

Modeling the role of specific inhibitory cells in the processing of visual stimuli

Fatemeh Bazeli Mahbob¹ , Mir Shahram Safari^{2*} , Abdolhossein Vahabi³

1. Cognitive Science Modeling, Shahid Beheshti University, Tehran, Iran

2. Neuroscience Research Center, Shahid Beheshti University of Medical Sciences, Tehran, Iran

3. Faculty Member of the Faculty of Electrical and Computer Engineering and the Faculty of Psychology and Educational Sciences, University of Tehran, Tehran, Iran

Abstract

Received: 10 Dec. 2024

Revised: 2 Mar. 2025

Accepted: 4 Mar. 2025

Keywords

Inhibitory neuron
Visual stimulation
Generalized leaky integrate-and-fire models

Corresponding author

Mir Shahram Safari, Neuroscience Research Center, Shahid Beheshti University of Medical Sciences, Tehran, Iran

Email: Safari@sbmu.ac.ir



 doi.org/10.30514/icss.27.1.73

Introduction: Gamma-aminobutyric acid (GABA) is an essential inhibitory neurotransmitter that regulates neural communication and reduces neuron activity, reducing stress and promoting sleep. Two main types of inhibitory interneurons, Parvalbumin (PV⁺) and Somatostatin (SST⁺), shape brain responses to visual stimuli, although their specific roles remain unclear. Gathering quantitative data on their connectivity with pyramidal cells is vital for deciphering complex brain circuits. Understanding these neurons aids in developing treatments for neurological disorders and advancing intelligent machine design.

Methods: The study utilized data from the Allen Institute, employing the patch clamp technique on mice to analyze the firing characteristics of central inhibitory neurons, PV⁺ and SST⁺. Mice aged 45-70 days were anesthetized, perfused, and brain slices prepared. Researchers enhanced generalized leaky integrate-and-fire (GLIF) models to simulate neuronal spiking behavior, optimizing them with electrophysiological data to accurately reproduce neural responses across various stimuli.

Results: The analysis of explained variance indicates that additional mechanisms are needed to replicate the spiking behavior of inhibitory and excitatory neurons. The GLIF₁ model achieved a variance of 65%, with inhibitory neurons showing higher performance (82%) compared to excitatory neurons (69%). The GLIF₂ model's reset rule reduced performance for both types, while the GLIF₃ model, including post-spike currents, improved inhibitory neuron variance to 84%. Introducing post-spike currents in the GLIF₄ model further enhanced performance for both neuron types. Overall, the findings highlight the distinct effects of post-spike currents on neuron types and the complexity of modeling their behaviors.

Conclusion: The present research concluded that GLIF models can accurately replicate biological neuron spike times with few adjustable parameters, simplifying input-output matching. More complex models, like GLIF₃, outperform simpler ones in differentiating transgenic lines. However, increasing complexity risks overfitting, complicating optimization. Clustering algorithms can classify cell types based on electrophysiological features, but identifying crucial features is essential for effective modeling.

Citation: Bazeli Mahbob F, Safari MSh, Vahabi A. Modeling the role of specific inhibitory cells in the processing of visual stimuli. *Advances in Cognitive Sciences*. 2025;27(1):73-90.

Extended Abstract

Introduction

Gamma-aminobutyric acid (GABA) is a vital neurotransmitter that plays a key role in regulating neural networks and communication between different brain regions. As an inhibitory neurotransmitter, it reduces neuron activity,

blocks stress signals, and promotes restful sleep. By studying GABA, researchers gain insights into how the brain functions and how it can be affected by neurological disorders. However, studying inhibitory neurons in the ce-

rebral cortex has posed challenges for researchers due to tool limitations. Imbalances in excitatory and inhibitory messages in the brain can cause harmful seizures. Hence, it is crucial to understand the connectivity between inhibitory interneurons and target cells. These interneurons are also responsible for regulating sensory stimuli processing and shaping visual stimulus responses, making it essential to study the tail. Two primary subtypes of inhibitory interneurons, Parvalbumin (PV+) and Somatostatin (SST+), play an essential role in shaping the brain's response to visual stimuli movement. However, whether these subtypes have distinct quantitative and functional roles in this process is still unclear. Accurate quantitative data on their connectivity with pyramidal target cells are crucial for understanding these complex brain circuits. Studying GABA and inhibitory interneurons provides crucial insights into the brain's functioning and the mechanisms behind neurological disorders. By understanding the intricate connectivity and roles of these neurons in regulating brain activity, researchers can develop effective treatments for various neurological conditions. Moreover, this research can potentially contribute to developing intelligent machines based on the human brain's functioning principles. Ultimately, a precise understanding and modelling of the nervous system can significantly advance the treatment of neurological conditions such as Alzheimer's, multiple sclerosis, and Parkinson's.

Methods

The information was sourced from the reputable Allen Institute and obtained through electrophysiological means using the patch clamp technique on mice. The Allen Institute database offers fundamental insights into cell firing characteristics, including those of the central inhibitory neurons PV+ and SST+, as observed in whole-cell patch clamp recordings. This study's recordings involved a variety of stimuli, such as short pulses, long steps, slow

ramps, and natural noise, to evaluate the inherent attributes of these neurons. The sections utilized in our research were extracted from adult mouse tissue.

Male or female mice between 45 and 70 days old were used to prepare brain slices. The mice were anaesthetized with 5% isoflurane and then perfused with 25 or 50 ml of oxygenated ice-cold artificial cerebrospinal fluid (ACSF.I) at a 9 ml/min flow rate via intracardiac injection. The perfusion continued until the liver was clear or all the perfusion solution was used. Afterwards, the brain was quickly dissected and placed on a Compressstome VF-300 (Precision Instruments) vibrating microtome chuck to create a coronal section.

In this study, researchers utilized GLIF models to mimic the firing patterns of neurons at varying levels of complexity, accurately replicating the spiking behaviour of both mouse and human neurons. By enhancing the leaky integral and fire (LIF) model with three generalizations-post-spike currents to represent ion channel activation effects, sub-threshold voltage and spike-dependent changes in the threshold, and voltage and threshold resetting rules based on electrophysiology data-the team optimized the model by maximizing the likelihood of a neuron with intrinsic noise that perfectly emulates a spike. The model's efficacy was tested by evaluating its ability to reproduce observed sequences on test stimuli, with the fraction of neural response variance explained by the model calculated across different time scales.

Results

The analysis of explained variance across various model levels suggests that additional mechanisms are necessary to achieve the spiking behaviour of inhibitory and excitatory neurons. Specifically, when considering all neurons, the GLIF₁ model (equivalent to the integral model and generalized leaky fire) demonstrated a variance above 65%. Notably, inhibitory neuron models outper-

formed excitatory models, with 82% and 69% variances, respectively. The reset rule in the GLIF₂ model reduced performance, or both inhibitory (76%) and excitatory (63%) neurons' post-spike currents in the GLIF₃ model contributed significantly to inhibitory neurons (84%) but not to excitatory neurons (65%). To further verify that the improvement in fit between the GLIF₂ and GLIF₃ models followed distinct trends for inhibitory and excitatory neurons, the researchers calculated the two-by-two explained variance ratio and found that differences in distribution were unique to each neuron type. These findings suggest that post-spike currents have distinct effects on inhibitory and excitatory neurons.

By introducing post-spike currents to the reset rule in the GLIF₄ model, excitatory (79%) and inhibitory (88%) neurons improved performance. However, the reset rule alone hurt performance. The difference in distribution between GLIF₃ and GLIF₄ was statistically significant for both types of neurons. Furthermore, adding adaptive thresholding to post-spike currents and baseline resets enhanced performance for excitatory cells (79%) but only slightly improved inhibitory cells (86%). Directly measuring the reset rule from the data hindered the ability of the GLIF₂ model to replicate the rise times of neural data. The present study explored the correlation between the model's ability to produce subthreshold behavior in the voltage waveform and the reproduction of neuronal spike times. While generally a correlation exists between the two, the average performance values indicate that although GLIF₂ outperforms GLIF₂ and GLIF₃ in replicating subthreshold voltage, it fares worse in producing spiky behaviour. Therefore, reproducing subthreshold voltage does not necessarily lead to better time performance.

When considering all the models, a correlation was generally found between the model's ability to reproduce subthreshold voltage and its ability to reproduce spike times. The model's ability to reproduce subthreshold voltage did

not necessarily lead to better spike timing performance. This study found that, out of 345 neurons, 167 from three transgenic lines matched the tuning criteria of the GLIF₄ and GLIF₅ models.

Conclusion

The current research demonstrated that GLIF models can recreate the spike times of biological neurons with only a few adjustable parameters while simplifying the matching complexity associated with the input-output spike sequence. Besides, the obtained findings on the relationship between complexity and ability reveal that these models can differentiate transgenic lines through unsupervised clustering while simultaneously reproducing spike times. Notably, more intricate models (such as GLIF₃ or higher) are superior to simpler models in producing biological spike times and accurately differentiating transgenic lines.

By optimizing all parameters at once, we can lower the optimized function's error by increasing the present model's complexity (by incorporating new mechanisms and variables into the previous model). However, if we continue to add parameters or state variables, we may reach a point where the model starts to overfit the training data. At this stage, optimization based on the present performance measure (variance explained) will no longer be convex, and additional optimization will require more significant computation and a higher level of uncertainty. As a result, there is no guarantee of convergence.

Electrophysiological feature sets can be analyzed through clustering algorithms to classify cell types associated with input/output conversion observed in electrophysiology experiments. However, it is essential to determine which features are crucial for consideration. Alternatively, one can integrate the input/output relationship with the model using the entire spike sequence information set and cluster based on the model parameters. For this purpose, a

model with biophysical detail is appropriate, as the parameters correspond well to biological mechanisms.

Ethical Considerations

Compliance with ethical guidelines

This study utilized publicly available animal data provided by the Allen Institute. The original data were generated in accordance with the National Institutes of Health (NIH) guidelines for the care and use of laboratory animals and approved by the relevant institutional ethics committees. No live animals were used or handled by the authors in the current study.

Authors' contributions

The first author wrote the draft paper. After reviewing

and applying some corrections by other authors, the second author compiled the final version.

Funding

This research received no specific grant from any funding agency in the public, commercial, or not-for-profit sectors.

Acknowledgments

This Article is a part of the first author's PhD dissertation. The authors appreciate all experts and professors for their advice.

Conflict of interest

The authors declared no conflicts of interest.

