

GÖRSEL ALGIYA İLİŞKİN BİR KORTEKS MODELİ

YÜKSEK LİSANS TEZİ

Mehmet Ali ANIL

Elektronik Mühendisliği Anabilim Dalı

Elektronik Mühendisliği Programı

HAZİRAN 2018

GÖRSEL ALGIYA İLİŞKİN BİR KORTEKS MODELİ

YÜKSEK LİSANS TEZİ

Mehmet Ali ANIL
(504141233)

Elektronik Mühendisliği Anabilim Dalı

Elektronik Mühendisliği Programı

Tez Danışmanı: Prof. Dr. Neslihan Serap ŞENGÖR

HAZİRAN 2018

İTÜ, Fen Bilimleri Enstitüsü'nün 504141233 numaralı Yüksek Lisans Öğrencisi Mehmet Ali ANIL, ilgili yönetmeliklerin belirlediği gerekli tüm şartları yerine getirdikten sonra hazırladığı “GÖRSEL ALGIYA İLİŞKİN BİR KORTEKS MODELİ” başlıklı tezini aşağıdaki imzaları olan jüri önünde başarı ile sunmuştur.

Tez Danışmanı : **Prof. Dr. Neslihan Serap ŞENGÖR**
İstanbul Teknik Üniversitesi

Jüri Üyeleri : **Prof. Dr. İnci Çilesiz**
İstanbul Teknik Üniversitesi

Dr. Öğr. Gör. Ali Bayram
İstanbul Üniversitesi

.....

Teslim Tarihi : **4 Mayıs 2018**

Savunma Tarihi : **Haziran 2018**





Başak'a.



ÖNSÖZ

Bana yıllardır ilginç bulduğumun peşinden koşturmamı teşvik etmiş, zorlandığım zamanlarda anlayış ile her zaman destek olmuş, zora girişmemizden hiç çekinmemiş hocam Prof. Dr. Neslihan Serap Şengör'e; bana dünyanın merak edilesi olduğunu tattıran ve keşfine destek çıkan anneme, babama ve kardeşime; bana sorgusuz sualsiz destek olan, ben vazgeçtiğimde beni ayaklandıran, beni benden hep daha öteye taşıyan biricik eşime teşekkür ederim.

Haziran 2018

Mehmet Ali ANIL





İÇİNDEKİLER

Sayfa

ÖNSÖZ	vii
İÇİNDEKİLER	ix
KISALTMALAR.....	xi
SEMBOLLER	xiii
ÇİZELGE LİSTESİ.....	xv
ŞEKİL LİSTESİ.....	xvii
ÖZET	xix
SUMMARY	xxi
1. GİRİŞ	1
2. NÖRONLAR VE DEVRELERİNE İLİŞKİN TEMEL BİLGİLER.....	3
2.1 Nöron	3
2.1.1 Membran potansiyeli	3
2.2 Aksiyon Potansiyeli ve Sinaptik Dinamik.....	4
2.3 Görsel Algı ve Dikkat ile İlgili Kavramsal Sinirbilim Tanımları	5
2.3.1 Görsel algı	5
2.3.2 Dipten yukarı ve üstten aşağı süreçler	6
2.3.3 Yanlı rekabet teorisi	7
2.3.4 Özellik entegrasyon teorisi	8
2.3.5 Görsel maskeleye	9
3. MODELLEME VE DOĞRULAMA	11
3.1 Brunel Wang Modeli	11
3.1.1 Modelin ayrıklaştırılması.....	16
3.1.2 Doğrulama benzetimi	17
4. DENEY.....	27
4.1 Görsel Dikkat Deneyi	27
5. SONUÇLAR	33
KAYNAKLAR.....	39
EKLER	41
EK A Brian ile İlgili Notlar	43
EK B Modelin sınanmasında kullanılan kaynak kodu	45
EK C Deneyde kullanılan kaynak kodu	49
ÖZGEÇMİŞ	59

KISALTMALAR

AMPA : α -amino-3-hidroksi-5-metil-4-isoxazolepropiyonik asit
GABA : γ -aminobutirik asit
NMDA : N-metil-D-aspartat





