

スパイクニューロンを用いたホップフィールドネットワークと

その連想記憶能力の向上

日大生産工 ○池田 英彬 日大生産工 山内 ゆかり

1 はじめに

スパイクニューラルネットワークとは生物学的な脳の仕組みにより近づけるため、活動電位（スパイク）を重視したニューラルネットワークである。人工ニューラルネットワークは3つの世代に分かれる。第1世代は任意のブール関数を計算できるもの、第2世代がシグモイド関数などを使用し、勾配法による学習を可能にしたもの、第3世代がスパイクニューラルネットワークである¹⁾。

スパイクニューラルネットワークの研究には関数近似を行ったもの²⁾、バックプロパゲーションを可能にしたもの³⁾⁴⁾、進化計算を用いて学習するもの¹⁾などがある。スパイクニューラルネットワークのためのバックプロパゲーションは初めネットワーク内の1つのニューロンがシミュレーション時間内に1回のみ発火するモデルで作成された³⁾。その後、複数の発火が可能なモデルに拡張された⁴⁾。スパイクニューロンのフィードバックネットワークへの応用は少ない⁵⁾。本論文ではスパイクホップフィールドネットワークについて述べる。

2 スパイクホップフィールドネットワーク

ここで扱うモデルはSASAKIら⁵⁾⁶⁾により提案されたモデルである。スパイクの発火タイミングの相対関係でアナログ情報を表現する。本研究で使用する積分発火（IF）ニューロンモデルを示す。ネットワーク構造は通

常のホップフィールドネットワークと同様であり図1、2に示す。

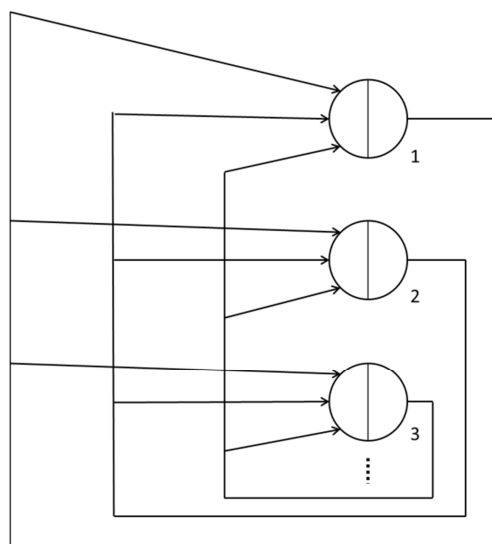


図1. ネットワーク

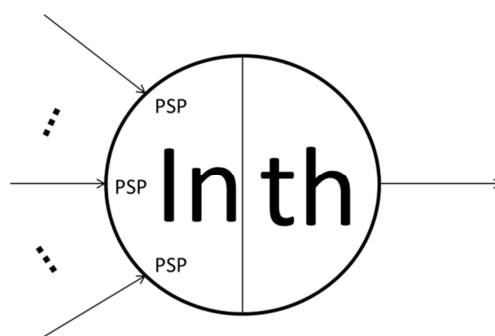


図2. スパイクニューロン

Inは内部電位、thは閾値である。

ニューロンがスパイクを受け取るとシナプス後電位（PSP）が式1により生成される。

$$PSP_{ni}(t) = \sum_{t_i^{(f)} \in \mathcal{F}_i} w_{ni} P(t - t_i^{(f)}) \quad (1)$$

ここで、 $\mathcal{F}_i = \{t_i^{(1)}, \dots, t_i^{(n)}\}$ はニューロンiが

発火した時間の集合、すなわち、スパイク列であり、 w_{ni} はニューロンnとニューロンiの間の重みである。

PSPを決めるカーネル関数を以下に示す。

$$P(s) = P_0 \left\{ \left(1 - e^{-\frac{s}{\tau}} \right) H(s) H(t_p - s) + (1 - e^{-\frac{t_p}{\tau}}) e^{-\frac{s-t_p}{\tau}} H(s - t_p) \right\} \quad (2)$$

ここで、 P_0 は定数、 τ は立ち上がり方を決める定数、 t_p はPSPが立ち上がり始めてからピークに達するまでの時間、 H はヘビサイドのステップ関数である。

ニューロンには内部電位 I_n が存在しPSPの総和として表される。

$$I_n(t) = \sum_{i \in \Gamma_n} PSP_{ni}(t) \quad (3)$$

ここで、 Γ_n はニューロンnへ入力するニューロンの集合である。

式3による内部電位 I_n が閾値 th を超えるとそのニューロンが発火しシナプスを通じ全てのニューロンにスパイクが伝わる。

ニューロンが発火すると、内部電位 I_n はリセットされ、期間 Tr の間発火しない。これを不応期と言う。同時に、生成されたスパイクは伝達遅れ Td の後、自身を含むすべてのニューロンに伝わる。不応期 Tr と伝達遅れ Td は同じ値であり200nsである。同時に、スパイクには幅が存在し、その幅の間、PSPを作り続けるとする。この幅は20nsとする。