مدل‌سازی نقش سلول‌های اختصاصی مهاري در پردازش محرک‌های بینایی

فاطمه باذلی محبوب^۱ ID، میرشهرام صفری^۲ ID*، عبدالحسین وهابی^۳

۱. مدل‌سازی شناختی، دانشگاه شهید بهشتی، تهران، ایران
۲. استادیار مرکز تحقیقات علوم اعصاب دانشگاه علوم پزشکی شهید بهشتی، تهران، ایران
۳. استادیار دانشکده مهندسی برق و کامپیوتر دانشگاه تهران، تهران، ایران

چکیده

مقدمه: گابا یک انتقال‌دهنده عصبی مهاري کلیدی در دستگاه عصبی مرکزی است که نقش مهمی در کاهش فعالیت نورونی، تنظیم ارتباطات عصبی، کاهش استرس و بهبود کیفیت خواب ایفا می‌کند. دو نوع اصلی از نورون‌های میانجی مهاري، پاروالومین (PV^+) و سوماتواستاتین (SST^+)، در تنظیم پاسخ مغز به محرک‌های حسی، به ویژه بینایی، نقش دارند. با این که اثرات کلی این نورون‌ها مشخص است، اما عملکرد دقیق و ارتباط آنها با سلول‌های پیرامیدال هنوز به‌درستی روشن نشده است. تحلیل کمی این ارتباطات برای رمزگشایی از مدارهای عصبی پیچیده و نیز پیشرفت در درمان بیماری‌های عصبی و طراحی سامانه‌های هوشمند اهمیت بالایی دارد.

روش کار: در این مطالعه از داده‌های مؤسسه Allen استفاده شد. با بهره‌گیری از تکنیک پیچ‌کلمپ، ویژگی‌های شلیک نورون‌های PV^+ و SST^+ در موش‌های ۴۵ تا ۷۰ روزه بررسی شد. حیوانات تحت بیهوشی قرار گرفتند، پرفیوژ شدند و برش‌های مغزی تهیه گردید. سپس از مدل‌های انتگرال و آتش‌نشتی تعمیم‌یافته (GLIF) برای شبیه‌سازی رفتار نورونی استفاده شد و این مدل‌ها با داده‌های الکتروفیزیولوژیکی بهینه‌سازی شدند.

یافته‌ها: تحلیل واریانس نشان داد که مدل $GLIF_1$ قادر به بازتولید ۶۵ درصد از واریانس بود و عملکرد نورون‌های مهاري (۸۲ درصد) بهتر از تحریکی (۶۹ درصد) بود. مدل $GLIF_2$ با قانون بازنشانی کاهش عملکرد نشان داد. در مدل $GLIF_3$ با افزودن جریان‌های پس‌شلیک، واریانس نورون‌های مهاري به ۸۴ درصد رسید. $GLIF_4$ با بهینه‌سازی بیشتر، عملکرد هر دو نوع نورون را بهبود داد.

نتیجه‌گیری: مدل‌های GLIF می‌توانند رفتار نورون‌های زیستی را با پارامترهای اندک و دقت بالا شبیه‌سازی کنند. مدل‌های پیچیده‌تر کارآمدترند، اما احتمال بیش‌برازش را افزایش می‌دهند. انتخاب ویژگی‌های کلیدی برای مدل‌سازی موفق ضروری است.

دریافت: ۱۴۰۳/۰۹/۲۰

اصلاح نهایی: ۱۴۰۳/۱۲/۱۲

پذیرش: ۱۴۰۳/۱۲/۱۴

واژه‌های کلیدی

نورون مهاري

تحریک بینایی

مدل‌های انتگرال و آتش‌نشتی تعمیم‌یافته

(GLIF)

نویسنده مسئول

میرشهرام صفری، استادیار مرکز تحقیقات عوم اعصاب دانشگاه علوم پزشکی شهید بهشتی، تهران، ایران

ایمیل: Safari@sbmu.ac.ir



doi.org/10.30514/icss.27.1.73

مقدمه

همچنین می‌تواند خواب‌آور باشد (۱). شناخت نقش نورون‌های مهاري در قشر مغز برای پژوهشگران چالش‌برانگیز است، زیرا ابزارهای خاص برای مطالعه این نورون‌ها کمبود دارد. بیشتر پژوهش‌ها به تأیید نقش مهار از طریق مدل‌های محاسباتی وابسته‌اند و تنها تعداد معدودی مطالعه بر ویژگی‌ها و مکانیسم‌های مهاري مدارهای عصبی انجام شده است (۲). پژوهش‌های بی‌سابقه‌ای در خصوصیات نورون‌های مهاري،

گابا (GABA) یا گاما آمینوبوتیریک اسید انتقال‌دهنده عصبی مهاري است که به هماهنگی شبکه‌های عصبی و ارتباطات نواحی مختلف مغز کمک می‌کند. بنابراین مطالعه آن به درک ناهماهنگی‌ها و اختلالات عصبی کمک می‌کند. گابا علاوه بر سیستم عصبی مرکزی، در بافت‌های غیرعصبی نیز وجود دارد و با کاهش فعالیت نورون‌ها، از عملکرد بیش از حد آنها جلوگیری و پیام‌های حاوی استرس را به مغز نمی‌رساند و

مدارها و سیناپس های آنها انجام شده که به درک نقش مهار در عملکرد مغز و اختلالات عصبی کمک می کند (۳). رویکردهای جدید ویژگی های متغیر مدارهای مهاری و نقش اینترنورون ها را در فهم فعالیت مغز و بهبود هدایت مدارها روشن می سازند. در مغز طبیعی، تعادل بین پیام های تحریکی و مهاری حفظ می شود. اختلال در این تعادل می تواند به حملات (Seizures) ناشی از ترکیبات ارگانوفسفره و اختلال در هوشیاری و حرکات غیرارادی منجر شود (۴، ۵). در نئوکورتکس، اینترنورون های مهاری گابا تنظیم شکل منحنی پاسخ و شدت پاسخ سلول های پیرامیدال را به محرک های حسی کنترل می کنند. این اینترنورون ها به دو زیرمجموعه اصلی آن پاروالبومین (PV^+) و Somatostation (SST^+) و Somatostation (SST^+) تقسیم می شوند که در مورفولوژی، خواص فیزیولوژیکی و هدف گیری پس سیناپسی تفاوت دارند. تأثیر مهاری گابا بر پاسخ به جهت حرکت محرک بینایی حائز اهمیت است، اما نقش های کمی و عملکردی اینترنورون های PV^+ و SST^+ هنوز مشخص نیست. درک این مکانیسم ها نیازمند داده های کمی درباره اتصال عملکردی آنها با سلول های پیرامیدال است که اطلاعات جامعی در این زمینه وجود ندارد (۶). در سال ۲۰۲۱، Thompson و همکاران به بررسی مدل های محاسباتی نورون های مهاری و تحریکی در اختلالات عصبی، به ویژه اسکیزوفرنی، پرداختند. این مطالعه نشان داد که تعادل میان فعالیت های تحریکی و مهاری چگونه می تواند پردازش بصری را مختل کند و این یافته ها برای درک بهتر مکانیسم های نورون های مهاری در مدارهای بینایی ضروری است. این رویکرد به ما کمک می کند تا ویژگی های شبیه سازی شده در نورون های مهاری را در سیستم های مختلف مغزی بررسی کنیم (۷). در سال ۲۰۱۷، Zhang و همکاران از طریق مدل سازی، نقش اساسی نورون های مهاری در پردازش زمینه ای بصری را مورد تأکید قرار دادند. این پژوهش نشان می دهد که نورون های مهاری می توانند به طور قابل ملاحظه ای در پردازش اطلاعات بصری و درک زمینه های محیطی، به ویژه در پردازش محرک های پیچیده، تأثیرگذار باشند. این مدل سازی ها به ما این امکان را می دهد که رفتار نورون های مهاری را در پاسخ به محرک های بینایی بهتر درک کنیم (۸). همچنین در سال ۲۰۲۱، Peterson و همکاران به تحلیل مدارهای بازخوردی در قشر بینایی اولیه پرداختند. نتایج این مدل سازی نشان داد که بازخوردهای عصبی تأثیرات معناداری در پردازش اطلاعات بصری دارند و به تنظیم دقیق پاسخ های بصری کمک می کنند. بنابراین، این مقاله اهمیت پیچیدگی های بازخوردی را در فرآیندهای بصری و نقش آن در نورون های مهاری مورد توجه قرار داد (۹). در ادامه، Li و همکاران در

سال ۲۰۲۱، به بررسی مدل سازی های میکرو مداری نورون های مهاری، شامل PV^+ ، SST^+ و VIP، در قشر بینایی پرداختند و نشان دادند که ترکیب نورون های مهاری در پردازش اطلاعات بصری و هماهنگی میان نورون های تحریکی و مهاری نقش کلیدی دارد. این مدل به درک دقیق تر از نقش نورون های مهاری در سیستم بینایی کمک می کند (۱۰). در سال ۲۰۲۳، Liu و همکاران به تحلیل نقش نورون های مهاری در کنترل شکل گیری و پایداری شبکه های عصبی از طریق هم زمانی فعالیت های خود پرداختند و بر اهمیت هماهنگی زمانی نورون های مهاری در پردازش بینایی تأکید کردند. این پژوهش ها می تواند به کمک شبیه سازی های مدل سازی و تحلیل دینامیک شبکه های عصبی به درک بهتری از تعاملات نورون های مهاری و تحریکی در پردازش اطلاعات بینایی کمک کند (۱۱). در راستای پیشینه پژوهشی قبلی ما نیز به بررسی نقش زیرگروه های مختلف اینترنورون ها در پردازش اطلاعات بصری پرداختیم. سه گروه از پژوهشگران شامل Atallah و همکاران، Lee و همکاران، Wilson و همکاران این مسئله را با دستکاری اپتوژنتیکی زیر مجموعه اینترنورون ها بررسی کردند و به نتایج متفاوتی دست یافتند. برخی از پژوهشگران مشاهده کردند که فعالیت نورون های SST^+ به وضوح موجب تشدید تنظیم جهت گیری می شود در حالی که فعالیت نورون های PV^+ اثر کمتری بر این فرآیند دارد (۱۲). در مقابل گروه دیگری از پژوهشگران دریافتند که فعالیت نورون های PV^+ تنها کارایی را افزایش می دهد و موجب تقویت تنظیم جهت گیری می شود و نورون های SST^+ تأثیری بر این عملکرد ندارد (۱۳). برای درک علت این تفاوت ها در نتایج Dan و همکارانش چندین آزمایش متفاوت را با تمرکز بر بیهوشی، سطح و مدت زمان تحریک اپتوژنتیک انجام دادند. آنها دریافتند که این اختلافات می تواند با تفاوت هایی در سطح و مدت فعالیت اینترنورون ها توضیح داده شود. این یافته ها نه تنها عمیق تر به پیچیدگی تعاملات میان زیرگروه های اینترنورون ها در پردازش اطلاعات بصری می پردازند بلکه فرصتی برای بررسی بیشتر رفتارهای نورون های مختلف در شرایط آزمایشگاهی فراهم می آورند.

