and SNN baselines.

Our plan has 5 distinct stages:

1. **Dataset Preprocessing:** At this stage, the input data is preprocessed to make it suitable for our analysis. This part is mainly focused on formatting the data and features.
2. **Processing of Input:** Most NLP tasks require further processing of the data. In addition, due to using a SNN model, the input needs to be further processed into discrete values instead of floating-points/actual values. Therefore, this particular stage of processing the data can be divided into several sub-stages:
  - **Splitting Dataset:** This involves separating the portion of the dataset the model will be trained on and the portion that will be used for testing and prediction.
  - **Embedding:** Converting both training and testing datasets into feature vectors.
  - **Rate-Coding:** Converting the embedded vectors into vectors consisting of discrete values (0/1), which can now be considered a spike train. For this part, we are using [43]’s custom ‘SNN-rate’ for rate-coding. This will be done for both training and testing datasets.
3. **Training:** In this stage, we implement the proposed reinforcement learning technique (R-STDP) to train the model on the training dataset.
4. **Prediction:** This stage focuses on using the proposed model to predict the classifications of the test dataset that we split earlier.
5. **Analysis and Comparison:** Finally, we analyze and compare the obtained results (e.g., accuracy, energy consumption) with different models that performed the same task.

## 3.2 Datasets

We used two benchmark datasets for the purpose of sentiment and affective computation analysis.

**IMDB Dataset:** It is a binary sentiment classification dataset with 50,000 entries. Given a string of text from a movie review, the target is to decide whether the text is expressing positive or negative sentiment [1].

CARER Dataset: It is a multi-class emotion dataset with 20,000 entries. Given a tweet as text, the target is to decide which of the 6 emotions- sadness, joy, love, anger, fear, or surprise, are being depicted in the text [19].

### 3.3 Data Preprocessing

In order to test our proposed SNN model, it is vital that we use datasets that are widely available and frequently used for benchmark testing. The datasets themselves have been explained in the Dataset section. However, in order to use these benchmark datasets, data preparation is necessary. The entire process of data preparation/ data preprocessing is explained below.

The data loading mechanism is invoked by the training module. In this stage of the preprocessing, the datasets library is used to load the pre-existing dataset. A custom clean-text function is used to change the text into lowercase, remove underscores, whitespaces and punctuations. The dataset is then randomly shuffled to ensure random distribution of output labels.

At this stage, the training module calls on the dataset preprocessing method which does 5 things:

- Initializes pre-trained GloVe embeddings.
- Since our model is spike event driven, just embeddings are not usable as inputs. In this regard, we use a slightly altered version of SNN-rate, the optimized deterministic rate-coding mechanism that converts these embeddings into a spike train. For more information on SNN-rate and how it works, we can refer to [43].
- Each token in the sentence is embedded into a binary vector. Bit positions with 1s indicate active neurons. These binary vectors are then encoded as spike trains over time using a configurable repetition factor and spike frequency.
- A spike train is constructed for each input sample with shape  $[T, \text{embed\_dim}]$ , where each word contributes to a specific time slice. An End-Of-Sequence (EOS) token is appended to ensure a fixed-length temporal window.
- All samples are padded to a uniform temporal dimension (512 for IMDB, 32 for CARER) and returned as `torch.Tensor` objects ready for spiking network training.

The spike trains are then split into 70% train, and 30% test sets.

| Dataset | Train  | Test   | Total  |
|---------|--------|--------|--------|
| IMDB    | 35,000 | 15,000 | 50,000 |
| CARER   | 14,000 | 6,000  | 20,000 |

Table 3.1: Dataset Split Distribution

This split dataset is finally returned back to the training module where it goes through further processing.

## 3.4 Proposed SNN model Architecture (SFF-RSTDP)

The architecture proposed in our work is a biologically-inspired SNN built on the SFF paradigm. It replaces traditional backpropagation with a decentralized, layer-wise R-STDP learning mechanism, along with surrogate loss calculation. This approach enables parallel and independent training of layers and aligns with the sparse, energy-efficient principles of neuromorphic computing. The structure of our model and its flow of operation is explained below.

### 3.4.1 Model Structure

The main spiking network, termed as *NetofLain* is a fully connected network where each layer is a spiking neural layer based on a custom class called *LayerofLain*. The overall architecture, shown in figure 3.2, has 3 layers. Our model architecture was inspired and influenced by [41]

- **Input Layer:** Accepts time-sequenced spike trains of shape [T, batch size, embed dimension], where each token in a sentence contributes to a portion of the temporal spike train. The dimension of this layer changes depending on the embedded dimension of the input spike trains.
- **Two hidden layers:** Each layer is defined as a fully connected linear spiking layer (*spikingjelly.layer.Linear*) followed by an Integrate-and-Fire neuron (*spikingjelly.neuron.IFNode*) with surrogate gradients (ATan). After fine tuning, each layer has a dimension of 500.
- **Output logits:** Since classification is made by calculating aggregating spike-based "goodness" across time and layers, our model does not have an output layer like conventional neural network models.

### 3.4.2 Parameters

For the parameters, we set the decay for pre-synaptic activity to 2.0. The decay for post-synaptic activity is 100.0. The learning rate is 0.003 for the first layer and 0.004 for the second layer. Both the positive and negative reward threshold are 0.005. Maximum timestep is 512 for IMDB and 32 for CARER emotion dataset. Lastly, the batch size is set to 15.

### 3.4.3 Decentralized structure

Our proposed model of network clones the same model twice. Each clone termed as *alice* and *alice2* has the same structure of [input dimension x 500 x 500]. In decentralized learning mechanism, *alice* separately trains only its second layer, using inputs that pass through the first layer. Similarly, *alice2* only trains its first layer using the raw spike encoded inputs. Finally, the learned weights from both the models are stacked into a new model called *alice3* which is an object of *NetofLain* class. This final *alice3* model is used to compute goodness scores across the test data and finally make predictions. Figure 3.3 shows an overview of how each model learns in a decentralized manner.

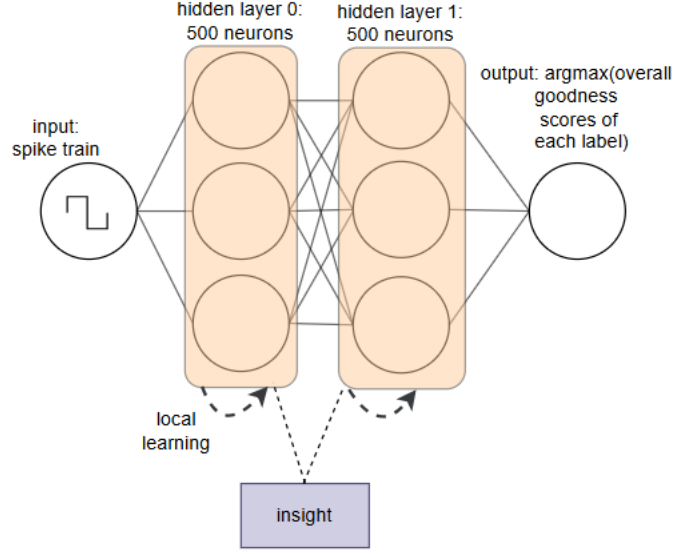


Figure 3.2: SFF-RSTDP Model (Note that the network doesn't output a class directly as it computes goodness per class label and chooses the class with maximum goodness)

### 3.4.4 Forward-Forward Mechanism

To implement the biologically plausible Forward-Forward training paradigm, instead of back propagation, our model corrects itself by forward passing positive and negative samples.

First, it creates 2 tensors which are respectively positive and negative feature-label samples. An example of positive and negative sample generation for a 10 class dataset is given by figure 3.4. These positive and negative samples are passed through the network (forward pass without learning) to produce corresponding spikes. From the obtained output, the goodness and negative goodness of each layer for this specific sample is calculated and the sum of mean goodness of all the layers for each sample is recorded.

The goodness is calculated by:

$$G = \frac{1}{T} \sum_{t=1}^T \left( \frac{1}{N} \sum_{i=1}^N x_i(t)^2 \right) \quad (3.1)$$

Where  $G$  is the goodness score of the layer for one input sample,  $T$  is the total number of time steps,  $N$  is the number of neurons in the layer,  $x_i(t)$  is the membrane potential of neuron  $i$  at time step  $t$ .

Finally, once the goodness measure for the sample is calculated and tracked, our model utilizes the goodness to calculate weight update for that layer only. A high goodness indicates that the layer is confidently excited by the input (positive sample) while a low goodness indicates that the layer is not excited (negative sample).