SEMBOLLER

C_m	: Membran kapasitansı
I_{syn}	: Sinaptik akım
I_x	: Sinaptik akıma x reseptörlerinin katkısı
C_{ext}	: Dış bağlantıların yapıldığı sinapsların kümesi
C_E	: İç bağlantılardan uyarıcılara olan sinapsların kümesi
C_I	: İç bağlantılardan baskılayıcılara olan sinapsların kümesi
g_x	: x reseptörlerinin iletkenliği
β	: NMDA reseptörünün nonlineer fonksiyonunun eksponenti, 0.062
γ	: NMDA reseptörünün nonlineer fonksiyonunun katsayısı, $[Mg^{2+}]/3.57$
s_{jx}	: x reseptörlerinin j sinapsındaki geçit değişkeni ya da kapı açıklık oranı
τ_C	: x reseptörünün geçitlerinin kapanma zaman sabiti
α	: NMDA geçitlerinin dinamiğinde model parametresi olan bir zaman sabiti



ÇİZELGE LİSTESİ

Sayfa

Çizelge 3.1: Benzetimde kullanılan katsayılar	19
---	----





ŞEKİL LİSTESİ

	<u>Sayfa</u>
Şekil 2.1 : Nöronlar üzerinde açıklama yaparken kullanılan bir takım tanım	4
Şekil 3.1 : AMPA ve GABA reseptörlerinin geçit katsayılarının zamana göre değişimi.....	14
Şekil 3.2 : AMPA ve NMDA reseptörlerinin geçit katsayılarının zamana göre değişimi.....	15
Şekil 3.3 : Nöronun elektriksel modeli	15
Şekil 3.4 : Kullanılan nöron modelinin Brian2 deki sinapslara tanımlanışı.	18
Şekil 3.5 : Deneyde kullanılan bağlantılandırma.....	21
Şekil 3.6 : Denge durumundaki nöronların tarama örüntüsü	22
Şekil 3.7 : Uyarılmış durumundaki nöronların tarama örüntüsü	23
Şekil 3.8 : Benzetimin denge durumundan alınan bir nöronun membran potansiyeli	25
Şekil 4.1 : Dikkat deneyinin süreci	28
Şekil 4.2 : Dikkat deneyinin bağlantı şeması.....	31
Şekil 5.1 : Odak S1 de iken V2 modülünde S1 ve S2 ye uyaranların etkisi.....	34
Şekil 5.2 : Odak S1 de iken V4 modülünde S1' ve S2' ye uyaranların etkisi.....	34
Şekil 5.3 : 8 tane deneyin sonucunda V4 grubundaki S1' ve S2' alt gruplarındaki aktivitenin ortalaması. Bu ortalama Deco ve Rolls (2005) dan farklı olarak odağa bağlı bir fenomen gözlemlenmemiştir.	37



GÖRSEL ALGIYA İLİŞKİN BİR KORTEKS MODELİ

ÖZET

Görsel algı, beyinde incelenen en etraflıca çalışılmış ve ulaşılabilir algısal konulardan biridir. Görsel içerik oldukça kolay manipüle edilebildiğinden, yanlış algılaması ilüzyon adı altında kültürel içerik bile olabilmektedir. Son zamanlarda görsel sanal gerçeklik insanı tamamen sarabilmeye başladığından, görsel algının diğer bilişsel konular ile ilintisini araştırabilmek için farklı imkanlar doğmaktadır.

Beyinin içindeki görsel devre, kısa dönem belleği, dikkat, alt katman görsel işleme gibi vasıflar arasında, obje farkındalığı gibi yüksek katman fonksiyonları oluşturacak şekilde bir köprü vasfı görür. Beyindeki belli yerlerin gerek deney ile gerek bir hastalık ile fiziksel olarak zarar görmesinin sonucunda elde edilen bilgiler ışığında fizyoloji, hangi bölgelerin ve sinir demetlerinin hangi işleve katkıda bulunduğu hakkında açıklamalara sahiptir. Bu bağlantıların görsel algının ve dikkatin temel yapı taşlarını sergileyebileceğinin alçak katmandaki açıklaması halen bir araştırma konusudur.

İhtiyacımız olan, süregelen etkinlik, çeldirici uyartılara bağışıklık gibi temel özelliklerin sağlanabilmesi için biyolojik akla yatkınlığa sahip bir modelle oluşturulmuş bir benzetim kullanılabilir. Doğrusal olmayan topla ve ateşle nöron modelleri ile bu fenomenlerin bir çoğu basite indirgenmiş halleri ile yakalanabilmektedir. Aynı nöron modellerinin görsel algı ve dikkat için de açıklayıcı güce sahip olma ihtimali bulunmaktadır. Bu çalışmada, bu ihtimal üzerinde durulacak ve uygunluğu teste tabii tutulacaktır.