Atallah و همکاران نشان دادند که فعالیت کم PV^+ هیچ تأثیری در کاهش عرض تنظیم جهت گیری (σ) ندارند ولی Lee و همکاران نشان دادند که کاهش σ به شدت با کاهش سرعت شلیک در ارتباط است. این نتایج متفاوت مربوط به درجه فعالیت PV^+ می باشد. برای تست وابستگی کاهش σ در سطح فعالیت اینترنورون، Dan و همکاران تنظیم منحنی هر نورون را در موش بیهوش در ناحیه V_1 در چندین شدت نور اندازه گیری کردند. با این حال اختلاف بین یافته های Lee و Wilson عمیق تر است و آزمایشات اصلی متفاوتی وجود دارد: انتخاب بیهوشی ها

پژوهش است که امکان مدل‌سازی سیستم‌های عصبی بزرگ مقیاس را فراهم می‌کند. این ابزارها می‌توانند در شبیه‌سازی بخش‌هایی از مغز و تحلیل عملکرد آنها مورد استفاده قرار گیرند. در نهایت، این اهداف به پیشبرد شبیه‌سازی‌های زیستی و محاسباتی کمک می‌کنند و پلی میان داده‌های تجربی و مدل‌سازی‌های نظری ایجاد می‌نمایند. این مدل‌ها نه تنها به درک بهتر اختلالات عصبی مانند صرع، پارکینسون و آلزایمر کمک می‌کنند بلکه می‌توانند در طراحی روش‌های درمانی جدید و هدفمند نیز مورد استفاده قرار گیرند. به این ترتیب مدل‌سازی شبکه‌های عصبی به عنوان یک ابزار کلیدی در علوم اعصاب محاسباتی نقش مهمی در پیشبرد دانش ما از سیستم‌های عصبی و بیماری‌های مرتبط با آنها ایفا می‌کند. هدف ما از ارائه مدل پیش‌بینی بهتر رفتار نورون‌ها در زمینه مدل‌سازی نقش نورون‌های مهارتی در تحریک بینایی است که از مدل GLIF استفاده شده است. برای تکمیل پژوهش‌های دکتر صفری و پاسخ به نتایج متناقض مقالات قبلی نیاز به مدل‌سازی و تحلیل دقیق‌تری داریم تا بتوانیم اطلاعات جامع‌تری در این حوزه به دست آوریم.

روش کار

جمع‌آوری داده‌های الکتروفیزیولوژیکی

داده‌های مورد استفاده در این پژوهش از منبع معتبر موسسه Allen استخراج شده که به صورت الکتروفیزیولوژیکی از موش با تکنیک پیچ کلمپ جمع‌آوری شده. ثبت‌های پیچ کلمپ تمام سلول (Whole cell patch clamp recordings) اطلاعات اولیه‌ای را در مورد خواص شلیک سلول ارائه می‌دهد. پایگاه داده موسسه Allen، بینش‌های بنیادی در مورد ویژگی‌های شلیک سلولی، از جمله نورون‌های مهارتی PV^+ و SST^+ ارائه می‌دهد. ثبت‌های ما شامل طیف وسیعی از محرک‌ها مانند پالس‌های مربعی کوتاه، مربعی بزرگ، رمپ‌های آهسته و نویز طبیعی برای ارزیابی ویژگی‌های ذاتی این نورون‌ها بود. بخش‌های مورد استفاده در پژوهش‌های ما از بافت مغز موش بالغ استخراج شدند. برای تهیه برش‌های مغزی از موش‌های نر یا ماده بین ۴۵ تا ۷۰ روز استفاده شد. موش‌ها با ۵ درصد ایزوفلوران بیهوش شدند و سپس با ۲۵ یا ۵۰ میلی‌لیتر مایع مغزی نخاعی مصنوعی یخ‌زده و اکسیژن‌دار (ACSF.I) با سرعت جریان ۹ میلی‌لیتر در دقیقه از طریق تزریق داخل قلبی پرفیوژن شدند. که این داده‌ها را می‌توانید در پایگاه داده موسسه Allen مشاهده کنید (۱۰، ۱۱). پتانسیل غشایی (membrane potential)، زمان‌بندی اسپایک‌ها (spike timing) و پاسخ به ورودی‌های مختلف اندازه‌گیری شد. خطوط ترانس‌ژنیک استفاده شده برای انتخاب انواع خاصی از

و مدت زمان تحریک لیزر و الگویی که برای تحریک استفاده شده (۱۷-۱۴). در تمامی این مطالعات فعالیت‌ها به صورت شبکه‌ای و اپتوژنتیک بررسی شدند که تحریک نورون‌های مهارتی در هر نوع به صورت دسته جمعی و بر روی نورون‌های تحریکی بود. در سال ۲۰۱۵ برای حل این اختلاف دکتر صفری و همکاران در موسسه دانش مغز ریکن ژاپن، این تحریک را به صورت تک نورونی و با تکنیک پیچ کلمپ سلول کامل، در حیوان زنده بررسی کردند که این پژوهش برای اولین بار در بین پژوهشگران صورت گرفته بود و در سال ۲۰۱۷ مقاله آن به چاپ رسید که تحلیل دقیق تاثیر تک نورونی می‌تواند به حل مساله کمک کند و آزمایش به این صورت می‌باشد که جریان‌های تک پس‌سیناپسی مهارتی (uIPSCs) را با جریان‌های کلی IPSCs به صورت درون‌تنی (in vivo) و برون‌تنی (in vitro) در موش‌های تراریخته با بیان پروتئین‌های حساس به نور کانال‌دوریدسین ۲ (ChR2) مقایسه کردند و دریافتند که حداقل حدود ۱۴ نورون PV^+ اتصالات قوی با سلول‌های پیرامیدال پس‌سیناپسی دارند اما تعداد بسیاری از نورون‌های SST^+ اتصالات ضعیفی داشتند. فعالیت یا مهار نورون‌های تک PV^+ پاسخ‌های بصری سلول‌های پیرامیدال پس‌سیناپسی را در ۶ از ۷ جفت‌ها تغییر می‌دهند اما نورون‌های تک SST^+ تغییر قابل توجهی را در ۸ از ۱۱ جفت‌ها نشان ندادند. همچنین نشان دادند که نورون‌های PV^+ می‌توانند به تنهایی فعالیت کنند و تاثیر بر منحنی پاسخ و میزان پاسخ نورون تحریکی داشته باشند اما ممکن است بیشتر نورون‌های SST^+ به صورت گروهی روی سلول پیرامیدال عمل کنند (۶). مدل‌سازی شبکه‌های عصبی به عنوان یک ابزار قدرتمند در درک رفتار نورون‌ها و سیستم‌های عصبی عمل می‌کند. این رویکرد محاسباتی امکان ایجاد مدل‌های عمومی و دقیق‌تری را فراهم می‌کند که قادر به شبیه‌سازی رفتار اسپایک انواع مختلف نورون‌ها با دقت بالا هستند. این مدل‌ها به تفاوت‌های ظریف در پاسخ‌های نورونی حساس بوده و درک بهتری از مکانیسم‌های عصبی ارائه می‌دهند. یکی از کاربردهای مهم استفاده از مدل‌های GLIF برای شناسایی و طبقه‌بندی نورون‌ها بر اساس رفتار دینامیک آنها مانند سرعت شلیک و پاسخ به ورودی‌هاست. این مدل‌ها امکان تحلیل دقیق‌تری از ویژگی‌های نورونی را فراهم کرده و رابطه میان ویژگی‌های زیستی (مانند زمان‌بندی اسپایک‌ها و جریان‌های یونی) و عملکرد شبکه‌های عصبی را کشف می‌کنند. این پژوهش همچنین به ساده‌سازی مدل‌های پیچیده‌ی نورونی برای شبیه‌سازی‌های بزرگ‌مقیاس با دقت مناسب و زمان کمتر می‌پردازد. علاوه بر این معیارهای کمی مانند زمان بازبایی و حساسیت به ورودی‌ها ارائه می‌شوند تا عملکرد نورون‌ها را دقیق‌تر تحلیل کنند. توسعه ابزارهای محاسباتی نیز بخشی از این

و مورفولوژی نورون ها) و عملکردی (مانند فعالیت اسپایک و پاسخ های نورونی به محرک های بینایی) در مدل های چندمقیاسی قشر بصری اولیه موش را نشان می دهد. داده ها پس از جمع آوری و پیش پردازش در یک چارچوب یکپارچه سازی قرار می گیرند تا مدل هایی از سطح تک نورونی تا شبکه های بزرگ مقیاس ایجاد شوند. این مدل ها تعاملات نورونی و جمعیت های عصبی را شبیه سازی کرده و با پارامترسازی، شبیه سازی و اعتبارسنجی تعاملات ساختار-عملکرد را بررسی می کنند. این رویکرد سیستماتیک به درک بهتر مکانیسم های عصبی پردازش بینایی کمک می کند. مدل استاندارد LIF نقطه شروعی بود که به مدل های انتگرال و آتش نشستی تعمیم یافت. در مدل استاندارد (GLIF₁) LIF، جریان تزریقی به سلول باعث افزایش ولتاژ به صورت خطی بود. وقتی ولتاژ به آستانه ثابت رسید (∞) مدل اسپایک ها و آستانه به پتانسیل استراحت نورون بازنشانی می شد. مدل GLIF₂، مدل GLIF₁ را با ترکیب قواعد بازنشانی واقعی تر (R) برای ولتاژ و آستانه ارتقا داد. تغییرات سریع پتانسیل های عمل با پویایی کندتر دنبال می شد که بر حالت نورون اثر می گذارد. ولتاژ پس از اسپایک به حالت استراحت بازنشانی نمی شد و آستانه در مقداری ثابت باقی نمی ماند. مدل GLIF₂ به این فرض ادامه می داد که اسپایک ها به اندازه کافی مشابه هستند به طوری که می توان رابطه ای بین ولتاژ و حالت آستانه قبل و بعد از اسپایک پیدا کرد.

نورون ها طراحی شده بودند که ویژگی های متنوعی دارند. از داده های سه زیرگروه اصلی نورون ها به عنوان نورون های پیرامیدال تحریکی و دو نورون اصلی مهارى به عنوان PV⁺ و SST⁺ استفاده کردیم (۵).