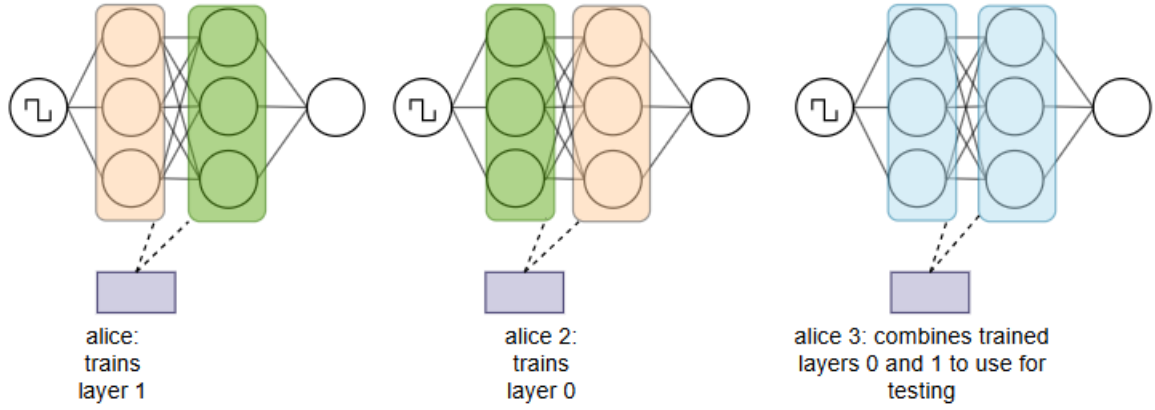


Figure 3.3: Decentralized Layer-wise Training Strategy

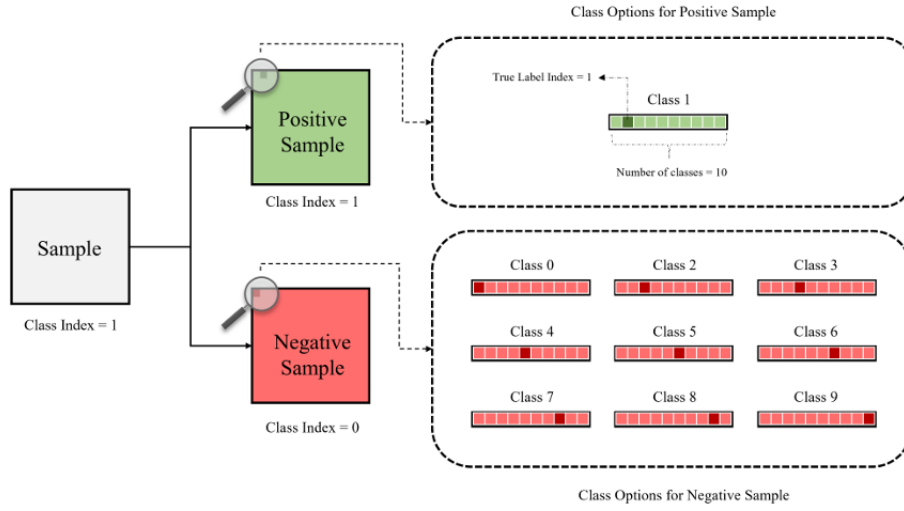


Figure 3.4: Different options for overlaying a label on a sample of a 10-class dataset for positive and negative sample generation [45]

### 3.4.5 Weight Update

In the second forward pass, each layer updates its weights using a combination of Reward-Modulated Spike-Timing-Dependent Plasticity (R-STDP) and a surrogate gradient-based loss function shown by figure 3.5. The mechanism of each is given below.

**Reward-Modulated STDP (R-STDP):** In conventional STDP, if a presynaptic neuron fires shortly before a postsynaptic neuron, the synaptic connection is strengthened; otherwise, it is weakened. In R-STDP, this plasticity rule is modulated by a scalar reward signal.

In our proposed model, the reward is derived from the layer's *goodness*, which is calculated during the initial (non-learning) forward pass. The reward is applied at each time step during the second forward pass and determines the magnitude and

direction of weight updates. It is defined as:

$$\text{Reward}_{\text{pos}} = \alpha_{\text{pos}} \cdot (G - \theta_{\text{pos}}) \quad (3.2)$$

$$\text{Reward}_{\text{neg}} = \alpha_{\text{neg}} \cdot (\theta_{\text{neg}} - G) \quad (3.3)$$

Where:

- $G$ : goodness score of the current layer
- $\theta_{\text{pos}}, \theta_{\text{neg}}$ : threshold parameters for positive and negative samples
- $\alpha_{\text{pos}}, \alpha_{\text{neg}}$ : reward scaling factors

If the layer fires strongly for the wrong input ( $R < 0$ ), synapses are punished (weakened). If it fires strongly for the correct input ( $R > 0$ ), synapses are reinforced (strengthened). These updates are handled by the *MSTDPLearner* object from the *SpikingJelly* library.

**Surrogate Gradient Loss:** To complement the STDP-based updates, a surrogate gradient-based loss is also used to optimize weights via gradient descent. This loss ensures separation between the positive and negative views by comparing their combined goodness values.

A collaborative signal, called *insight*, is computed as the sum of the goodness from all other layers (excluding the one currently being trained). For a layer  $i$ , the combined scores are:

$$\text{Combined}_{\text{pos}} = G_i^+ + \sum_{j \neq i} G_j^+ \quad (\text{insight\_pos}) \quad (3.4)$$

$$\text{Combined}_{\text{neg}} = G_i^- + \sum_{j \neq i} G_j^- \quad (\text{insight\_neg}) \quad (3.5)$$

- $G_i^+$ : Goodness score of the current layer  $i$  for a positive (correctly labeled) input.
- $G_i^-$ : Goodness score of the current layer  $i$  for a negative (incorrectly labeled) input.
- $G_j^+$ : Goodness score of layer  $j \neq i$  for the same positive input.
- $G_j^-$ : Goodness score of layer  $j \neq i$  for the same negative input.

The final loss is calculated using a smooth log-sum-exp function that promotes contrast between the two views.

$$\mathcal{L}_{\text{mixed}} = \log(\exp(\text{Combined}_{\text{pos}}) + \exp(\text{Combined}_{\text{neg}}) + 1) \quad (3.6)$$

Where  $L$  is the surrogate loss. This loss allows for gradient-based fine-tuning using optimizers like Adam, with gradient flow enabled by surrogate activation functions (e.g., ATan).

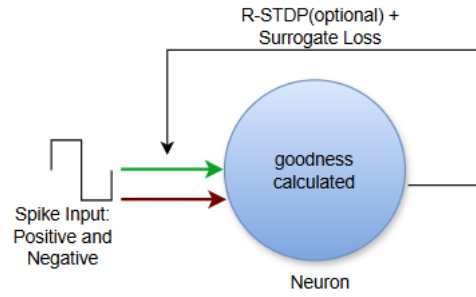


Figure 3.5: Neuron-level learning loop: Each neuron receives spike input from both positive and negative views. It computes its activation goodness (based on membrane potential), then adjusts synaptic weights using R-STDP and surrogate loss to favor correct labels and suppress incorrect ones.

# Chapter 4

## Experimental Results

In order to achieve the objectives of this research paper, some metrics and goals had to be set for proper evaluation of the results. The purpose of our model, when broken down, ultimately strives to achieve the following:

1. **Improvement Over Existing SNN Models:** Enhance the performance of SNN models that rely solely on gradient-based backpropagation. Spiking Neural Networks with non-deterministic rate-encoding were the baseline for this metric. In the result section, these models are referred to as “SNN” and ”SNN-rate”.
2. **Efficiency:** Significantly reduce the total number of operations required to train a model that achieves acceptable levels of accuracy. Existing ANN models were the baseline for this metric. In the result section, this is referred to as “ANN”.
3. **Optimization Without Gradient Backpropagation:** Outperform or match the performance of SNN-rate encoded optimized SNN models while avoiding the complexities and memory constraints associated with gradient-based backpropagation. SNN models optimized through gradient backpropagation and custom rate-encoding technique, as referenced in [43] was considered to be the baseline for this metric. In the result section, this is referred to as “SNN-rate-op”.
4. **Research Gap Reduction:** Contribute to the research on using Reward Modulated STDP (Spike-Timing-Dependent Plasticity) and Forward-Forward mechanism on SNN models for NLP tasks such as sentiment analysis and addressing the current lack of research in this area.

For benchmark testing, the number of epochs for all models except our proposed model were set to 10. All 5 models (ANN, SNN, SNN-rate (Non-deterministic SNN), SNN-rate-op (Deterministic SNN) and SFF-RSTDP [ this paper’s proposed model] were trained on 2 datasets, IMDB and CARER emotion dataset. For comparison, 5 metrics were focused on: Accuracy, mean reciprocal rank (MRR), precision, recall, and total number of operations. How these metrics were calculated have already been explained in details in section 4.2

## 4.1 Training Module

The training module is the core of our system where everything is sequentially called on and executed. The entire flow of operation of the training module is explained below:

The training module can be separated into 5 operations:

1. Calls on the Dataloader to retrieve the appropriate dataset (IMDB or CARER) and split the dataset as required (Explained in Dataset section). Finally return the training and validation data.
2. Invokes the SNN-rate encoding mechanism and obtains the optimized rate-encoded embeddings which will be used as the input for training the Network.
3. Sets the amount of neurons in the first layer of the SNN model based on the size of the input embedding.
4. Training Loop: Loops through the training process for 1 epoch.
5. The training loop runs independently for each clone of the network (*alice* and *alice2*). For each input, the model creates positive and negative samples by combining the same input spike with the correct and incorrect labels. Both samples are passed through the network in a forward pass (without learning). From the output spikes, the model calculates the goodness for each layer. The average goodness from the preceding layers is also recorded and returned as "insight". Reward is computed and based on this reward, the model applies R-STDP updates (if enabled) and computes surrogate loss to update the synaptic weights locally for the active layer. The first clone, *alice*, trains the second hidden layer while keeping the first layer frozen. The second clone, *alice2*, trains the first layer using raw input spikes, keeping the second layer inactive. After training, both layers are combined into a new model named *alice3*, which is used for inference and classification.