全てのニューロンは互いに接続し合っており、全ての接続には重みが存在する。その重みは式4により記憶されるパターンから決定される。結合には興奮性の結合と抑制性の結合があり、重みの符号の正負でそれぞれ決定される。

$$w_{ij} = \sum_{k=1}^N (2I_i^k - 1)(2I_j^k - 1), \text{ for } i \neq j \quad (4)$$

ただし、 $w_{ii} = 0$ であり、 I_i^k はk番目の保存されたパターンのi番目の要素である。パター

ンの内、黒を0、白を1とする。

使用した5つのパターンを図3に示す。ニューロンは左上から右に向かって順に1番、2番、3番、…、36番とする。

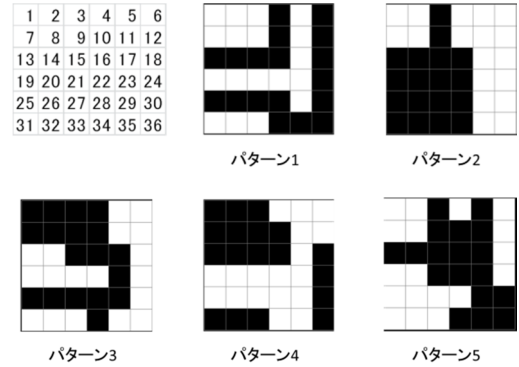


図3. 記憶するパターン

積分発火ニューロンモデルでは内部電位が閾値を超えない場合、ニューロンがスパイクを生成しなくなるという問題がある。これを解決するためにnegative thresholdingを使用する⁵⁾⁶⁾。それを行うための方法がGlobal Excitatory Unit (GEU) である。GEUはスパイクを受け取ると直ちに発火し、全てのニューロンに一定の興奮を与え、その結果、閾値を下げることに同じ効果を与える。GEUから与えられる内部電位は以下の式5で与えられる。

$$G_n(t) = w_{nG} \left\{ 1 - e^{-\left(\frac{t-t_i^{(1)}}{\tau_G} \right)} \right\} H(t - t_i^{(1)}) \quad (5)$$

w_{nG} は全てのニューロンに共通するGEUからの重みである。 τ_G は立ち上がりを決める定数であり、 H はヘビサイドのステップ関数である。

スパイクを用いて白と黒のパターンを表現する方法について以下の方法を採用。最初にパターンを入力する場合、相当するニューロンを、白なら0ns、黒なら100nsで発火させることにより入力する。

3 実験内容

SASAKIら⁵⁾による研究においてノイズの入

ったパターンを入力し正しいパターンを想起するかについて実験がなされ、2節で提示したモデルは正しく想起することが分かっている。しかし、このモデルに何度も信号を与えた場合の振る舞いについては実験が行われていない。そこで、筆者はこのモデルに対して断続的に新しいスパイクを入力し、その後のスパイクの振る舞いを観察した。

4 実験結果

図4、5の縦軸は上から順に1番から12番までのニューロンのスパイク列 $\mathcal{F}_1, \mathcal{F}_2, \dots, \mathcal{F}_{12}$ 、つまり、図3の入力パターンの例における上部2段分のニューロンのスパイク列である。横軸は時間である。パターン1を入力し、それを想起させた状態で始め、1000nsから500nsごとに5回、パターン2を入力した。その結果、パターン1から、パターン2へ想起結果が変化していることが図4、5より分かる。

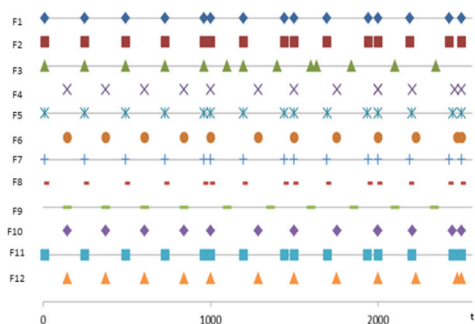


図4. 0ns以降、2500ns

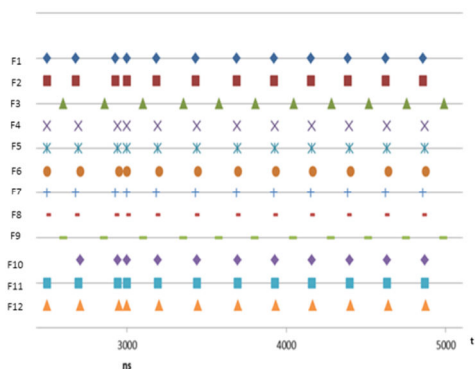


図5. 2500ns以降、5000nsまで

加えて、パターン2を1回入力した場合から5回入力した場合までについて、5000nsまでシミュレーションを行い、その最新のスパイク発火時間を計測した。最新のスパイクの発火時間の標準偏差について、パターン2において白を示すべきニューロンと黒を示すべきニューロンそれぞれについて入力回数に応じてどのように変化するか図6に示す。図6より入力回数を多くするほど発火時間が一致していく様子が分かる。

