

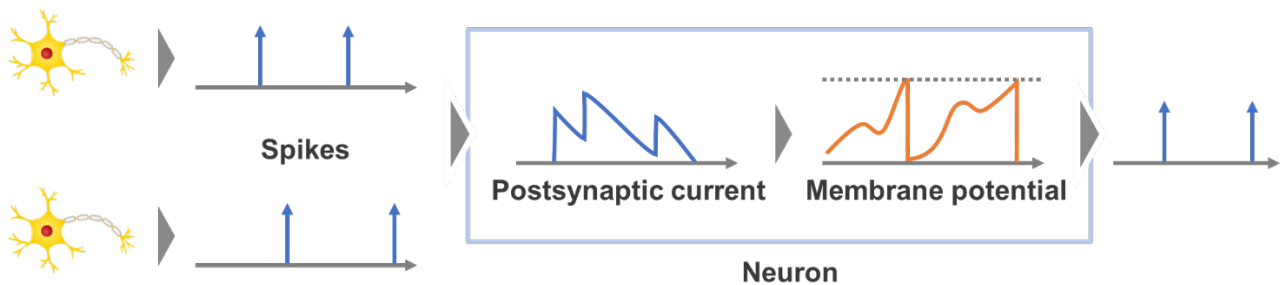
2024 年 11 月 27 日

報道機関 各位

千葉工業大学 数理工学研究センター

## スパイキングニューラルネットワークによる テンポラルコーディングの実現 ～単一スパイク制限を越えた誤差逆伝搬則～

キーワード：スパイキングニューラルネットワーク、テンポラルコーディング



### [ 発表者 ]

- ・山本 かけい（マサチューセッツ工科大学）
- ・酒見悠介（千葉工業大学 数理工学研究センター 上席研究員／東京大学国際高等研究所ニューロインテリジェンス国際研究機構（WPI-IRCN） 連携研究者）
- ・合原 一幸（東京大学 特別教授・名誉教授／東京大学国際高等研究所ニューロインテリジェンス国際研究機構（WPI-IRCN） エグゼクティブディレクター／千葉工業大学 数理工学研究センター 主席研究員）

### [ 概要 ]

山本かけい（MIT）、酒見悠介（千葉工大）、合原一幸（東大）は、脳のようにスパイク信号で情報処理を行うスパイキングニューラルネットワーク(SNN)に対し、新たな学習手法を提案しました。本手法は、スパイクの発生時刻を直接学習に組み込むことで、発火タイミングに基づくテンポラルコーディングをより理想的な形で実現します。従来の類似手法では、各ニューロンの発火回数が1回以下に制限されていましたが、この新しい手法は複数回の発火を可能にする

ことで、柔軟な情報処理が可能になります。これにより、脳の情報処理メカニズムの研究や、脳型ハードウェアの高効率化が期待されます。この成果は、2024 年 7 月 2 日に開催された査読付き国際学会「2024 International Joint Conference on Neural Networks (IJCNN)」で発表されました。

## ■ 背景

スパイクニューラルネットワーク(SNN)は、生物脳の神経回路にインスパイアされた人工知能モデルで、スパイクと呼ばれる不連続な信号を用いて情報を処理します。SNN は既存のディープラーニングモデルと比べ、エネルギー効率に優れており、特にエッジ AI としての応用が期待されています。SNN でより効率的に情報処理を行うには、スパイク発火時刻に基づく情報処理、すなわちテンポラルコーディングを実現する学習方法が必要です。

テンポラルコーディングを実現する主な手法は Time-to-First-Spike (TTFS) コーディングに基づいており、各ニューロンが各入力に対して一度しか発火できない制約がありました。この制約により、スパイク数の削減が図られる一方で、各ニューロンの情報処理能力が限定され、学習性能が低下する可能性が指摘されていました。

## ■ 内容

本研究では、従来の TTFS コーディングの限界を克服し、ニューロンが複数回発火できる multi-spike SNN を実現する、発火時刻に基づいた誤差逆伝搬アルゴリズムを開発しました。この新しいアルゴリズムは、ニューロンの発火時におけるシナプス電流を保存し、ニューロンが複数回発火した場合の、各ニューロンの膜電位および発火時刻の解析的な解を導出することで得られました。

またテンポラルコーディングの誤差逆伝播を実装する際に、初期値に依存して一度も発火しないニューロン (dead neurons) が学習に参加できないという問題が指摘されていました。我々はこの問題に対処するために新たなペナルティ正則化を提案しました。これにより初期値に依存せずに全てのニューロンが学習に組み込まれ、SNN の重みの初期値に対してロバストな学習を実現しました。

加えて従来の single-spike SNN と比較して大きくなった計算量に対処するために、訓練時の損失関数の計算に二段階計算アルゴリズムを導入しました。これは一つの入力サンプルに対しては複数回発火するニューロンの数が限定される事実を利用し、より効率的な計算を実現しました。