## 4.2 Analysis Metrics

During the entire training process, some metrics are calculated by the training module which is later used for model analysis. They are as follows:

1. Accuracy: Accuracy is calculated by:

$$\text{Accuracy} = \frac{\text{Number of Correct Predictions}}{\text{Total Number of Predictions}} \quad (4.1)$$

2. Mean Reciprocal Rank: The ability of the model to rank the correct labels among all possible labels. It is calculated by:

$$\text{MRR} = \frac{1}{N} \sum_{i=1}^N \frac{1}{\text{rank}_i}, \quad (4.2)$$

where  $\text{rank}_i$  is the position of the correct label in the sorted prediction scores for sample  $i$  and  $N$  is total sample.

3. Precision and Recall:

Both metrics are computed per class and then averaged (e.g., macro-average) as done using `sklearn.metrics`.

**Precision (for class  $c$ ):**

$$\text{Precision}_c = \frac{TP_c}{TP_c + FP_c} \quad (4.3)$$

**Recall (for class  $c$ ):**

$$\text{Recall}_c = \frac{TP_c}{TP_c + FN_c} \quad (4.4)$$

**Where:**

- $TP_c$  = True Positives for class  $c$
- $FP_c$  = False Positives for class  $c$
- $FN_c$  = False Negatives for class  $c$

**In macro-averaging:**

$$\text{Precision}_{\text{macro}} = \frac{1}{C} \sum_{c=1}^C \text{Precision}_c, \quad \text{Recall}_{\text{macro}} = \frac{1}{C} \sum_{c=1}^C \text{Recall}_c \quad (4.5)$$

4. Confusion Matrix: The confusion matrix is a tabular representation:

$$\mathbf{M}_{i,j} \in \mathbb{Z}^{C \times C}$$

where:

$$M_{i,j} = \text{Number of samples with true label } i \text{ and predicted label } j \quad (4.6)$$

This allows insight into which classes are being confused with others.

5. Average Spikes Count: Calculated by:

$$\text{Average Spikes}_l = \frac{1}{B} \sum_{b=1}^B S_{l,b} \quad (4.7)$$

where  $l$  is the layer index,  $B$  is the total number of batches (or samples),  $S_{l,b}$  is the number of spikes produced by layer  $l$  on batch  $b$ .

6. F-1 score:

$$F_1 = 2 \cdot \frac{\text{Precision} \cdot \text{Recall}}{\text{Precision} + \text{Recall}}$$

## 4.3 Analysis of proposed model's results

The proposed model, referred to as ‘SFF-RSTD’ from here on, gave overall positive results when tested on the 2 provided datasets. It is to be mentioned that the data provided below is not of one specific run but average of multiple rounds of training. This applies for all the models mentioned from here on.

| Dataset          | Accuracy | MRR   | Precision | Recall | F1 Score |
|------------------|----------|-------|-----------|--------|----------|
| CARER<br>emotion | 0.485    | 0.713 | 0.468     | 0.624  | 0.535    |
| IMDB             | 0.732    | 0.866 | 0.747     | 0.732  | 0.739    |

Table 4.1: SFF-RSTDP Performance Metrics

### 4.3.1 Initial Results

The table 4.1 shows the initial performance of our proposed model in all the relevant metrics. The significance of these values have been discussed in details for each respective dataset in section 4.3.2 and 4.3.3

### 4.3.2 CARER Emotion Dataset

For the CARER emotion dataset, SFF-RSTDP achieved an accuracy of 48.5% and MRR (Mean Reciprocal Rank) of 71.3% while having a precision and recall score of 0.47 and 0.62 respectively.

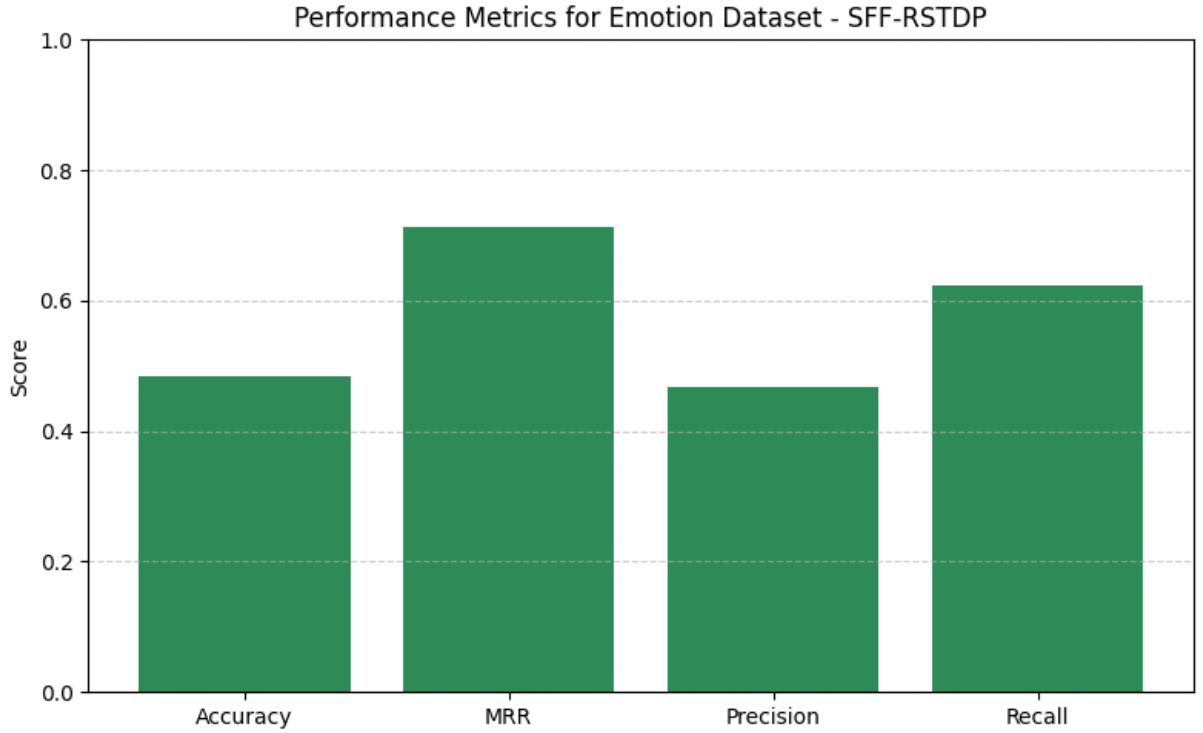


Figure 4.1: Accuracy, MRR, Precision and Recall data for CARER dataset

The figure 4.1 above states 3 things. Firstly, the model can achieve exceptionally high MRR for an SNN model. Secondly, the proposed model almost achieves an accuracy close to 50% which is a notable achievement when being trained on a multiclass dataset. Lastly, the precision and recall scores are moderate yet consistent with the accuracy achieved by the model. The confusion matrix given below will attest to that fact in figure 4.2