A CORTEX MODEL OF VISUAL COGNITION

SUMMARY

Visual perception is one of the most widely studied and accessible cognitive subjects in the brain. Since visual sensory stimuli are easily craftable, the misperception of them have been subject to human content such as illusions. It also has been lately possible to encapsulate a person within a visual virtual reality, to dive deeper into the endeavor of understanding visual cognition, and how it is connected to other subjects of scientific interest. The visual circuitry within brain exhibits a tight gameplay between short term memory, attention, and low-level visual processing, facilitating higher level functions such as object recognition. Physiology has explanations on how this circuitry exhibits these higher-level functions by singling out zones within the brain and bundles of neurons through losses of functionality by either experimental means or by an intrusive malady. The low-level explanation of how an interconnect of neurons can exhibit the primitives of visual cognition and attention is still under research.

Simulations with marginal biological plausibility can be used in order to implement core features that might be required, such as persistence of activity, distractor immunity for cortical networks. Many of these phenomena are exhibited in a simplistic manner with integrate and fire neurons with nonlinearity. There is a possibility that same neuron models might be explanatory for the case of visual attention and perception. In this work, this possibility will be pursued and put to test.

In this work, we construct a simulation of the visual cortex that relies on the Brunel and Wang integrate and fire neurons and put them into a trial to whether they can exhibit some higher neuroscientific phenomena, such as backward masking, or biased competition. The neuron is modeled with an electrical model, in which every neuron has a leaking membrane capacitance, individual voltage-dependent non-linear channel conductivities, and uses integrated spike trains in as excitation for channel gating variables, a figure that is based on the ratio of open channels. Several receptors play role on the dynamics, where particularly NMDA, an excitatory receptor that is activated nonlinearly with respect to the membrane potential, which the consensus is that it plays key role in mechanisms of attention within the brain.

The experiment uses a group of 1000 neurons that are all modeled as mentioned, and all connected to each other (except self-connections), making up close to 1000000 connections, thus synapses, to compute. 200 of these neurons are interneurons that suppress excitation, which excites only GABA receptors that depolarize, and move the postsynaptic neuron away from excitation. The rest are excitatory neurons, which use AMPA and NMDA receptors that hyperpolarize, thus initiate excitation in the postsynaptic neuron. This interaction is modeled with delta functions which signify the events of spiking in the presynaptic neuron. These delta function are integrated, thus instantaneously increment postsynaptic channel activation variables, and this activity is

damped down with a measured time constant of that particular channel. Each activated postsynaptic channel lets through a current that depolarizes the membrane potential. When these excitations accumulate in the postsynaptic neuron such that it results in a membrane potential higher than a threshold value of -50mV, the postsynaptic neuron "fires" and instantaneously polarizes back its membrane to a steady state of -70mV. This is the singular phenomenon of how a spike or an excitation is generated.

The experiments include several neuron groups that have been assumed to have undergone a process of Hebbian learning, which has resulted in a static nonhomogenous connectivity between these groups. It is important to emphasize that the learned information, or association is a given trait in this experiment. Execution of a learning process might be a way to extend the scope of this experiment. Inhibitory neurons are connected to every excitatory neuron without differentiating which group it belongs to, because it is necessary for reasons of general stability that every excitatory neuron be repressed when excited. On the other hand, excitatory neurons are separated into 5 groups of 80 (S1, S2, S3, S4 and S5) and a group of 400 (Sn). The groups of 80 each represent a trait, which might be a shape or position, whereas the group of 400 is in place just to keep the excitation present in the population.

All weights (that multiply the gating variables thus affect the induced current) of these 1000000 synapses are predetermined representative of this phase of a Hebbian training. The neurons in the same groups (e.g. S1 to S1) that have object representation are connected to each other in an elevated weight. The groups that represent different representations have a diminished weight (e.g. S1 to S2). This way, it becomes advantageous for groups to exhibit recurrent spiking, and when one group is excited the others are effectively inhibited, therefore the overall collective dynamic forms attractors (steady states) in which one of the representative groups are in recurrent spiking. In order to keep the population in the verge of excitation, a noise with a Poisson distribution of 2.4 kHz is supplied to every neuron. This elevates the rate of excitation to a level that only a few incoming spikes is enough to initiate a post-synaptic spike. This frequency of cumulative spontaneous excitation is found to be biologically plausible, shown to be consistent with in-vivo measurements.