この導出した学習則によって得られる SNN の性質を調べるために、隠れ層を一つ持つ SNN (784-400-10)を MNIST データセットで学習させました。図 1 に学習後の SNN のダイナミクスを示しています。従来の single-spike SNN の場合と異なり、複数回発火している様子がわかり

ます。さらに、提案した multi-spike SNN では、single-spike SNN に比べて、より高い認識精度を達成しました（図 2 参照）。表 1 に single-spike SNN を基にした先行研究との比較も行っています。また、ニューロンのリーク時定数（膜電位の減衰速度）が発火回数と学習性能に大きな影響を与えることが確認されました。具体的には、single-spike SNN の場合には、リーク時定数が実験時間（図 2 においては 1）より大きいときは学習性能への影響が全くないのに対して、multi-spike SNN では、ある特定のリーク定数で学習性能が最大化することがわかりました。これらの結果から、multi-spike SNN は複数スパイクを生成するニューロンを許容することでより高い学習性能を有しており、さらに、リークのダイナミクスを有効に活用することができていると考えられます。

今回提案されたマルチスパイク SNN の学習手法は、従来の SNN に比べ、より効率的なテンポラルコーディングを実現するだけでなく、脳型ハードウェアやエネルギー効率の高い AI システムへの応用が期待されます。特に、今回の研究成果は、SNN を利用した次世代の超省電力デバイスやエッジコンピューティングにおける AI 技術の発展に貢献する可能性があります。また、今回の研究で示されたリーク時定数と発火回数の関係は、ニューロモルフィックコンピューティング分野においても重要な知見を提供し、今後のハードウェア設計やアルゴリズム開発に活かされと考えられます。今後は、より大規模で複雑なタスクに対する SNN の学習性能を調べ、さらに高度な脳型計算モデルの実現を目指します。

## ■ 論文情報

国際学会: 2024 International Joint Conference on Neural Networks (IJCNN)

論文題目: Can Timing-Based Backpropagation Overcome Single-Spike Restrictions in Spiking Neural Networks?

著者: Kakei Yamamoto, Yusuke Sakemi, and Kazuyuki Aihara

URL: <https://ieeexplore.ieee.org/document/10651135>

DOI: 10.1109/IJCNN60899.2024.10651135

発表日時: 2024 年 7 月 2 日

## ■ 謝辞

本研究の一部は、JST さきがけ JPMJPR22C5、セコム科学技術振興財団、JST Moonshot R&D Grant Number JPMJMS2021、AMED under Grant Number JP23dm0307009、Institute of AI and Beyond of UTokyo、the International Research Center for Neurointelligence (WPI-IRCN) at The University of Tokyo Institutes for Advanced Study (UTIAS)、JSPS KAKENHI Grant Number JP20H05921、Takenaka Scholarship Foundation から助成を受けて行われました。