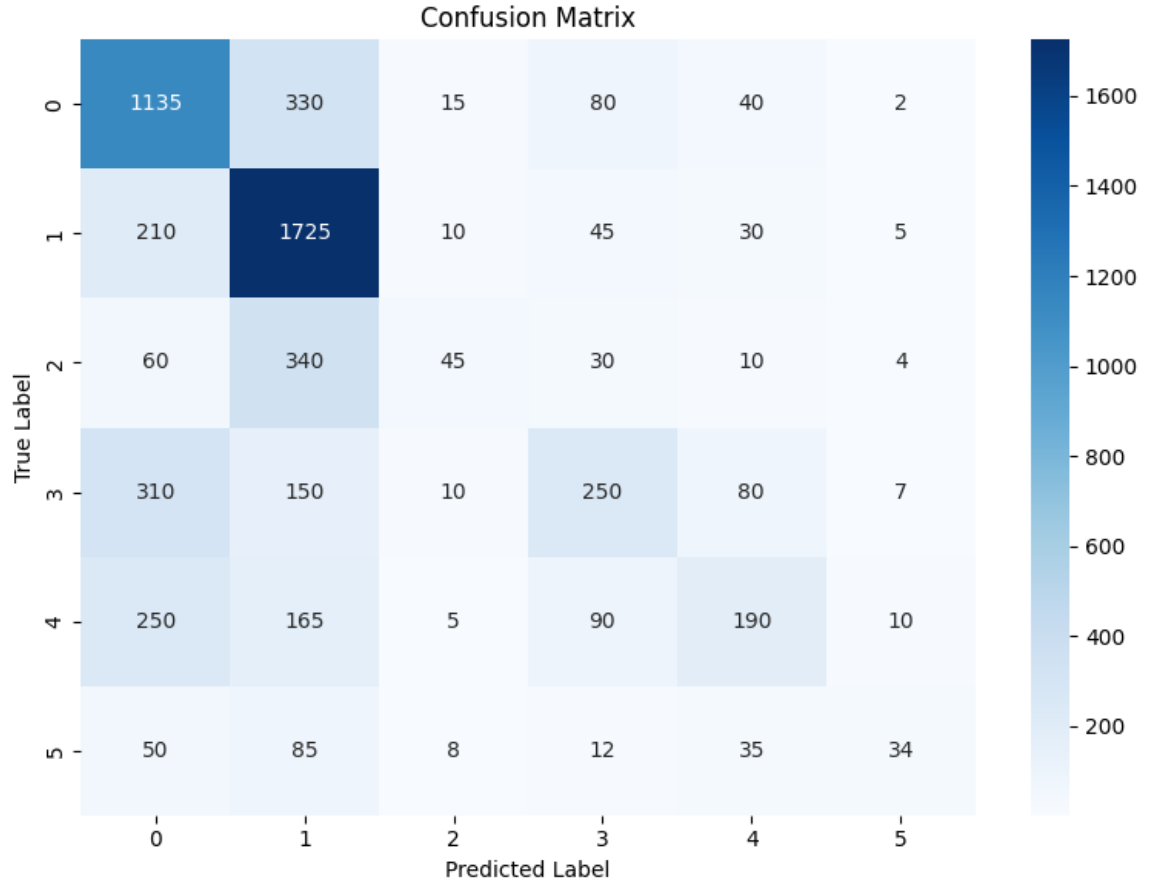


Figure 4.2: SFF-RSTD P Confusion Matrix for CARER Dataset

### 4.3.3 IMDB dataset

For the IMDB dataset, SFF-RSTD P achieved an accuracy of 73.2% and MRR (Mean Reciprocal Rank) of 86.6% while having a precision and recall score of 0.74 and 0.73 respectively.

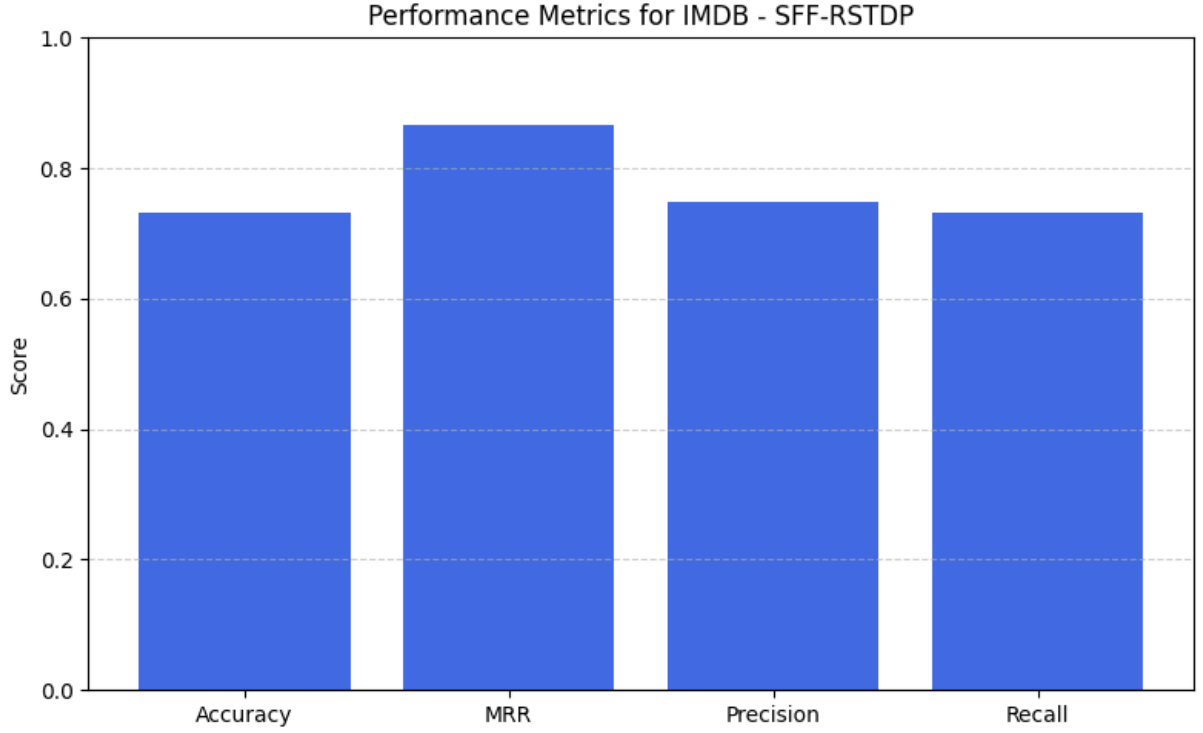


Figure 4.3: Accuracy, MRR, Precision and Recall Data for IMDB dataset

In the case of the IMDB dataset, graphical representation is somewhat similar to CARER dataset. The figure 4.3 given above indicates the following, that similar to the CARER dataset, the model achieves a very high MRR score with the IMDB dataset. This demonstrates that our model's ability to rank the correct label is exceptionally high for an SNN model. Unlike the CARER dataset, the IMDB dataset is binary-classed. Consequently, it is natural for the model to achieve higher accuracy scores. The figure 4.3 demonstrates that successfully. Also, unlike the CARER dataset, the precision score of 0.74 is actually slightly higher than its recall score of 0.73. In addition, both the precision and recall score is comparatively high in this case, which results in a more balanced confusion matrix as shown in figure 4.4

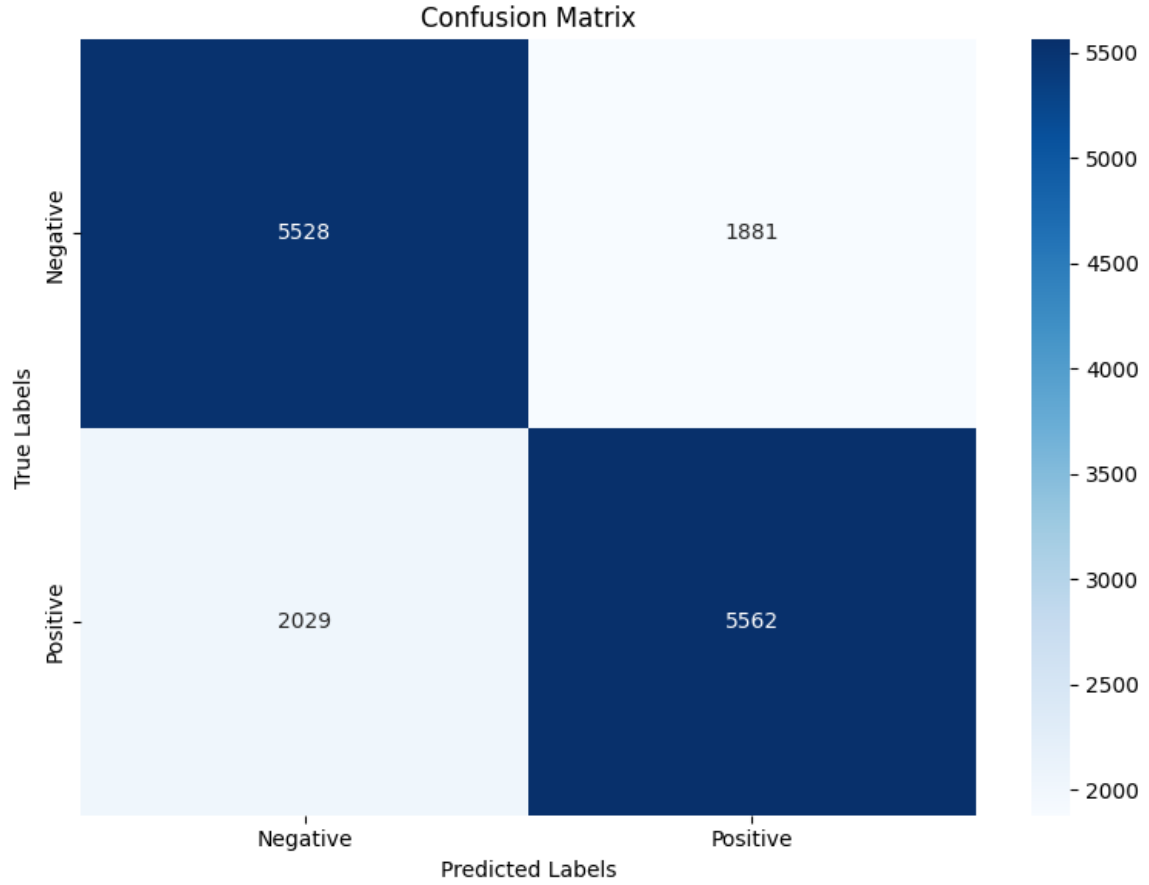


Figure 4.4: SFF-RSTD P Confusion Matrix for IMDB Dataset

## 4.4 Comparison with other models

Given below are the tabulated results for all 5 models and their corresponding graphs. The other SNN models and their benchmark numbers were collected from [37]

| Model       | Accuracy | MRR   | Precision | Recall | F1 Score |
|-------------|----------|-------|-----------|--------|----------|
| ANN         | 0.567    | 0.732 | 0.462     | 0.354  | 0.401    |
| SNN         | 0.340    | 0.570 | 0.290     | 0.350  | 0.317    |
| SNN-rate    | 0.380    | 0.570 | 0.300     | 0.510  | 0.377    |
| SNN-rate-op | 0.420    | 0.608 | 0.350     | 0.560  | 0.429    |
| SFF-RSTD P  | 0.485    | 0.713 | 0.468     | 0.624  | 0.535    |