از موش های تراریخته زیر استفاده کردیم:

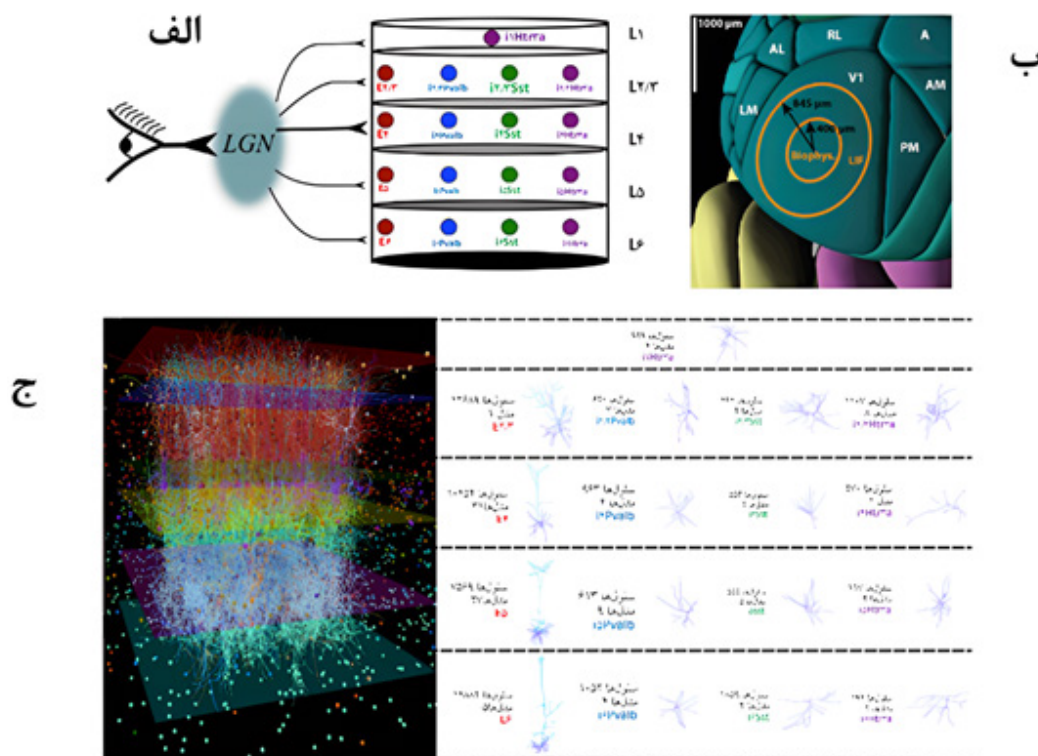
Pvalb-IRES-Cre

Sst-IRES-Cre

Nr5a1-Cre

طراحی مدل انتگرال و آتش نشستی تعمیم یافته GLIF

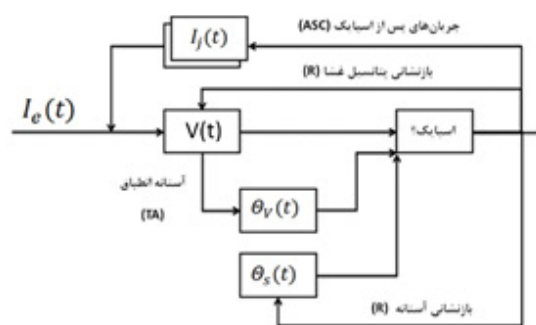
در این پژوهش از مدل های GLIF استفاده شده بود که رفتار شلیک نورون ها را در پنج سطح پیچیدگی شبیه سازی می کردند که برای بازتولید رفتار اسپایک نورون های موش و انسان استفاده می شدند. با شروع مدل انتگرال و آتش نشستی (LIF) سه تعمیم اضافه شد: الف) جریان های پس از اسپایک نشان دهنده اثرات کندتر کانال های یونی بودند. ب) ولتاژ زیر آستانه و تغییرات وابسته به اسپایک در آستانه ناشی از فعال سازی و غیرفعال سازی کانال های یونی بودند و ج) قاعده بازنشانی ولتاژ و آستانه که از داده های الکتروفیزیولوژی به دست آمده بود. روش بهینه سازی احتمال نورون مدل را با نویز ذاتی حداکثر می کرد تا اسپایک ها را دقیقاً بازتولید کند و ارزیابی می شد (شکل ۱ الف و ج) (۱۲). شکل ۱ ب نمای شماتیک از یکپارچه سازی داده های ساختاری (مانند اتصالات سیناپسی



شکل ۱. مروری بر مدل های V₁. الف: مدل ها از یک دسته تحریکی و سه دسته مهارى (Htr3a, Sst, Pvalb) ب: تصویر قشر خلفی موش، که V₁ و نواحی قشر بینایی بالاتر و ناحیه تحت پوشش مدل ها را نشان می دهد. ج: تجسم شبکه با جزئیات بیوفیزیکی

ایجاد کنیم. مدل $GLIF_4$ هم قواعد بازنشانی $GLIF_2$ و هم جریان‌های پس از اسپایک $GLIF_3$ را در خود جای می‌داد. در نهایت دیپلریزاسیون آهسته می‌توانست منجر به غیرفعال شدن جزئی جریان سدیم وابسته به ولتاژ شود که باعث ایجاد اسپایک می‌شد. در $GLIF_5$ این مکانیسم را در آستانه تطبیقی (AT) که به پتانسیل غشاء وابسته بود، وارد کردیم. توجه داشته باشید که تعداد متغیرها از یک $V(t)$ برای $GLIF_1$ به دو $\theta_s(t)$ و $V(t)$ برای $GLIF_2$ و غیره تا پنج برای $GLIF_5$ افزایش یافت (۱۳).

رابطه خطی خاص ولتاژ بعد از اسپایک به عنوان تابعی از ولتاژ قبل از اسپایک مستقیماً از داده‌های الکتروفیزیولوژیکی دریافت می‌شد. این رابطه همچنین عرض اسپایک یا "دوره بی‌پاسخی" را تعریف می‌کرد و زمانی اجرا می‌شد که در آن مدل نمی‌توانست اسپایک دیگری تولید کند. علاوه بر دوره بی‌پاسخی به دلیل مکانیسم‌هایی مانند غیرفعال شدن آهسته جریان‌های وابسته به ولتاژ که در طول اسپایک فعال می‌شدند، اغلب دشوارتر است که نوروں را دوباره بعد از اولین اسپایک



شکل ۲. خانواده مدل‌های GLIF به صورت شماتیک ارائه شده است و تعامل بین متغیرهای حالت در سطوح مختلف را نشان می‌دهد.

در زیر به اختصار تعریفی از معادله پنج مدل مختلف GLIF به صورت یکجا ارائه شده است:

$$X = \{V(t), \theta_s(t), I_{(j=1,2)}(t), \theta_v(t)\}$$

متغیرهای حالت برای سطوح مختلف معادلات مدل را نشان داده می‌شد. فرض شده بود که این متغیرهای حالتی به صورت خطی بین اسپایک‌ها هستند.

این معادله به ترتیب از چپ به راست نشان‌دهنده پتانسیل غشاء نوروں، جزء وابسته به اسپایک آستانه، دو جریان پس از اسپایک و جزء وابسته به پتانسیل غشاء آستانه بود. در شکل ۲ تعاملات بین

$$V'(t) = \frac{1}{C} \left(I_e(t) + \sum_j I_j(t) - \frac{1}{R} (V(t) - E_L) \right) \quad (1)$$

$$\theta_s'(t) = -b_s \theta_s(t) \quad (2)$$

$$I_j'(t) = -k_j I_j(t) : j = 1, 2 \quad (3)$$

$$\theta_v'(t) = a_v (V(t) - E_L) - b_v (\theta_v(t)) \quad (4)$$

جریان‌های پس از اسپایک، اسپایک و وابستگی ولتاژ آستانه بودند و a_v نیز پتانسیل غشا را با آستانه متصل می‌کرد.

که C ظرفیت نوروں، R مقاومت غشاء، E_L پتانسیل غشاء در حال استراحت، I_e جریان خارجی، $1/b_s$ ، $1/b_v$ ، $1/K_j$ ثابت‌های زمانی

اگر $V(t) > \theta_v(t) + \theta_s(t) + \theta_{\alpha}$ اسپایک تولید می شود. پس از یک بی پاسخی،

δ_i ، متغیرهای حالت با وابستگی خطی به حالت قبل اسپایک آپدیت می شدند:

$$V(t_+) \leftarrow E_L + f_v \times (V(t_-) - E_L) - \delta V \quad (5)$$

$$\theta_s(t_+) \leftarrow \theta_s(t_-) + \delta \theta_s \quad (6)$$

$$I_j(t_+) \leftarrow f_j \times I_j(t_-) + \delta I_j \quad (7)$$

$$\theta_v(t_+) \leftarrow \theta_v(t_-) \quad (8)$$

پاسخ اسپایک به دسته های مختلف طبقه بندی شدند. نورون ها بر اساس پارامترهای استخراج شده از مدل ها (مانند زمان ثابت، آستانه شلیک و جریان تطبیق) طبقه بندی شدند. از روش های تحلیل آماری برای شناسایی شباهت ها و تفاوت ها میان نورون ها استفاده شد. طبقه بندی نورون ها کمک کرد تا ویژگی های متمایزکننده انواع مختلف نورون ها مشخص شود (۱۵).

شبیه سازی شبکه های نورونی

مدل های کالیبره شده در شبیه سازی های پیچیده تری به منظور مطالعه دینامیک شبکه های عصبی و ویژگی های فردی نورون ها مورد استفاده قرار گرفته اند. هدف این مرحله بررسی چگونگی تعامل بین نورون های مختلف با ویژگی های شبیه سازی شده بود. از مدل های بهینه شده GLIF در شبیه سازی های بزرگ تر استفاده شد و تأثیر ویژگی های خاص هر نورون بر رفتار کلی شبکه بررسی شد. رفتار شلیک نورون ها را در پنج سطح پیچیدگی با ماژول شبیه سازی Allen SDK GLIF شبیه سازی کردیم. این ماژول آستانه اسپایک شبیه سازی شده نورون را هر زمان که ولتاژ از آستانه بالاتر می رفت، پتانسیل های عمل را ثبت می کرد. پتانسیل های عمل قوانین تنظیم مجدد را آغاز می کردند که ولتاژ و آستانه را بروزرسانی و (به صورت اختیاری) جریان های پس از اسپایک را تحریک می کردند. شبیه ساز Allen SDK GLIF با Python 2.7.9 که به عنوان بخشی از Anaconda Python نصب شده بود، توسعه و آزمایش می شد (۲۱-۱۸). مدل های $GLIF_1$ تا $GLIF_5$ که بر اساس اصول ریاضی و داده های تجربی توسعه یافته اند. این مدل ها به صورت سلسله مراتبی طراحی شده اند:

$GLIF_1$: شبیه سازی دینامیک پایه ای ولتاژ غشا با ورودی جریان و آستانه اسپایک.