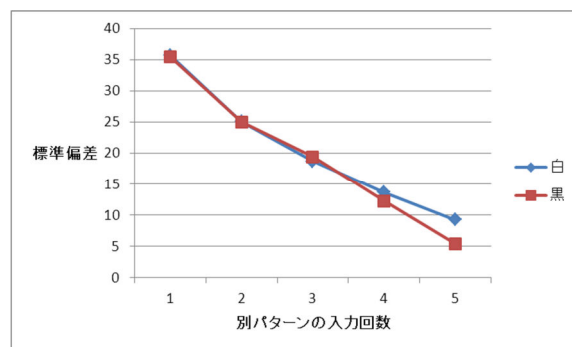


図6. パターンの入力回数とスパイク発火時間の標準偏差

パターン1を想起しているネットワークにおいて、パターン2を500ns間隔で新たに入力した場合、どのパターンに収束したかを表1に示す。加えて、パターン1を想起しているネットワークにおいて、パターン2を1回入力した後にパターン1を再び入力した場合、どのパターンに収束したかを表2に示す。

表1

パターン入力	回数	結果
P1の後、P2	1	P1とP2の間
P1の後、P2	2	P2に収束
P1の後、P2	3	P2に収束
P1の後、P2	4	P2に収束
P1の後、P2	5	P2に収束

パターン1の後、P2を複数回入力した

表2

パターン入力	回数	結果
P1、P2の後P1	1	P1に収束
P1、P2の後P1	2	P1に収束
P1、P2の後P1	3	P1に収束

パターン1、パターン2の後、パターン1を複数回再入力した

ネットワークに1回だけ別パターンを入力した場合はそのパターンを想起することができず、最初に想起していたパターンとの中間を想起することが分かった。

ネットワークがあるパターンを想起している状況から他のパターンを想起させるためには何回か他のパターンを断続的に入力しなければならないことが分かった。

最初に想起していたパターンと追加したパターンとの中間を想起している場合に、最初に想起していたパターンを再入力することにより、ネットワークが最初に想起していたパターンに戻る事が分かった。

パターンの入力回数を多くすることで標準偏差が減少する。すなわち、より入力したパターンに収束することが分かった。

5 おわりに

本論文ではスパイクングホップフィールドネットワークについて実験を行い、追加パターンの提示によるネットワークの振る舞いを解析した。

スパイクングホップフィールドネットワークに追加パターンを提示することはSASAKIらによる既存のモデルを改変していると考えられる。従って、今後の課題として、追加パターンの提示を外部刺激としたモデルを提案し、これらの現象を解析したい。

参考文献

- 1) N.G. Pavlidis¹, D.K. Tasoulis¹, V.P. Plagianakos¹, G. Nikiforidis and M.N. Vrahatis “Spiking Neural Network Training Using Evolutionary Algorithms” , Proceedings of International Joint Conference on Neural Networks, Montreal, Canada, 2005
- 2) Nicolangelo Iannella, Andrew D. Back “A spiking neural network architecture for nonlinear function approximation” , Neural Networks, 14, 2001, p.933-939
- 3) Sander M. Bohte^a, Joost N. Koko^a, Han La Poutr “Error-backpropagation in temporally encoded networks of spiking neurons” , Neurocomputing, 48, 2002, p. 17-37
- 4) Samanwoy Ghosh-Dastidar , Hojjat Adeli “A new supervised learning algorithm for multiple spiking neural networks with application in epilepsy and seizure detection” , Neural Networks, 22, 2009, p. 1419-1431
- 5) Kan'ya SASAKI, Takashi MORIE, Atsushi IWATA “A VLSI Spiking Feedback Neural Network with Negative Thresholding and Its Application to Associative Memory” , IEICE TRANS.ELECTRON., VOL. E89-C, NO. 11, 2006, p. 1637-1644
- 6) Kan'ya SASAKI, Takashi MORIE, Atsushi IWATA “A spiking neural network with negative thresholding and its application to associative memory” , The 47th IEEE International Midwest Symposium on Circuits and Systems, 2004