Table 4.2: CARER Emotion Dataset Results

The table 4.2 above contains the 4 mentioned metrics for each of the model when the CARER dataset is used. The contents of the table has been analyzed below:

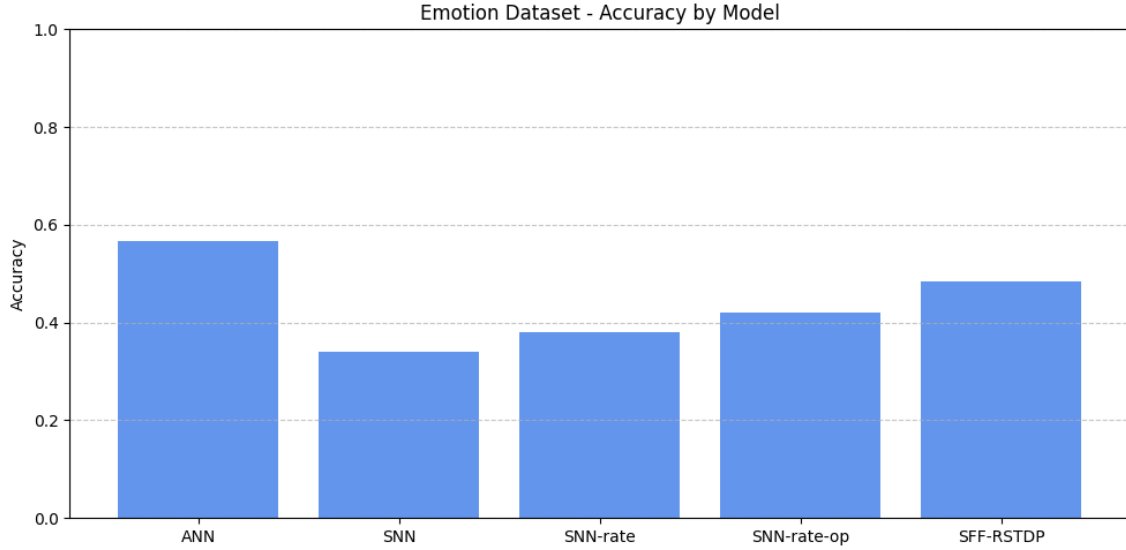


Figure 4.5: Accuracy data for CARER

Figure 4.5 shows that SFF-RSTDTP (48.5% ACC) managed to beat all the other SNN based models [SNN (34% ACC), SNN-rate (38% ACC)] and SNN-rate-op (42% ACC) while failing to compete with the ANN model (56.7% ACC).

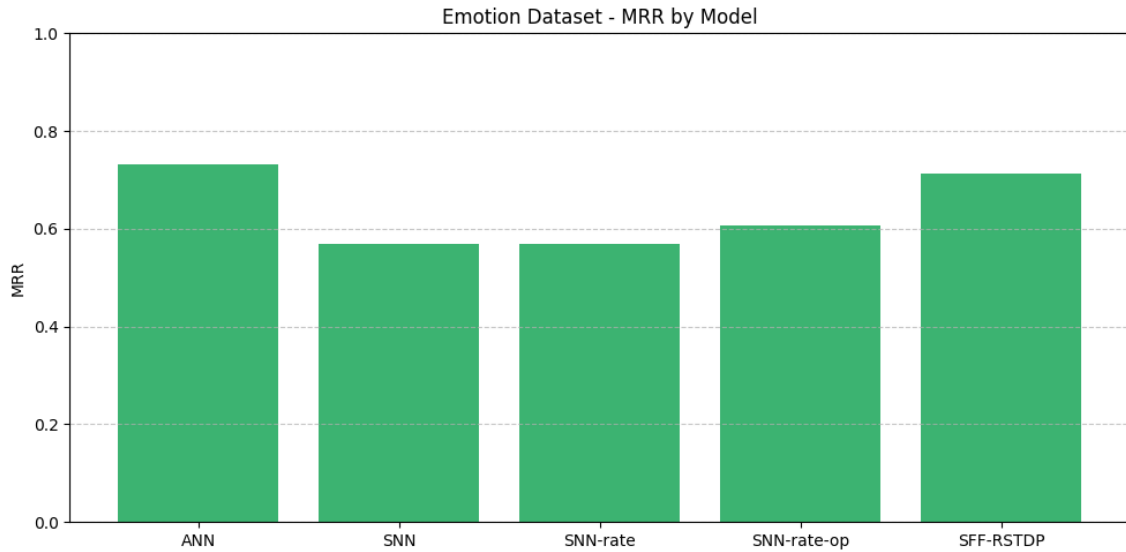


Figure 4.6: MRR for CARER

From the figure 4.6 above, we can deduce a pattern that while our proposed model SFF-RSTDTP (71.3% MRR) consistently beats conventional SNN (57% MRR), SNN-rate (57%) and SNN-rate-op (61.73% MRR), it can't match the results of a conventional ANN model's accuracy or MRR (73.2% MRR) despite getting close to within 2% in this metric of MRR.

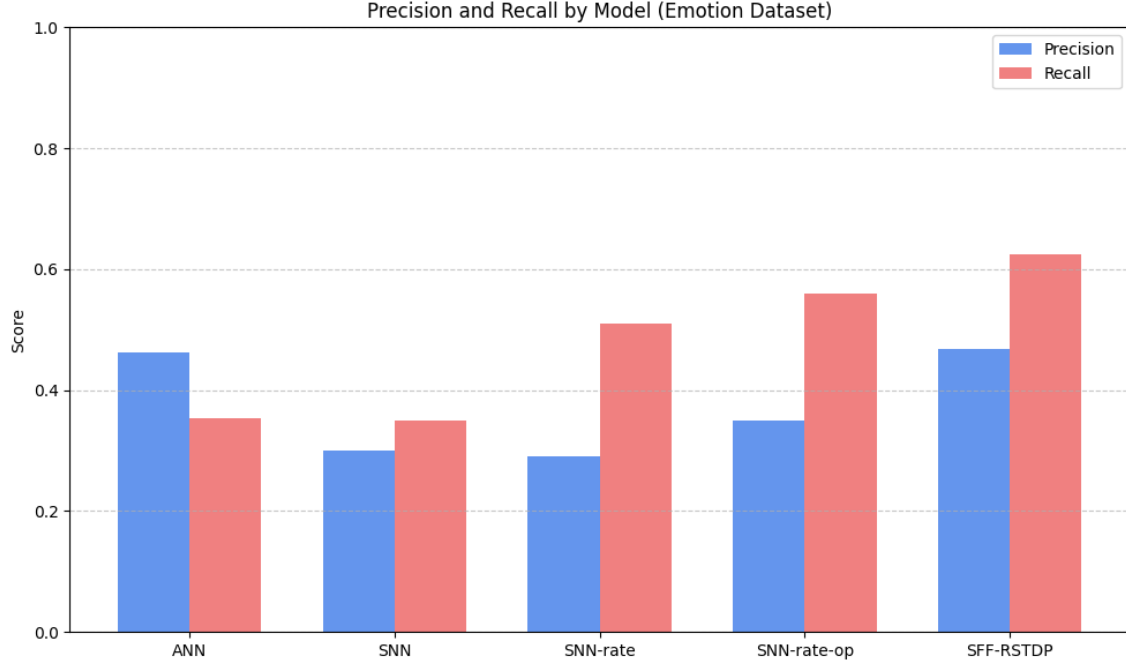


Figure 4.7: Precision and Recall Score for CARER

The figure 4.7 demonstrates 2 things. Firstly, it demonstrates SFF-RSTDTP’s superior ability to be correct when it makes a positive prediction (Precision). Secondly, it shows our models ability to find all the correct positives (Recall). Compared to the values of other models SNN (0.29,0.35), SNN-rate (0.3,0.51), SNN-rate-op (0.35, 0.56) and ANN (0.46, 0.35), only ANN comes close to beating SFF-RSTDTP.