$GLIF_2$ تا $GLIF_5$: اضافه کردن مکانیسم های بیوفیزیکی مانند تطبیق

جایی که t_+ و t_- به ترتیب نشان دهنده زمان بعد و قبل از اسپایک بودند. δV و f_v شیب و نقطه رابطه خطی ولتاژ قبل و بعد از اسپایک بودند. $\delta \theta_s$ دامنه آستانه ناشی از اسپایک پس از اسپایک زدن بود. f_j ، کسری از جریان پس از اسپایک بود که همیشه روی ۱ تنظیم می شد. δI_j ، دامنه جریان های ناشی از اسپایک بود.

تنظیم و بهینه سازی مدل ها

هر مدل GLIF با الگوریتم های محاسباتی پیشرفته بهینه سازی پارامتری تنظیم شد تا پارامترهای مدل به گونه ای بهینه شوند که با داده های واقعی مطابقت بهتری داشته باشند. این پارامترها شامل ثابت های زمانی نورون، آستانه اسپایک و جریان های تطبیقی هستند. بهینه سازی با استفاده از روش های کالیبراسیون پارامتری مانند کمینه سازی ریشه میانگین مربعات خطا (RMSE) انجام شد تا تفاوت بین ولتاژ غشایی مدل شده و داده های تجربی کاهش یابد و رفتار اسپایکینگ نورون ها به خوبی بازتولید شود.

ارزیابی عملکرد مدل ها

مدل های مختلف از نظر دقت در بازتولید پاسخ اسپایکینگ نورون ها با داده های تجربی ارزیابی شدند تا میزان تطابق رفتار پیش بینی شده با رفتار واقعی روشن شود. معیارهای ارزیابی شامل RMSE (ریشه میانگین مربعات خطا) برای اندازه گیری دقت ولتاژ پیش بینی شده، شباهت زمانی اسپایک ها برای بررسی دقیق زمان بندی اسپایک ها و دقت بازتولید الگوهای اسپایک بود که باید توانایی مدل در شبیه سازی ویژگی های خاص هر نوع نورون را منعکس کند (۱۴).

طبقه بندی انواع نورون ها

با استفاده از مدل های GLIF نورون ها بر اساس ویژگی های بیوفیزیکی و

اسپایک و بازیابی.

شبیه‌سازی بر اساس معادلات دیفرانسیل مرتبه اول انجام شد:

$$\tau m \frac{dv}{dt} = -(V - E_{rest}) + I_{input}$$

مقایسه دقت مدل‌های مختلف ($GLIF_1$ تا $GLIF_5$) و بررسی تأثیر پارامترها و پیچیدگی مدل‌ها بر عملکرد آنها به کار رفت. همچنین، روش Cross-Correlation برای مقایسه زمان‌بندی اسپایک‌های مدل و داده‌های واقعی به کار گرفته شد که نشان‌دهنده میزان شباهت رفتار دینامیکی مدل با نوروں واقعی است. توزیع فاصله بین اسپایک‌ها (ISI Distribution) برای تحلیل زمان‌بندی اسپایک‌ها و پایداری پاسخ‌های نوروں لحاظ گردید و مدل‌ها از نظر بازتولید الگوی فاصله بین اسپایک‌ها بررسی شدند. در نهایت، آزمون‌های آماری غیرپارامتریک، مانند آزمون Mann-Whitney U، برای مقایسه ویژگی‌های آماری نوروں‌ها بین دسته‌های مختلف (نوروں‌های تحرکی در مقابل نوروں‌های مهارتی) مورد استفاده قرار گرفت. در این پژوهش با ترکیب روش‌های تجربی، محاسباتی و آماری رفتار نوروں‌ها را شبیه‌سازی و طبقه‌بندی کردیم. روش کار از داده‌محوری به مدل‌محوری حرکت کرده و ویژگی‌های واقعی نوروں‌ها را با مدل‌سازی‌های دقیق مطابقت می‌دهد (۲۲).

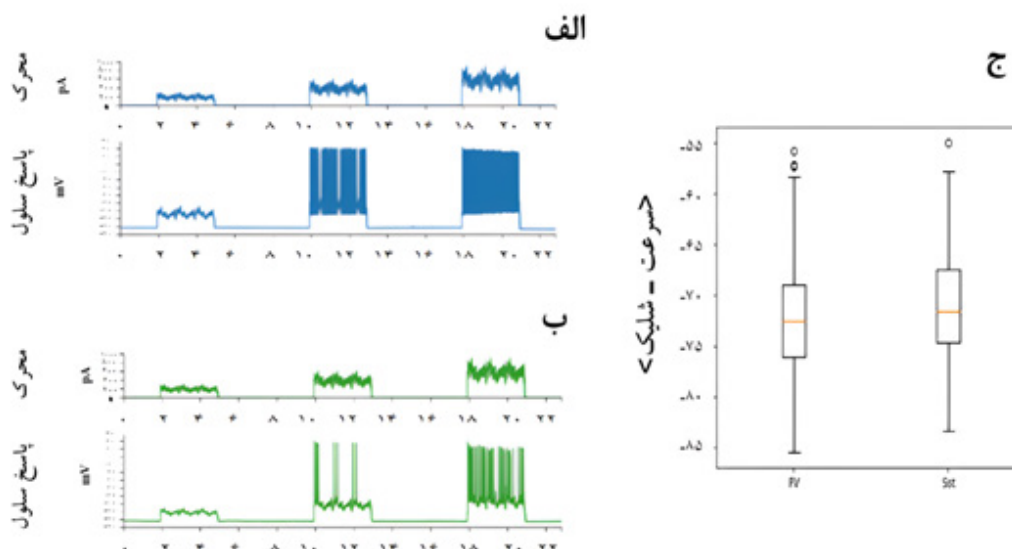
یافته‌ها

• خصوصیات نوروں‌های PV^+ و SST^+ در طول UP – state

در این بخش نتایج رفتار نوروں‌ها در حالات مختلف (UP و DOWN) و مقایسه انواع نوروں‌ها را بررسی کردیم. برای توصیف سرعت شلیک و ویژگی‌های ذاتی نوروں‌های PV^+ و SST^+ ثبت‌ها انجام شد. هر دو نوروں PV^+ و SST^+ بین حالت‌های Up و Down با هماهنگی منتقل شدند (شکل ۳ الف و ب). نوروں‌های PV^+ سرعت شلیک بالاتری نسبت به نوروں‌های SST^+ نشان دادند (شکل ۳ ج، $P < 0.001$). سرعت شلیک متوسط دو کلاس نوروں عبارت بود از: $PV^+ = 15$ ، $IQR = 9.55$ Hz، $SST^+ = 12$ ، $IQR = 13.55$ Hz.

این معادلات در مدل‌های پیچیده‌تر با تطبیق جریان‌ها و آستانه دینامیکی تکمیل شدند. جریان‌های مصنوعی و طبیعی (بر اساس داده‌های تجربی ثبت شده) به مدل اعمال شدند تا رفتار اسپایک نوروں‌ها شبیه‌سازی شود. پارامترهای مدل برای هر نوروں بهینه شدند مانند: زمان ثابت غشا (τm)، آستانه شلیک ($V_{threshold}$)، جریان‌های تطبیق ($I_{adaptation}$)، جریان‌های بازیابی ($I_{recovery}$). الگوریتم‌های بهینه‌سازی Least Squares Fitting برای تطبیق پارامترها با داده‌های واقعی استفاده شدند و بهینه‌سازی با هدف کاهش خطا بین ولتاژ مدل و ولتاژ ثبت شده نوروں انجام شد. مراحل شبیه‌سازی به این صورت انجام شد که جریان‌های پالس‌های الکتریکی با دامنه‌ها و الگوهای متنوع به مدل‌ها اعمال شدند تا واکنش نوروں‌ها تحت شرایط مختلف بررسی شود. معادلات دیفرانسیل مرتبط با هر مدل با استفاده از روش‌های عددی از جمله روش اویلر حل شدند تا دینامیک سیستم به دقت شبیه‌سازی گردد. خروجی شبیه‌سازی شامل ولتاژ غشای نوروں‌ها در طول زمان، زمان‌های شلیک اسپایک و رفتار تطبیقی نوروں‌ها بود. مدل‌های بهینه‌شده برای هر نوع نوروں در قالب شبکه‌های کوچک‌تر شبیه‌سازی شدند و ارتباطات بین نوروں‌ها با استفاده از مدل‌های ساده سیناپسی شبیه‌سازی گردید. در این فرآیند هر نوروں به عنوان یک گره در شبکه در نظر گرفته شد و نحوه ارسال و دریافت سیگنال‌ها مورد بررسی قرار گرفت. ورودی‌ها شامل پالس‌های تصادفی و ورودی‌های کنترل شده بودند و خروجی شبکه برای تحلیل دینامیک جمعی نوروں‌ها مورد بررسی قرار گرفت.

برای تحلیل نتایج شبیه‌سازی از چندین روش آماری به کار گرفته شد. RMSE (ریشه میانگین مربعات خطا) به منظور بررسی دقت مدل در بازتولید ولتاژ غشا و محاسبه تفاوت بین ولتاژ پیش‌بینی‌شده مدل و ولتاژ ثبت شده استفاده گردید. تحلیل واریانس (ANOVA) برای



شکل ۳. نشانه های UP-state به صورت تجربی در نورون های PV^+ و SST^+ . (الف) نمونه ای از Up-state در نورون های PV^+ که به طور همزمان ثبت شدند. (ب) نمونه ای از Up-state در نورون های SST^+ که به طور همزمان ثبت شدند. (ج) میانگین سرعت شلیک نورون های PV^+ و SST^+ در طول Up-state.

مهارى (۸۴ درصد) کمک می کند اما به نورون های تحریکی (۶۵ درصد) کمک نمی کند. برای تأیید این که بهبود برازش بین $GLIF_1$ و $GLIF_3$ از روندهای مختلف برای نورون های مهارى و تحریکی پیروی می کند، نسبت واریانس توضیحی دو به دو $GLIF_1$ را از مدل های $GLIF_3$ کم کردیم و دریافتیم که توزیع این تفاوت ها برای نورون های تحریکی و مهارى متفاوت است، نشان می دهد که جریان های پس از اسپایک در نورون های مهارى و تحریکی تأثیر متفاوتی دارند.