## <お問い合わせ先>

### 【本研究内容に関する問い合わせ先】

千葉工業大学 数理工学研究センター 上席研究員

酒見 悠介

HP: <https://sites.google.com/view/rcme-cit/>

TEL: 047-478-0345

E-mail: yusuke.sakemi@p.chibakoudai.jp

### 【取材・大学広報関連に関する問い合わせ先】

千葉工業大学 入試広報部

大橋 慶子

TEL: 047-478-0222 FAX: 047-478-3344

E-mail: ohhashi.keiko@it-chiba.ac.jp

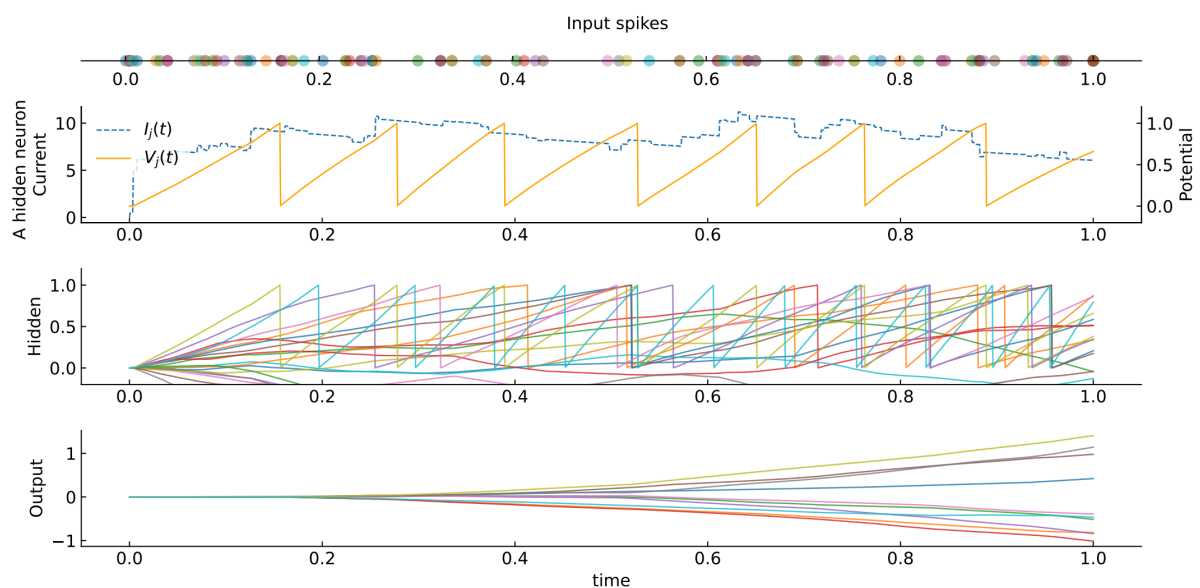


図1 学習後のSNNの動作

隠れ層を一層持つ学習後のSNNの時間発展を示している。最上段の図は発火時刻に符号化された入力スパイクの時刻をプロットしている。上から2番目の図は、ある隠れ層ニューロンの膜電位（橙色実線）とシナプス電流（青色破線）の時間発展を示している。ニューロンは発火後もシナプス電流を保存し、複数回発火している。上から3番目の図は、複数の隠れ層ニューロンの膜電位の時間発展を示している。最下段の図は、出力層ニューロンの膜電位の時間発展を示しており、処理終了時刻（この場合  $t=1.0$ ）の各出力層ニューロンの膜電位を出力値としている。このモデルでは、出力層が発火しないことに留意されたい。

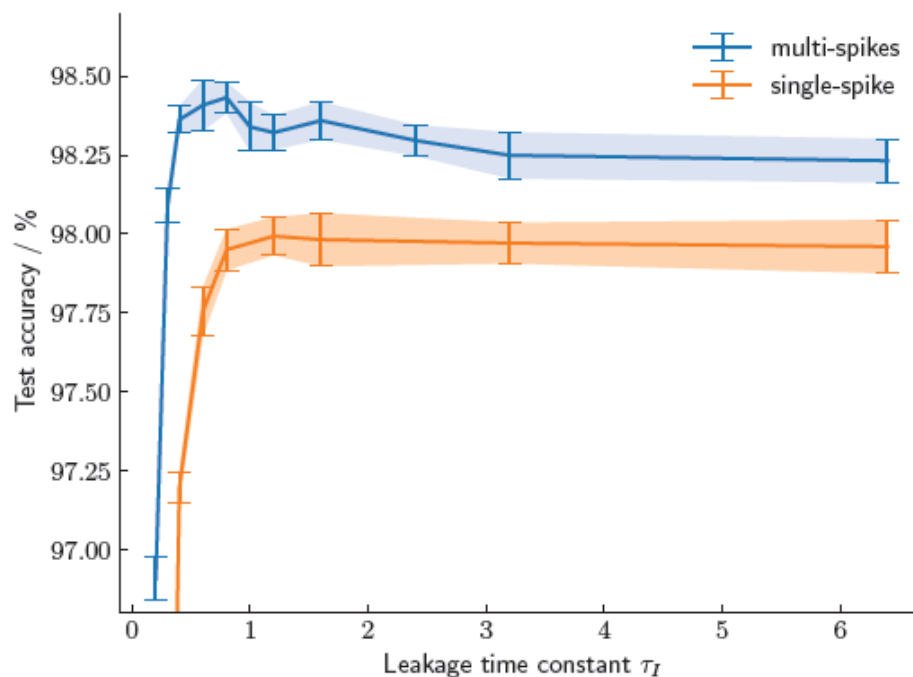


図2 提案モデル (multi-spikes)と従来モデル (single-spike)との性能比較

提案した学習手法によって実現される複数回発火可能なモデル (multi-spikes)と、単一発火に制限された従来モデル (single-spike)の性能を比較している。特に、リーク定数の変化による各モデルの性能を比較している。

表1 類似研究との性能比較

2層のSNNを用いた場合のMNISTデータセットにおける認識性能。Spike Typeについては、Singleは各ニューロンあたり最大一回の発火制約 (TTFS コーディング)を表し、Multiは複数回発火が可能な場合を表す。

| Author                   | Spike Types | Leakage   | Accuracy             |
|--------------------------|-------------|-----------|----------------------|
| Mostafa (2018)           | Single      | Non-leaky | 97.55%               |
| Comsa et al. (2020)      | Single      | leaky     | 97.96%               |
| Sakemi et al. (2021)     | Single      | Non-leaky | 97.99 ± 0.07%        |
| Wunderlich et al. (2021) | Multi       | leaky     | 97.6 ± 0.1%          |
| <b>This work</b>         | Multi       | Non-leaky | 98.23 ± 0.07%        |
| <b>This work</b>         | Multi       | leaky     | <b>98.43 ± 0.05%</b> |