Table 4.3: Number of Operations for Each Model (CARER Dataset)

| Model       | Total Operations |
|-------------|------------------|
| ANN         | 84,736           |
| SNN         | 1,487            |
| SNN-rate    | 1,506            |
| SNN-rate-op | 1,471            |
| SFF-RSTDTP  | 10,726           |

At a glance, the table 4.5 shows that SFF-RSTDTP easily beats ANN in terms of reducing energy consumption while getting significantly beaten by the other SNN models. However, the data presented above can be slightly misleading at initial judgement. The SFF-RSTDTP and the other SNN models have significant difference in neuron structure and the time steps required to train the model. Due to time constraints and lack of machine resources, the SFF-RSTDTP model was run for only one epoch. To compensate for the reduced learning, the time-step applied by the model was 512 compared to 10 of the other models, which significantly bloated the number for average number of operations. However, for a fair comparison, the SNN-rate-op model’s time step was increased to 512 for one epoch and the result we obtained was 98,220 which is significantly higher than 10,726. It can be assumed that, if training parameters and conditions were same, the SFF-RSTDTP would beat all the SNN based models.

| Model       | Accuracy | MRR    | Precision | Recall | F1 Score |
|-------------|----------|--------|-----------|--------|----------|
| ANN         | 0.783    | 0.891  | 0.7700    | 0.8000 | 0.784    |
| SNN         | 0.510    | 0.7553 | 0.5000    | 0.5700 | 0.532    |
| SNN-rate    | 0.490    | 0.7470 | 0.4900    | 1.0000 | 0.657    |
| SNN-rate-op | 0.648    | 0.8240 | 0.5900    | 0.9100 | 0.714    |
| SFF-RSTDTP  | 0.7322   | 0.8661 | 0.7473    | 0.7327 | 0.740    |

Table 4.4: IMDB Dataset Results

The table 4.4 above contains the 4 mentioned metrics for each of the model when the IMDB dataset is used. The contents of the table has been analyzed below:

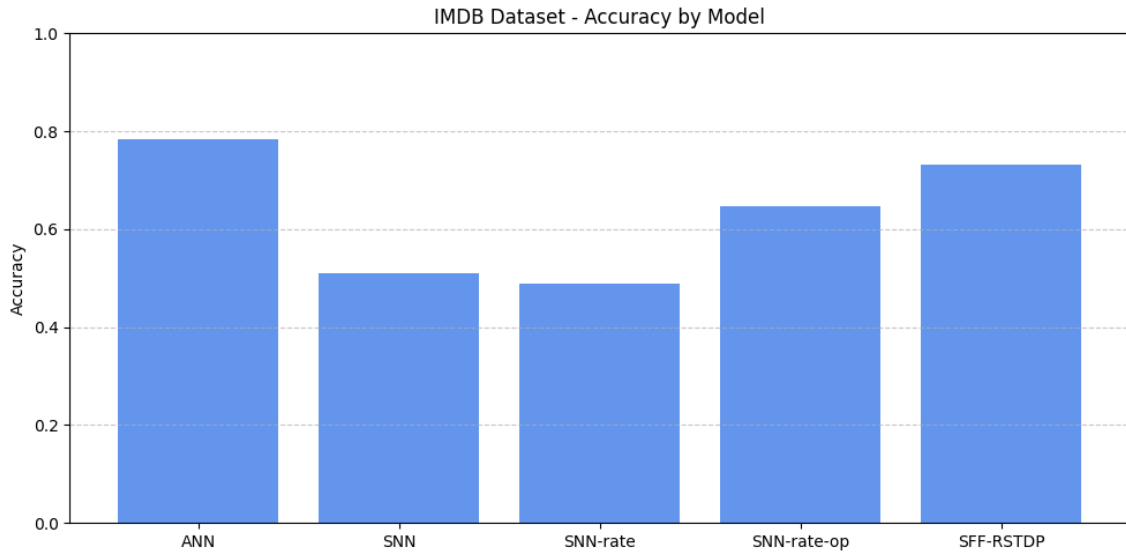


Figure 4.8: Accuracy data for IMDB

The pattern continues in figure 4.8 as we can see ANN models have the highest accuracy (78.3% ACC) while conventional SNN (51% ACC), SNN-rate (49% ACC) and SNN-rate-op (64.8% ACC) struggles to keep up with SFF-RSTDTP (73.2% ACC).

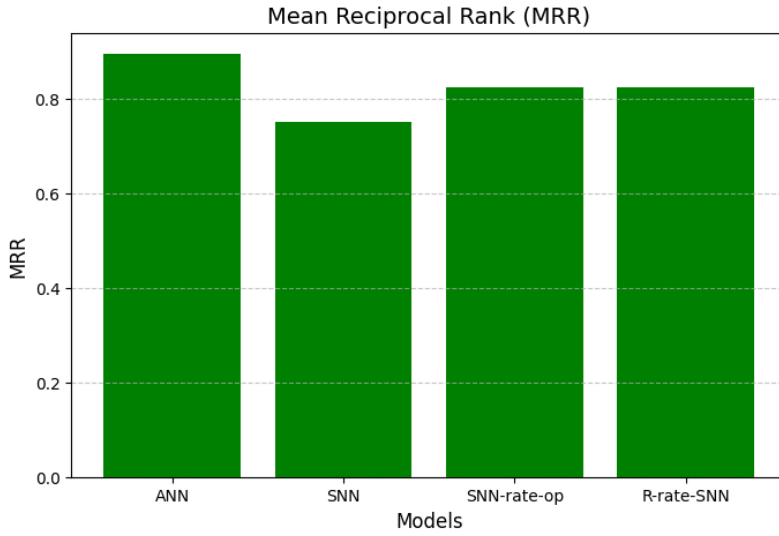


Figure 4.9: MRR for IMDB

Similar to accuracy, figure 4.9 shows that our model (86.66% MRR) is outperformed by ANN (89.1% MRR). However, it manages to beat the rest of the SNN models very easily.

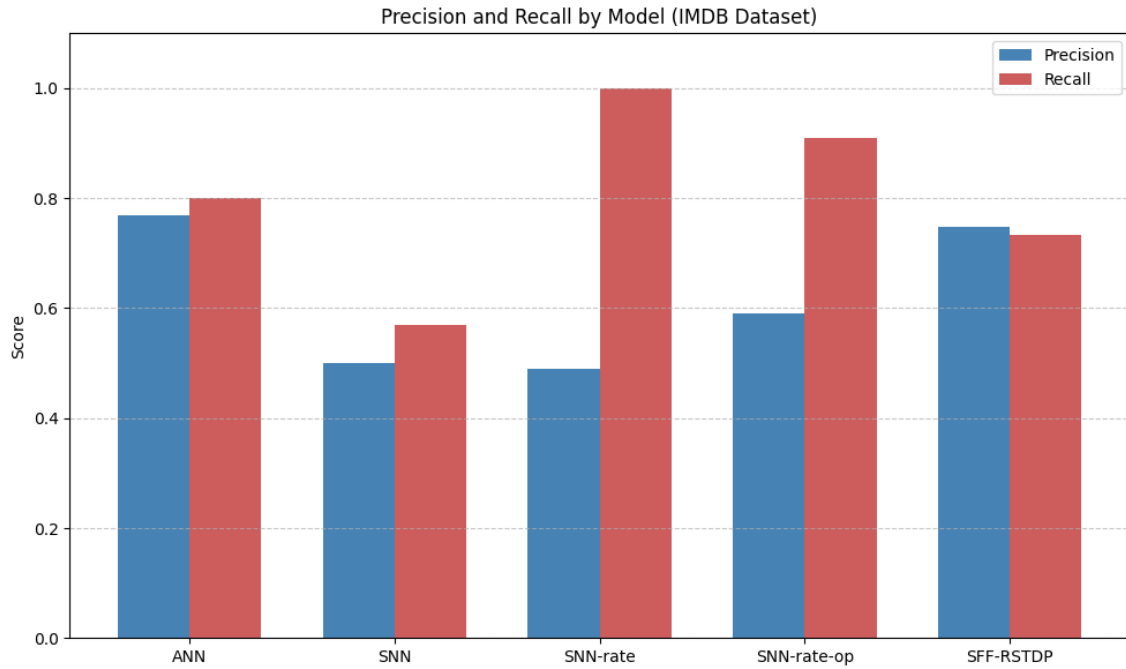


Figure 4.10: Precision and Recall Score for IMDB

Figure 4.10 shows slightly different results compared to the CARER dataset. In case of the CARER dataset, SFF-RSTDTP outperformed all the other models across the board. In this case however, SFF-RSTDTP's precision (0.74) is beat by ANN while its recall score of 0.73 is beat by ANN, SNN-rate and SNN-rate-op.

Table 4.5: Number of Operations for Each Model (IMDB Dataset)

| <b>Model</b> | <b>Total Operations</b> |
|--------------|-------------------------|
| ANN          | 84,244                  |
| SNN          | 1,692                   |
| SNN-rate     | 1,862                   |
| SNN-rate-op  | 1,789                   |
| SFF-RSTDTP   | 15,337                  |

Despite seemingly not being energy efficient enough to compete with other SNN models, theoretically, its use of Forward-Forward algorithm will significantly reduce the total number of operations per sample if similar training parameters and conditions were met.

After extensive analysis it can be concluded that even though SFF-RSTDTP cannot compete with the accuracy and MRR of the ANN model, it significantly reduces the number of operations from 84736 to 10,726 for CARER dataset and from 84224 to 15,337 for IMDB dataset while maintaining respectable scores in the other metrics.

Secondly, compared to non-optimized back propagated SNN and SNN-rate, which are considered to be the most widely used version of the model in current research fields, SFF-RSTDTP outperforms them in almost every metric provided. Despite being a deterministic and optimized version of conventional SNN, SNN-rate-op still struggles to compete with the proposed SFF-RSTDTP model. The only metric it seemingly falls short in is the Average Number of Operations Per Sample. However, in the analysis section, a detailed explanation has been provided to mitigate this tradeoff. Moreover, if the model was deployed on a neuromorphic hardware, theoretically, the energy consumption reduced by SFF-RSTDTP would be much more significant than conventional SNN.