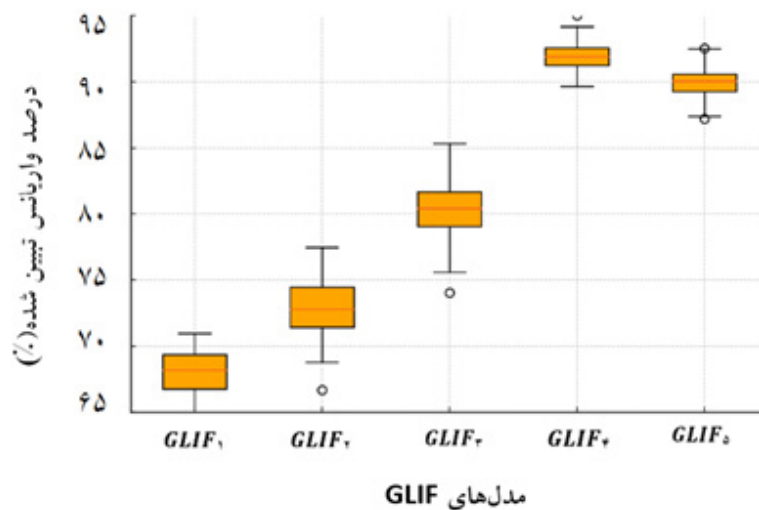
قرارگیری قاعده بازنشانی با جریان های پس از اسپایک در $GLIF_4$ عملکرد نورون های تحریکی (۷۹ درصد) و مهارى (۸۸ درصد) را بهبود می بخشد حتی اگر قاعده بازنشانی به تنهایی به عملکرد آسیب برساند. توزیع تفاوت بین $GLIF_3$ و $GLIF_4$ از نظر آماری برای نورون های مهارى و تحریکی متفاوت است. در نهایت افزودن آستانه تطبیقی همراه با جریان های پس از اسپایک و قاعده بازنشانی، بهبودی را برای سلول های تحریکی (۷۹ درصد) ایجاد می کند اما برای سلول های مهارى (۸۶ درصد) کار کمی انجام می دهد. افزودن قاعده بازنشانی که مستقیماً از داده ها اندازه گیری می شوند در واقع مانعی از توانایی مدل $GLIF_2$ برای بازآفرینی زمان های افزایش داده های عصبی می شود. ما بررسی کردیم که آیا رابطه ای بین توانایی مدل برای بازتولید رفتار زیر آستانه شکل موج ولتاژ و بازتولید زمان های اسپایک عصبی وجود دارد یا خیر. وقتی همه مدل ها را در نظر می گیریم به طور کلی بین توانایی مدل برای بازتولید ولتاژ زیر آستانه و توانایی آن برای بازتولید زمان های اسپایک همبستگی وجود دارد. با این حال زمانی که مقادیر عملکرد

• روش های تنظیم مدل ها، ارزیابی عملکرد آنها و تحلیل تأثیر مکانیسم های مختلف در بازتولید رفتار نورونی

پس از تنظیم و بهینه سازی مدل بر روی مجموعه ای از محرک های آموزش، عملکرد مدل را بر روی محرک نويز "hold-out" آزمایش می کنیم که برای آموزش استفاده نشده و این کار انجام می شود تا اطمینان حاصل کنیم مدل های ما بیش برآزش نشده اند. استفاده از داده های hold-out برای ارزیابی عملکرد مدل «استاندارد طلایی» برای انتخاب مدل است. عملکرد را بر اساس این که چقدر از واریانس زمانی در زمان های ایجاد ضربان داده ها توسط مدل در مقیاس زمانی ۱۰ میلی ثانیه قابل توضیح است، اندازه گیری کردیم. ابتدا با آزمون Friedman نشان دادیم که با مقدار P برابر با $2/56e-45$ ، اختلافات کلی بین سطوح وجود دارد. بنابراین با آزمون Wilcoxon sign-ranked ادامه دادیم که در آن مقدار P با استفاده از روش Bejamini-Hochberg برای مقایسه های چندگانه اصلاح شدند. پیشرفت واریانس توضیحی از طریق سطوح مختلف مدل نشان داد که مکانیسم های مختلف برای دستیابی به رفتار اسپایک نورون های مهارى و تحریکی مهم هستند. به طور کلی $GLIF_1$ (معادل مدل انتگرال و آتش نشستی تعمیم یافته) واریانس بالای ۶۵ درصد داشتند که همه نورون ها در نظر گرفته شدند. مدل های نورون مهارى بهتر از مدل های تحریکی عمل کردند: ۸۲ درصد در مقابل ۶۹ درصد. معرفی قوانین بازنشانی در $GLIF_2$ عملکرد مدل ها را برای نورون های تحریکی (۶۳ درصد) و مهارى (۷۶ درصد) کاهش می دهد. جریان های پس از اسپایک به تنهایی در $GLIF_3$ به نورون های

تنظیم مدل $GLIF_1$ و $GLIF_3$ مطابقت داشتند. تنها برخی از نورون‌ها تحریک لازم برای ایجاد مدل‌های $GLIF_2$ ، $GLIF_4$ و $GLIF_5$ را دریافت کردند. ۵۱ نورون الزامات اضافی مدل‌های با قوانین بازنشانی ($GLIF_2$) و $GLIF_4$ را برآورده کردند و ۴۹ نورون تمام معیارها را برای یک مدل $GLIF_5$ داشتند. میانه‌های واریانس توضیحی همه داده‌ها را می‌توانید در شکل ۴ و ۵ مشاهده کنید.

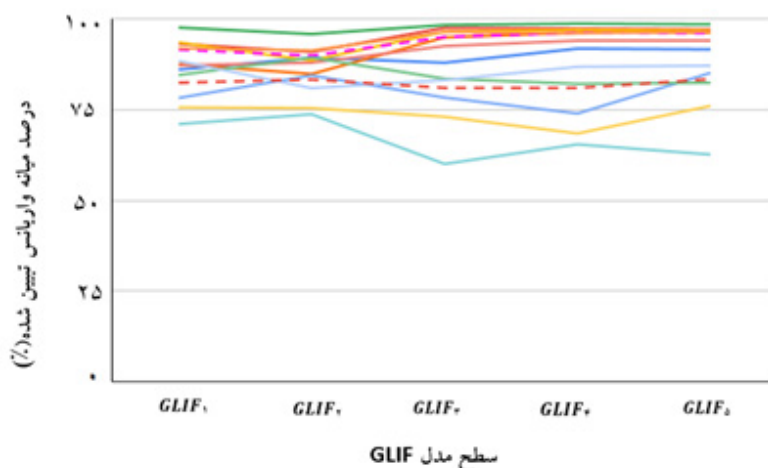
متوسط را در نظر می‌گیریم که توانایی سطوح مختلف مدل را برای بازتولید ولتاژ زیرآستانه و زمان‌های اسپایک در نظر می‌گیرند اگرچه $GLIF_2$ ولتاژ زیرآستانه را بهتر از $GLIF_1$ و $GLIF_3$ تولید می‌کند اما در بازتولید رفتار اسپایکی بدتر عمل می‌کند. بنابراین توانایی مدل برای بازتولید ولتاژ زیر آستانه لزوماً به عملکرد بهتر زمان افزایش نمی‌یابد. به طور کلی مجموعاً ۱۶۷ نورون از ۳ خط ترانسژنیک با معیارهای



شکل ۴. نمودار جعبه‌ای از پراکندگی واریانس برای مدل‌های GLIF نورون‌های مهاری.

نورون در توضیح واریانس داده‌ها در مدل‌های مشخص شده GLIF موثر هستند. مدل $GLIF_3$ بالاترین واریانس را در مقایسه با سایر مدل‌ها نشان می‌دهد که بیانگر توانایی برتر این مدل در بازتولید (capturing) دینامیک پیچیده فعالیت‌های عصبی است. این عملکرد برتر به دلیل توانایی مدل در شبیه‌سازی دقیق‌تر تغییرات زمانی و الگوهای پاسخ نورونی به ورودی‌های سیناپسی است که برای درک رفتارهای متنوع نورون‌ها در شرایط مختلف ضروری است. عملکرد مدل $GLIF_4$ نیز برای هر دو نوع نورون قابل توجه است که نشان می‌دهد تغییرات در ساختار مدل می‌تواند تأثیر قابل توجهی بر عملکرد داشته باشد.

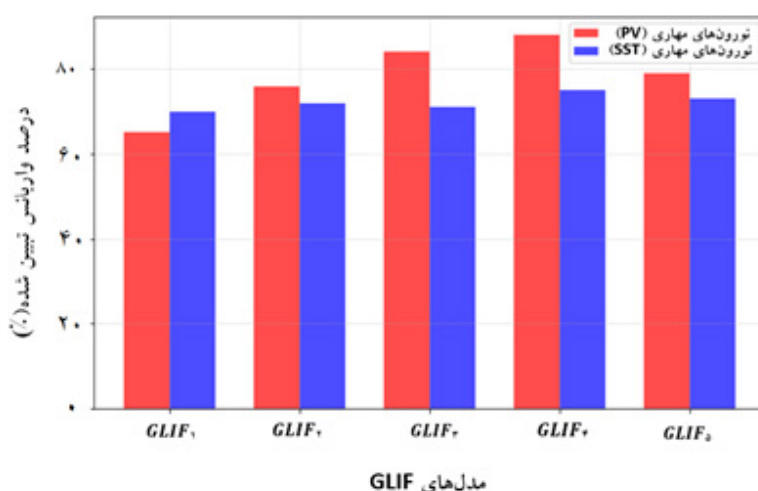
عملکرد مدل‌ها بر اساس درصد واریانس اندازه‌گیری شده است. این معیار نشان می‌دهد که مدل‌ها چقدر از رفتار زمانی نورون‌ها (مانند زمان‌های اسپایک) را توضیح می‌دهند. تفاوت بین مدل‌ها (مانند $GLIF_1$ تا $GLIF_4$) از طریق مقایسه واریانس توضیحی تحلیل شده است و همچنین تأثیر مکانیسم‌های مختلف (مانند قوانین بازنشانی، جریان‌های پس از اسپایک و آستانه تطبیقی) بر عملکرد مدل‌ها بررسی شده است. همان‌طور که در شکل ۶ می‌بینید، هر دو نوع نورون PV^+ (رنگ قرمز) و SST^+ (رنگ آبی) درصد بالایی از واریانس را در تمامی مدل‌ها نشان می‌دهند. این موضوع حاکی از این است که هر دو نوع



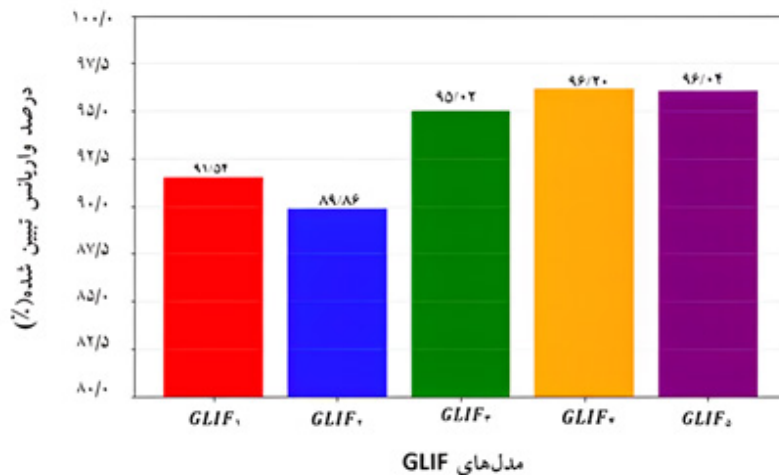
شکل ۵. مکانیسم های مختلف، عملکرد مدل را برای نورون های مهارتی و تحریکی بهبود می بخشد

جریان های پس از شلیک و قوانین بازنشانی ($GLIF_4$) برای بهبود عملکرد مدل های تحریکی ضروری بود. آستانه سازگار وابسته به ولتاژ ($GLIF_5$) عملکرد مدل های تحریکی را بیشتر بهبود بخشید اما تأثیر محدودی بر نورون های مهارتی داشت. خط قرمز نقطه چین نشان دهنده همه نورون های تحریکی، خط بنفش نقطه چین نشان دهنده همه نورون های مهارتی است

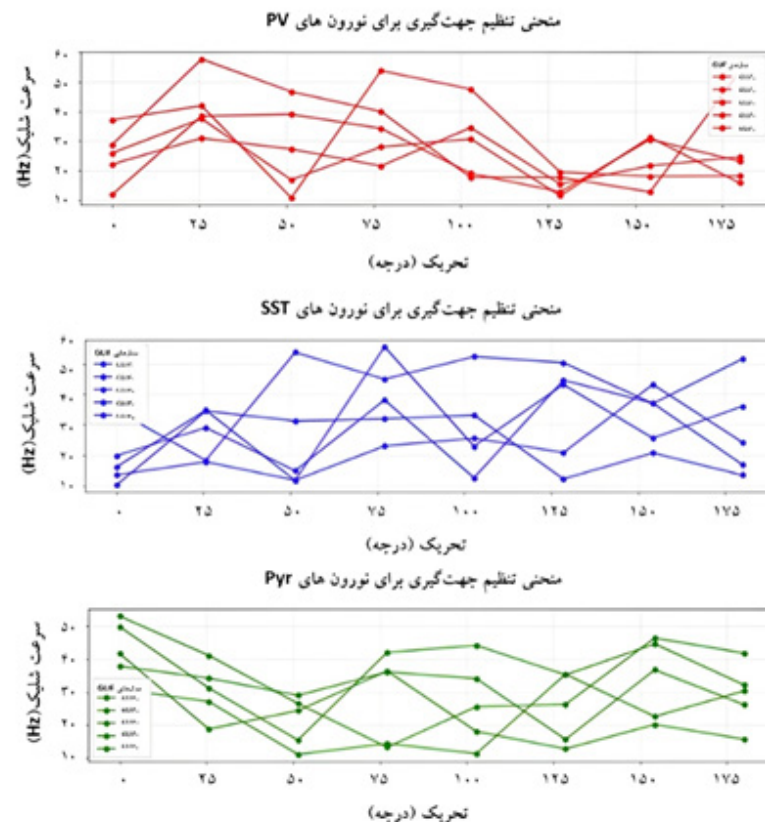
همان طور که در شکل ۵ مشاهده می شود، مدل پایه ($GLIF_1$) عملکرد قابل توجهی نشان داد. به طور کلی، مدل های مهارتی (PV و SST) در بازتولید زمان های شلیک موفق تر از مدل های تحریکی بودند. قوانین بازنشانی ($GLIF_2$) عملکرد مدل را کاهش دادند، در حالی که جریان های پس از شلیک ($GLIF_3$) باعث بهبود عملکرد مدل های مهارتی شدند. ترکیب



شکل ۶. مقایسه درصد واریانس برای مدل های GLIF. در این شکل عملکرد دو نوع نورون مهارتی (PV^+ و SST^+) در پنج مدل مختلف GLIF (از $GLIF_1$ تا $GLIF_5$) نشان داده شده



شکل ۷. نمودار میله‌ای مقایسه عملکرد مدل‌های مختلف GLIF را بر اساس درصد واریانس



شکل ۸. نمودار منحنی‌های تنظیم (Tuning Curves) برای نورون‌های مختلف (PV^+ , SST^+ و Pyr) تحت تأثیر محرک‌های زاویه‌ای متفاوت از ۰ تا ۱۸۰ درجه و بر اساس چندین مدل هستند

نورون‌های پیرامیدال را تنظیم کرده و الگوهای نوسانی شبکه را شکل می‌دهند. در مقابل اینترنورون‌های SST^+ با سرعت شلیک پایین‌تر و مقاومت ورودی بالاتر، مهار پایدار و کششی ارائه می‌دهند که تحرک نورون‌های پیرامیدال را در مقیاس‌های زمانی طولانی‌تر تعدیل می‌کنند.

اینترنورون‌ها با ویژگی‌های فیزیولوژیکی و عملکردی متمایز، تعاملات با نورون‌های پیرامیدال و دینامیک شبکه را تحت تأثیر قرار می‌دهند. اینترنورون‌های PV^+ با سرعت شلیک بالا و مقاومت ورودی پایین، مهار سریع و قوی اعمال می‌کنند. سینتیک سریع آنها زمان‌بندی شلیک

استفاده از اطلاعات دنباله اسپایک و مدل سازی روابط ورودی/خروجی و سپس خوشه بندی پارامترها است. این پژوهش با ارائه رویکردی نوین در مدل سازی نورون، بهبودهایی در دقت و کارایی مدل ها را به همراه داشته است. برخلاف مدل های پیشین که نورون ها را به صورت گروهی بررسی می کردند، این مطالعه به تحلیل فردی نورون ها پرداخته و می کوشد تا دقت مدل سازی را بهبود بخشد. با افزودن قوانین بازنشانی و جریان های پس از اسپایک به مدل های سنتی LIF، عملکرد مدل ها در پیش بینی الگوهای اسپایک ارتقا یافته است. این پژوهش بر نورون های مهاری تمرکز دارد و رویکردی جامع برای مقایسه این نورون ها با نورون های تحریکی ارائه می کند. استفاده از روش hold-out برای ارزیابی مدل ها نیز می تواند به دقت بیشتر کمک کرده و از بیش برآزش جلوگیری کند. با این حال، مدل های GLIF هنوز با محدودیت هایی مواجه هستند؛ از جمله عدم در نظر گرفتن ویژگی های پیچیده تر مانند پلاستیسیته سیناپسی و وابستگی های طولانی مدت. همچنین تمرکز بر نورون های ثبت شده از خطوط ترانس ژنیک ممکن است به تعمیم پذیری نتایج به تمام نورون های مغزی محدودیت هایی وارد کند. در آینده، پیشنهاد شده است که مدل های GLIF با افزودن مکانیسم هایی مانند پلاستیسیته سیناپسی و وابستگی زمانی به فرکانس اسپایک ها بهبود یابند. همچنین بررسی این مدل های نورونی در شرایط پاتولوژیک (مانند بیماری های عصبی) و استفاده از محرک های پیچیده تر، می تواند به درک بهتر فرآیندهای عصبی و بهبود دقت مدل های شبیه سازی کمک کند. مطالعه مدارهای عصبی به جای نورون های منفرد و مدل سازی تعاملات بین نورون های مهاری و تحریکی نیز می تواند درک بهتری از شبکه های عصبی و پردازش اطلاعات به همراه داشته باشد. در نهایت، نتایج این پژوهش نشان می دهد که دقت مدل سازی می تواند افزایش یابد و به ویژه در زمینه توسعه مدل های دقیق تر برای شبیه سازی مغز و بهبود عملکرد شبکه های عصبی مصنوعی کمک رسان باشد. با این وجود، توجه به شواهد و محدودیت های ذکر شده در بررسی نتایج ضروری است تا دقت علمی مقاله به خوبی حفظ شود.

نتیجه گیری

مطالعه حاضر با بررسی مدل های GLIF، ویژگی های اسپایک نورون های تحریکی و مهاری را تحلیل کرد. نتایج نشان داد که افزایش پیچیدگی مدل ها مانند افزودن جریان های پس از اسپایک و قوانین بازنشانی، توانایی آنها را در بازتولید الگوهای اسپایک واقعی به طور معناداری بهبود می بخشد. مدل های پیچیده تر مانند GLIF₄ تا ۸۸ درصد از واریانس داده های نورونی را توضیح دادند در حالی که مدل های ساده تر

این سلول ها به حفظ پایداری شبکه و جلوگیری از تحرک بیش از حد کمک می کنند. در مدل های GLIF اینترنورون های PV⁺ اثر مهاری قوی و لحظه ای اعمال می کنند که برای همگام سازی فعالیت های عصبی و حفظ انسجام زمانی شبکه حیاتی است. از سوی دیگر اینترنورون های SST⁺ نقش تعدیلی بیشتری دارند و بر نرخ کلی شلیک و واریانس فعالیت نورون های پیرامیدال تأثیر می گذارند. تأثیر آنها اغلب از طریق حلقه های بازخوردی است که تعادل بین تحریک و مهار را تنظیم کرده و پایداری شبکه را در طول زمان مدیریت می کنند. این تفاوت ها نشان می دهد که هر زیرگروه اینترنورونی نقش منحصر به فردی در تنظیم دینامیک شبکه های عصبی ایفا می کند.