#### 4.4.1 Ablation Study

Table 4.6: Ablation Study Results on IMDB Dataset

| <b>Ablation Setting</b>          | <b>Accu-<br/>racy</b> | <b>MRR</b> | <b>Preci-<br/>sion</b> | <b>Recall</b> | <b>F1<br/>Score</b> |
|----------------------------------|-----------------------|------------|------------------------|---------------|---------------------|
| <b>None(Full<br/>Model)</b>      | 0.732                 | 0.866      | 0.747                  | 0.732         | 0.739               |
| <b>No R-STDTP</b>                | 0.674                 | 0.810      | 0.692                  | 0.683         | 0.687               |
| <b>No Surrogate<br/>Gradient</b> | 0.608                 | 0.750      | 0.631                  | 0.588         | 0.609               |
| <b>Binary<br/>Embedding</b>      | 0.570                 | 0.710      | 0.550                  | 0.660         | 0.600               |

In order to understand the individual contribution of the components within our biologically inspired SNN model, due to resource limitation, we conducted an ablation study on the IMDB dataset only by disabling some of the key mechanisms like

R-STDP, the use of real valued embeddings, and surrogate gradient-based loss. The entire model (SFF-RSTDP) had achieved an accuracy of 73.2%, F1 score of 73.9%, and MRR of 86.6%, while serving as the baseline.

The removal of R-STDP demonstrated a moderated drop in performance with an accuracy of 67.4% and F1 score of 68.7%. This proves the significance of reward-modulated synaptic updates, as there was a loss in the model’s ability to dynamically reinforce the beneficial neuron activations based on the feedback signals. However, the model still retained a certain degree of functionality due the surrogate gradient-based training.

When the surrogate gradient loss was disabled it had a more severe impact, which was an accuracy of 67.4% and a F1 score of 68.7%. This result indicates that while R-STDP supports in guiding the synaptic updates, the surrogate is also very significant in providing consistent gradient-based separation between the positive and negative samples, especially for the deeper layer.

On the other hand, using binary embeddings instead of the real-valued GloVe vectors had extended reduction in the performance with an accuracy of 57.0% and F1 score of 60.0%, which indicates that the reduced semantic resolution of binary encoding leads to a less informative pattern for the spikes. Despite, the model still produced reasonable results, which indicates the robustness of our rate-coding and SNN architecture.

In conclusion, the ablation study validates our design choices, whilst highlighting that every component, R-STDP, real-valued embedding, and surrogate gradient optimization is crucial for achieving a biologically plausible yet a competitive NLP performance as well.

#### 4.4.2 Discussion

Table 4.7: Evaluation Metrics and Operation Count for CARER and IMDB Datasets

| Dataset | Model       | Accuracy | MRR    | Precision | Recall | F1 Score | Total Operations |
|---------|-------------|----------|--------|-----------|--------|----------|------------------|
| CARER   | ANN         | 0.567    | 0.732  | 0.462     | 0.354  | 0.401    | 84,736           |
|         | SNN         | 0.340    | 0.570  | 0.290     | 0.350  | 0.317    | 1,487            |
|         | SNN-rate    | 0.380    | 0.570  | 0.300     | 0.510  | 0.377    | 1,506            |
|         | SNN-rate-op | 0.420    | 0.608  | 0.350     | 0.560  | 0.429    | 1,471            |
|         | SFF-RSTDP   | 0.485    | 0.713  | 0.468     | 0.624  | 0.535    | 10,726           |
| IMDB    | ANN         | 0.783    | 0.891  | 0.7700    | 0.8000 | 0.784    | 84,244           |
|         | SNN         | 0.510    | 0.7553 | 0.5000    | 0.5700 | 0.532    | 1,692            |
|         | SNN-rate    | 0.490    | 0.7470 | 0.4900    | 1.0000 | 0.657    | 1,862            |
|         | SNN-rate-op | 0.648    | 0.8240 | 0.5900    | 0.9100 | 0.714    | 1,789            |
|         | SFF-RSTDP   | 0.7322   | 0.8661 | 0.7473    | 0.7327 | 0.740    | 15,337           |

Following the objectives we set for ourselves in this research, the model we proposed has nearly achieved them all quite successfully. SNNs, especially in the field of NLP, have a long way to go in order to compete with the ANN models. However, the performance of our model effectively shows that it is possible to mitigate that gap in performance with additional use of Reward Modulation and Forward-Forward training paradigms. The table 4.6 further demonstrates the significance of each part of our proposed pipeline and their contribution to enhancing the performance of SNN. In conclusion, the results of our proposed model lays a solid foundation and baseline for improving SNN model’s performance. The final tabulated metric values are shown in table 4.7.

# Chapter 5

## Conclusion

Our research presents the R-STDP SFF model, a biologically inspired, decentralized learning framework that combines Reward-Modulated Spike-Timing Dependent Plasticity (R-STDP) with the Forward-Forward (FF) algorithm for natural language classification tasks. By applying this model to sentiment analysis and affective computing datasets like IMDB and CARER, the study demonstrates a promising intersection of biological plausibility, computational efficiency, and layer-wise local learning without reliance on global backpropagation. The FF algorithm enables each layer to evaluate its output independently using a spike-based "goodness" metric, mimicking cortical-level autonomy in the brain. Coupled with R-STDP's dopamine-modulated reward system, the model learns adaptively through time-sensitive weight updates that reflect synaptic plasticity in biological neurons. The architecture also improves energy efficiency by promoting sparse, binary spiking activity, making it potentially well-suited for future deployment on neuromorphic hardware.

However, the approach comes with several tradeoffs and limitations. The most prominent is the performance gap compared to state-of-the-art artificial neural networks, particularly Transformer-based models. These outperform SNNs in accuracy but lack energy efficient training paradigms. The scope of this study was also limited to specific datasets, which restricts broader generalization. Additionally, due to lack of access to neuromorphic platforms like Intel Loihi or IBM TrueNorth, hardware-level validation of spike efficiency and latency remains unexplored. Lastly, the model structure is still in a foundational stage, and there is room for architectural improvements.

Looking forward, future research should focus on expanding the dataset diversity to include more complex, multilingual, or multi-domain text sources. There is also strong potential in hybrid models that integrate SNN layers with Transformer components, aiming to combine the benefits of biological plausibility and state-of-the-art accuracy. Finally, deploying the model on actual neuromorphic hardware would provide a more accurate assessment of its energy and latency advantages, moving closer to practical, brain-like cognitive systems.

# Bibliography

- [1] A. Maas, R. Daly, P. Pham, D. Huang, A. Ng, and C. Potts, “Learning word vectors for sentiment analysis,” Jan. 2011, pp. 142–150.
- [2] L. N. Andersen and T. S. Haus, “Dopamine modulated stdp and reinforcement learning in an embodied context,” 2013. [Online]. Available: <https://api.semanticscholar.org/CorpusID:58927504>.
- [3] T. Mikolov, “Efficient estimation of word representations in vector space,” *arXiv preprint arXiv:1301.3781*, 2013.
- [4] T. Mikolov, I. Sutskever, K. Chen, G. S. Corrado, and J. Dean, “Distributed representations of words and phrases and their compositionality,” *Advances in neural information processing systems*, vol. 26, 2013.
- [5] J. Pennington, R. Socher, and C. D. Manning, “Glove: Global vectors for word representation,” in *Proceedings of the 2014 conference on empirical methods in natural language processing (EMNLP)*, 2014, pp. 1532–1543.
- [6] S. K. Esser, R. Appuswamy, P. Merolla, J. V. Arthur, and D. S. Modha, “Backpropagation for energy-efficient neuromorphic computing,” in *Advances in Neural Information Processing Systems*, C. Cortes, N. Lawrence, D. Lee, M. Sugiyama, and R. Garnett, Eds., vol. 28, Curran Associates, Inc., 2015. [Online]. Available: [https://proceedings.neurips.cc/paper\\_files/paper/2015/file/10a5ab2db37feedfdeaab192ead4ac0e-Paper.pdf](https://proceedings.neurips.cc/paper_files/paper/2015/file/10a5ab2db37feedfdeaab192ead4ac0e-Paper.pdf).
- [7] T. Iakymchuk, A. Rosado-Muñoz, J. F. Guerrero-Martínez, M. Bataller-Mompeán, and J. V. Francés-Víllora, “Simplified spiking neural network architecture and stdp learning algorithm applied to image classification,” *EURASIP Journal on Image and Video Processing*, vol. 2015, pp. 1–11, 2015.
- [8] O. Y. Sinyavskiy, “Training of spiking neural networks based on information theoretic costs,” 2016. arXiv: 1602.04742 [cs.NE]. [Online]. Available: <https://arxiv.org/abs/1602.04742>.
- [9] P. Bojanowski, E. Grave, A. Joulin, and T. Mikolov, “Enriching word vectors with subword information,” *Transactions of the association for computational linguistics*, vol. 5, pp. 135–146, 2017.
- [10] A. Conneau, D. Kiela, H. Schwenk, L. Barrault, and A. Bordes, “Supervised learning of universal sentence representations from natural language inference data,” *arXiv preprint arXiv:1705.02364*, 2017.
- [11] A. Vaswani, “Attention is all you need,” *Advances in Neural Information Processing Systems*, 2017.

- [12] Z. Bing, C. Meschede, K. Huang, *et al.*, “End to end learning of spiking neural network based on r-stdp for a lane keeping vehicle,” Jun. 2018. DOI: 10.1109/ICRA.2018.8460482.
- [13] Z. Bing, C. Meschede, K. Huang, *et al.*, “End to end learning of spiking neural network based on r-stdp for a lane keeping vehicle,” in *2018 IEEE international conference on robotics and automation (ICRA)*, IEEE, 2018, pp. 4725–4732.
- [14] D. Cer, “Universal sentence encoder,” *arXiv preprint arXiv:1803.11175*, 2018.
- [15] S. R. Kheradpisheh, M. Ganjtabesh, S. J. Thorpe, and T. Masquelier, “Stdp-based spiking deep convolutional neural networks for object recognition,” *Neural Networks*, vol. 99, pp. 56–67, 2018.
- [16] M. Mozafari, S. R. Kheradpisheh, T. Masquelier, A. Nowzari-Dalini, and M. Ganjtabesh, “First-spike-based visual categorization using reward-modulated stdp,” *IEEE Transactions on Neural Networks and Learning Systems*, vol. 29, no. 12, pp. 6178–6190, Dec. 2018, ISSN: 2162-2388. DOI: 10.1109/tnnls.2018.2826721. [Online]. Available: <http://dx.doi.org/10.1109/TNNLS.2018.2826721>.
- [17] M. Mozafari, S. R. Kheradpisheh, T. Masquelier, A. Nowzari-Dalini, and M. Ganjtabesh, “First-spike-based visual categorization using reward-modulated stdp,” *IEEE transactions on neural networks and learning systems*, vol. 29, no. 12, pp. 6178–6190, 2018.
- [18] M. Pfeiffer and T. Pfeil, “Deep learning with spiking neurons: Opportunities and challenges,” *Frontiers in Neuroscience*, vol. 12, 2018, ISSN: 1662-453X. DOI: 10.3389/fnins.2018.00774. [Online]. Available: <https://www.frontiersin.org/journals/neuroscience/articles/10.3389/fnins.2018.00774>.
- [19] E. Saravia, H.-C. Liu, Y.-H. Huang, J. Wu, and Y.-S. Chen, “Carer: Contextualized affect representations for emotion recognition,” Jan. 2018, pp. 3687–3697. DOI: 10.18653/v1/D18-1404.
- [20] M. Mozafari, M. Ganjtabesh, A. Nowzari-Dalini, S. J. Thorpe, and T. Masquelier, “Bio-inspired digit recognition using reward-modulated spike-timing-dependent plasticity in deep convolutional networks,” *Pattern recognition*, vol. 94, pp. 87–95, 2019.
- [21] E. O. Neftci, H. Mostafa, and F. Zenke, “Surrogate gradient learning in spiking neural networks: Bringing the power of gradient-based optimization to spiking neural networks,” *IEEE Signal Processing Magazine*, vol. 36, no. 6, pp. 51–63, 2019.
- [22] E. O. Neftci, H. Mostafa, and F. Zenke, “Surrogate gradient learning in spiking neural networks: Bringing the power of gradient-based optimization to spiking neural networks,” *IEEE Signal Processing Magazine*, vol. 36, no. 6, pp. 51–63, 2019. DOI: 10.1109/MSP.2019.2931595.
- [23] E. Strubell, A. Ganesh, and A. McCallum, *Energy and policy considerations for deep learning in nlp*, 2019. arXiv: 1906.02243 [cs.CL]. [Online]. Available: <https://arxiv.org/abs/1906.02243>.

- [24] J. C. V. Tieck, P. Becker, J. Kaiser, *et al.*, “Learning target reaching motions with a robotic arm using brain-inspired dopamine modulated stdp,” in *2019 IEEE 18th International Conference on Cognitive Informatics & Cognitive Computing (ICCI\* CC)*, IEEE, 2019, pp. 54–61.
- [25] J. Tissier, C. Gravier, and A. Habrard, “Near-lossless binarization of word embeddings,” in *Proceedings of the AAAI Conference on Artificial Intelligence*, vol. 33, 2019, pp. 7104–7111.
- [26] Y. Wu, X. Hu, L. Liang, *et al.*, “Rethinking the performance comparison between snns and anns,” *Neural Networks*, vol. 121, Sep. 2019. DOI: 10.1016/j.neunet.2019.09.005.
- [27] M. Białas, M. M. Mironczuk, and J. Mańdziuk, “Biologically plausible learning of text representation with spiking neural networks,” in *International Conference on Parallel Problem Solving from Nature*, Springer, 2020, pp. 433–447.
- [28] Y. Hao, X. Huang, M. Dong, and B. Xu, “A biologically plausible supervised learning method for spiking neural networks using the symmetric stdp rule,” *Neural Networks*, vol. 121, pp. 387–395, 2020.
- [29] Q. Zhou, C. Ren, and S. Qi, “An imbalanced r-stdp learning rule in spiking neural networks for medical image classification,” *IEEE Access*, vol. 8, pp. 224 162–224 177, 2020.
- [30] M. Yao, H. Gao, G. Zhao, *et al.*, “Temporal-wise attention spiking neural networks for event streams classification,” 2021. arXiv: 2107.11711 [cs.CV]. [Online]. Available: <https://arxiv.org/abs/2107.11711>.
- [31] Y. Chen, Z. Yu, W. Fang, Z. Ma, T. Huang, and Y. Tian, “State transition of dendritic spines improves learning of sparse spiking neural networks,” in *International Conference on Machine Learning*, PMLR, 2022, pp. 3701–3715.
- [32] J. Ding, B. Dong, F. Heide, *et al.*, “Biologically inspired dynamic thresholds for spiking neural networks,” *Advances in Neural Information Processing Systems*, vol. 35, pp. 6090–6103, 2022.
- [33] G. Hinton, *The forward-forward algorithm: Some preliminary investigations*, 2022. arXiv: 2212.13345 [cs.LG]. [Online]. Available: <https://arxiv.org/abs/2212.13345>.
- [34] J. D. Nunes, M. Carvalho, D. Carneiro, and J. S. Cardoso, “Spiking neural networks: A survey,” *IEEE Access*, vol. 10, pp. 60 738–60 764, 2022.
- [35] Y. Dong, D. Zhao, Y. Li, and Y. Zeng, “An unsupervised stdp-based spiking neural network inspired by biologically plausible learning rules and connections,” *Neural Networks*, vol. 165, pp. 799–808, 2023.
- [36] S. Gandhi, R. Gala, J. Kornberg, and A. Sridhar, *Extending the forward forward algorithm*, 2023. arXiv: 2307.04205 [cs.LG]. [Online]. Available: <https://arxiv.org/abs/2307.04205>.
- [37] A. Knipper, *Snnlp: Spiking neural networks for natural language processing*, <https://github.com/alexknipper/SNNLP>, Accessed: 2025-06-18, 2023.

- [38] E. Lemaire, L. Cordone, A. Castagnetti, P.-E. Novac, J. Courtois, and B. Miramond, “An analytical estimation of spiking neural networks energy efficiency,” in *Neural Information Processing*. Springer International Publishing, 2023, pp. 574–587, ISBN: 9783031301056. DOI: 10.1007/978-3-031-30105-6\_48. [Online]. Available: [http://dx.doi.org/10.1007/978-3-031-30105-6\\_48](http://dx.doi.org/10.1007/978-3-031-30105-6_48).
- [39] L.-Y. Niu, Y. Wei, W.-B. Liu, J.-Y. Long, and T.-h. Xue, “Research progress of spiking neural network in image classification: A review,” *Applied Intelligence*, vol. 53, pp. 1–25, Mar. 2023. DOI: 10.1007/s10489-023-04553-0.
- [40] Z. Shao, X. Fang, Y. Li, C. Feng, J. Shen, and Q. Xu, “Eicil: Joint excitatory inhibitory cycle iteration learning for deep spiking neural networks,” *Advances in Neural Information Processing Systems*, vol. 36, pp. 32 117–32 128, 2023.
- [41] J. Vitay, *Ic\_irp\_2023\_spikingforwardforward*, [https://github.com/JulienInWired/IC\\_IRP\\_2023\\_SpikingForwardForward](https://github.com/JulienInWired/IC_IRP_2023_SpikingForwardForward), Accessed: 2025-06-18, 2023.
- [42] Y. Zeng, D. Zhao, F. Zhao, *et al.*, “BrainCog: A spiking neural network based, brain-inspired cognitive intelligence engine for brain-inspired AI and brain simulation,” *en, Patterns (N Y)*, vol. 4, p. 100 789, Jul. 2023.
- [43] R. A. Knipper, K. Mishty, M. Sadi, and S. K. K. Santu, “Snnlp: Energy-efficient natural language processing using spiking neural networks,” *arXiv preprint arXiv:2401.17911*, 2024.
- [44] R.-J. Zhu, Q. Zhao, G. Li, and J. K. Eshraghian, *Spikegpt: Generative pre-trained language model with spiking neural networks*, 2024. arXiv: 2302.13939 [cs.CL]. [Online]. Available: <https://arxiv.org/abs/2302.13939>.
- [45] M. Ghader, S. R. Kheradpisheh, B. Farahani, and M. Fazlali, *Backpropagation-free spiking neural networks with the forward-forward algorithm*, 2025. arXiv: 2502.20411 [cs.NE]. [Online]. Available: <https://arxiv.org/abs/2502.20411>.
- [46] M. Ghader, S. R. Kheradpisheh, B. Farahani, and M. Fazlali, *Backpropagation-free spiking neural networks with the forward-forward algorithm*, 2025. arXiv: 2502.20411 [cs.NE]. [Online]. Available: <https://arxiv.org/abs/2502.20411>.