بحث

مدل های GLIF می توانند به طور همزمان زمان های اسپایک نورون های بیولوژیکی را با تعداد کمی از پارامترها که قابل تغییرند، تولید کنند و همچنین تطبیق پیچیدگی مرتبط با ورودی جریان تا خروجی دنباله اسپایک را کاهش دهند و همچنین فهمیدیم که چگونگی پیچیدگی بر توانایی مدل ها برای همزمان بازتولید زمان های اسپایک و تمایز خطوط ترانس ژنیک از طریق خوشه بندی بدون نظارت تأثیر می گذارد. در نتیجه مدل های پیچیده تر (GLIF₃ یا بیشتر) نسبت به مدل های کمتر پیچیده در بازتولید زمان های اسپایک بیولوژیکی و تمایز خطوط ترانس ژنیک بهتر هستند. هنگامی که همه پارامترها به طور همزمان بهینه سازی شوند، افزودن مکانیسم های جدید به مدل پیشین می تواند خطای تابع را کاهش دهد؛ اما خطر overfitting وجود دارد. همچنین بهینه سازی بر اساس واریانس، دیگر محدب نبوده (این ویژگی به این معنی است که منحنی تابع هیچ گاه به سمت بالا نخواهد رفت و به همین دلیل هر حد که با معادله خطی بین دو نقطه متصل شود همیشه بالاتر از منحنی تابع قرار دارد) و نیاز به محاسبات بیشتر دارد. پیشرفت توضیح واریانس برای سطوح مختلف GLIF نشان می دهد که جریان های پس از اسپایک برای رفتار اسپایکی نورون های مهاری حیاتی اند. ترکیب این جریان ها و قوانین بازنشانی عملکرد نورون های تحریکی را بهبود می بخشد در حالی که ادغام آستانه ولتاژ فعال نیز تأثیر مثبتی بر نورون های تحریکی دارد. نورون های مهاری عملکرد بهتری در بازسازی زمان های اسپایک دارند. طول برش و تعداد اسپایک ها به عنوان پیش بینی کننده های کلیدی توانایی مدل در بازتولید زمان های اسپایک تأثیر گذارند. برای توصیف انواع سلول های مرتبط با تبدیل ورودی/خروجی در آزمایش های الکتروفیزیولوژی می توان الگوریتم های خوشه بندی را روی ویژگی های الکتروفیزیولوژیک اجرا کرد. روش دیگر

ملاحظات اخلاقی

پیروی از اصول اخلاق در پژوهش

در این مطالعه از داده‌های حیوانی منتشر شده توسط مؤسسه آلن (Allen Institute) استفاده شده است. این داده‌ها مربوط به مطالعاتی هستند که توسط این مؤسسه انجام شده و مطابق با دستورالعمل‌های اخلاقی سازمان ملی سلامت ایالات متحده (NIH Guidelines for the Care and Use of Laboratory Animals) و با تأیید کمیته اخلاقی مربوطه جمع‌آوری شده‌اند. در پژوهش حاضر، هیچ‌گونه مداخله یا استفاده مستقیم از حیوانات صورت نگرفته است.

مشارکت نویسندگان

پیش‌نویس مقاله را نویسنده اول نوشته است. پس از بازبینی و اعمال برخی اصلاحات توسط سایر نویسندگان، نسخه نهایی با مسئولیت نویسنده دوم تدوین شد.

تشکر و قدردانی

مقاله حاضر بخشی از رساله دکتری است. از کلیه اساتید و متخصصانی که ما را با نظراتشان یاری کردند سپاس‌گزاریم.

منابع مالی

این مطالعه بدون دریافت هیچ‌گونه حمایت مالی از سوی نهادهای، مؤسسات یا سازمان‌های تأمین‌کننده بودجه انجام شده است.

تعارض منافع

نویسندگان اظهار می‌دارند که هیچ‌گونه تعارض منافع مالی یا غیرمالی در ارتباط با این مقاله وجود ندارد.

مانند GLIF₁ تنها ۶۵ درصد را پوشش دادند. این بهبود به ویژه در نورون‌های مهارتی مشهود بود جایی که مدل‌ها ۸۲ درصد از واریانس را توضیح دادند در مقایسه با ۶۹ درصد برای نورون‌های تحریکی. این تفاوت نشان می‌دهد که مکانیسم‌های نورون‌های مهارتی ممکن است منظم‌تر و قابل پیش‌بینی‌تر باشند. یکی از یافته‌های کلیدی این بود که بازتولید دقیق ولتاژ زیرآستانه لزوماً به بهبود پیش‌بینی زمان‌های اسپایک منجر نمی‌شود. برای مثال مدل GLIF₂ ولتاژ زیرآستانه را بهتر بازتولید کرد اما در پیش‌بینی زمان‌های اسپایک عملکرد ضعیف‌تری داشت. این نتیجه نشان می‌دهد که مکانیسم‌های اسپایکینگ مانند قوانین بازنشانی و جریان‌های پس از اسپایک نیز باید در نظر گرفته شوند. پژوهش همچنین تفاوت‌های بین زیرگروه‌های نورون‌های مهارتی مانند اینترنورون‌های PV⁺ و SST⁺ را بررسی کرد. اینترنورون‌های PV⁺ با مهار سریع و دقیق خود از پردازش سریع سیگنال‌ها حمایت می‌کنند در حالی که اینترنورون‌های SST⁺ مهار کندتر و متعادل‌تری را فراهم می‌کنند که فعالیت شبکه را در طول مدت تثبیت می‌کند. این تفاوت‌ها اهمیت درک نقش زیرگروه‌های نورونی در تنظیم فعالیت شبکه‌های عصبی را نشان می‌دهد. در نهایت این مطالعه تأکید می‌کند که مدل‌سازی نورونی باید فراتر از بازتولید ویژگی‌های زیرآستانه‌ای باشد و الگوهای اسپایک را نیز به دقت در نظر گیرد. یافته‌های این پژوهش نشان می‌دهند که تفاوت‌های عملکردی بین نورون‌های مهارتی و تحریکی و تأثیر مکانیسم‌های مختلف بر این دو نوع نورون یکی از نقاط برجسته این مقاله است. این نتایج می‌توانند به طراحی مدل‌های دقیق‌تر برای تحلیل و شبیه‌سازی فعالیت نورونی در سیستم‌های عصبی کمک کنند و درک ما را از نقش نورون‌های مهارتی در دینامیک شبکه‌های عصبی بهبود می‌بخشند. این بهبودها ممکن است منجر به توسعه مدل‌های واقع‌بینانه‌تر برای مطالعه پردازش اطلاعات در مغز شوند.

References

1. Mao R, Schummers J, Knoblich U, Lacey CJ, Van Wart A, Cobos I, et al. Influence of a subtype of inhibitory interneuron on stimulus-specific responses in visual cortex. *Cerebral Cortex*. 2012;22(3):493-508.
2. Ascoli GA, Alonso-Nanclares L, Anderson SA, Barriónuevo G, Benavides-Piccione R, Burkhalter A, et al. Petilla terminology: Nomenclature of features of GABAergic interneurons of the cerebral cortex. *Nature Reviews Neuroscience*. 2008;9:557-568.
3. Mann EO, Paulsen O. Role of GABAergic inhibition in hippocampal network oscillations. *Trends in Neurosciences*. 2007;30(7):343-349.
4. Isaacson JS, Scanziani M. How inhibition shapes cortical activity. *Neuron*. 2011;72(2):231-243.

5. Bergoin R, Torcini A, Deco G, Quoy M, Zamora Lopez G. Inhibitory neurons control the consolidation of neural assemblies via adaptation to selective stimuli. *Scientific Reports*. 2023;13:6949.
6. Safari MS, Mirnajafi-Zadeh J, Hioki H, Tsumoto T. Parvalbumin-expressing interneurons can act solo while somatostatin-expressing interneurons act in chorus in most cases on cortical pyramidal cells. *Scientific Reports*. 2017;7:12764.
7. Qian N, Lipkin RM, Kaszowska A, Silipo G, Dias EC, Butler PD, et al. Computational modeling of excitatory/inhibitory balance impairments in schizophrenia. *Schizophrenia Research*. 2022;249:47-55.
8. Lee JH, Koch C, Mihalas S. A computational analysis of the function of three inhibitory cell types in contextual visual processing. *Frontiers in Computational Neuroscience*. 2017;11:28.
9. Oldenburg IA, Hendricks WD, Handy G, Shamardani K, Bounds HA, Doiron B, et al. The logic of recurrent circuits in the primary visual cortex. *Nature Neuroscience*. 2024;27(1):137-147.
10. Wagatsuma N, Nobukawa S, Fukai T. A microcircuit model involving parvalbumin, somatostatin, and vasoactive intestinal polypeptide inhibitory interneurons for the modulation of neuronal oscillation during visual processing. *Cerebral Cortex*. 2023;33(8):4459-4477.
11. Dzyubenko E, Fleischer M, Manrique-Castano D, Borbor M, Kleinschmitz C, Faissner A, et al. Inhibitory control in neuronal networks relies on the extracellular matrix integrity. *Cellular and Molecular Life Sciences*. 2021;78(14):5647-5663.
12. Dong Y, Mihalas S, Russell A, Etienne-Cummings R, Niebur E. Estimating parameters of generalized integrate-and-fire neurons from the maximum likelihood of spike trains. *Neural Computation*. 2011;23(11):2833-2867.
13. Mensi S, Naud R, Pozzorini C, Avermann M, Petersen CC, Gerstner W. Parameter extraction and classification of three cortical neuron types reveals two distinct adaptation mechanisms. *Journal of Neurophysiology*. 2012;107(6):1756-1775.
14. Paninski L, Pillow JW, Simoncelli EP. Maximum likelihood estimation of a stochastic integrate-and-fire neural encoding model. *Neural Computation*. 2004;16(12):2533-2561.
15. Pozzorini C, Naud R, Mensi S, Gerstner W. Temporal whitening by power-law adaptation in neocortical neurons. *Nature Neuroscience*. 2013;16(7):942-948.
16. Atallah BV, Bruns W, Carandini M, Scanziani M. Parvalbumin-expressing interneurons linearly transform cortical responses to visual stimuli. *Neuron*. 2012;73(1):159-170.
17. Lee SH, Kwan AC, Dan Y. Interneuron subtypes and orientation tuning. *Nature*. 2014;508(7494):E1-2.
18. Lee SH, Kwan AC, Zhang S, Phoumthipphavong V, Flannery JG, Masmanidis SC, Taniguchi H, Huang ZJ, Boyden ES, Deisseroth K, Dan Y. Activation of specific interneurons improves V1 feature selectivity and visual perception. *Nature*. 2012;488(7411):379-383.
19. Wilson NR, Runyan CA, Wang FL, Sur M. Division and subtraction by distinct cortical inhibitory networks in vivo. *Nature*. 2012;488(7411):343-348.
20. Allen Institute for Brain Science. Allen Cell Types Database, Technical White Paper: Electrophysiology. 2016. Available from: http://help.brainmap.org/display/celltypes/Documentation?preview=/8323525/108135/CellTypes_Ephys_Overview.pdf
21. Allen Institute for Brain Science. Allen Cell Types Database. 2016. Available from: <http://celltypes.brain-map.org/>.
22. Schreiber S, Fellous JM, Whitmer D, Tiesinga P, Sejnowski TJ. A new correlation-based measure of spike timing reliability. *Neurocomputing*. 2003;52:925-931.