After the model is constructed, several computational complications are eradicated, and seemingly, a dynamical plausibility is achieved, a test experiment is executed. In this test experiment, a stripped down version of binary behavior is exhibited. When the aforementioned setup is built, when S1 is excited by an elevated noise source, the excitation rate of S1 elevates, and remains elevated after the excitation is lifted. When a distracting excitation floods every single neuron within the setup, this persistent behavior is eliminated, and the system goes back to its initial state. Therefore, the population can exhibit the function of a memory, and the behavior is in line with the literature that focuses on the dynamical aspect of these neurons.

To exhibit biased competition, two groups of 1000 neurons with the same arrangement, but with 2 groups of 80 (S1 and S2) and a group of 640 (Sn). One group represents V2, and the other V4 in the visual neural pathway. The V4 group of 1000 has the same arrangement, but the groups are denoted with S1', S2' and Sn'. In this arrangement, individual groups are connected in a similar manner with the previous test experiment, and the groups are connected that favor associated groups over cross-associated groups (e.g. $S1 \rightarrow S1'$ to $S1 \rightarrow S2'$) and favors forward connections to backward connections (e.g. $S1 \rightarrow S1'$ to $S1' \rightarrow S1$)

In this experiment, it is investigated whether a bias in one of the representations (e.g. S1) creates easier excitability in the associated V4 group, S1'. We compared this case with the case where a bias in one of the representations creates easier excitability in the cross-associated V4 group. Although there were single occurrences of induced biases, we failed to create a statistically relevant case for biased competition in this arrangement. When an average is taken over multiple instances of this experiment, these two cases fail to exhibit a difference.

In this work, the simulations are executed with Brian2 libraries, which opens up possibilities for further work to be done with similar experimental arrangements, with similar underlying biological principals. The work can be downloaded and reproduced by downloading the digital notebook, and the work can be seamlessly improved. Possible improvements might be inclusion of a Hebbian training phase, use of real images with their Gabor representations, or further work in order to reveal top-down biasing with the models that have been introduced and coded in this work.





1. GİRİŞ

İnsanın etrafındaki dünyayı kavrayışı, birbirlerine ters yönelimlere sahip bilişsel süreçleri birlikte yaşaması sayesinde mümkün olmaktadır. Doğduğundan itibaren bireyin yarattığı öz-imesi, gerçeklik algısı, anıları gibi yapıları maruz kalınan duyusal bilginin anlamlandırılmasında tamamlayıcı rol oynamaktadır. Bu kavramsal ikililik, duyular yolu ile gerçekliğin algılanışı ve deneyim ile oluşturulmuş modellerin gerçekliği sınaması, insanın değişen çevre koşullarında ve sosyal yapılarda bir birey olarak kalabilmesine ve karmaşık davranışlar sergilemesinde merkezlidir. Bu ikililiğin birlikte çalışamamaya başladığı durumlar, mesela nörolog Oliver Sacks'ın "Karısını Şapka Sanan Adam" Sacks (2009) gibi eserlere konu olmuş görsel agnozi vakaları, sinirbilim ile psikolojinin kavuştuğu, etraflı açıklamasının iki daldan da gelişme ile mümkün olacağı, merak uyandıran konulardır. Bunun yanı sıra öz imgenin, benlik algısının kaybedildiği durumlar neredeyse ilham kaynağı kabul edilebilecek şekilde güncel birçok içeriğe konu olmaktadır. (Aviv, 2018; Taylor, 2008) Aynı zamanda zaman geçtikçe gerçeklik algısını daha başarılı kandırabildiğimiz teknoloji araçları vasıtası ile soyut kavramlarla çalışabilecek ortamlar elde edilmektedir ve yeni psikolojik deneylere kapılar açılmaktadır. (Rothman, 2018, April 2, 2018)

Sinir hücresinin temel birim olarak alınıp popülasyonlarını bağlantılandırılarak temel bazı işlevsel davranışların gösterimi, bunların biyolojik görüntülendirme, deneysel ölçümler ve sinir ağı haritalandırma eforları ile temas halinde ve farklı anlayışları birleştirecek şekilde yapılması, bilişin anlamlandırılmasında tümevarımcı bir yaklaşımdır. Bu çalışmadaki tutum da bu yönde olacaktır, ama filizlenmekte olan bu tutumun açıklayıcılığının sınırlarını unutulmamasına özen gösterilecektir.

Görsel biliş, bahsi edilen ikililiğin gözlemlenebildiği, daha karmaşık bilişsel süreçlere nazaran daha kesinlikli tanımlanabilmiş, görme duyusunun sağladığı arayüz ile oldukça üzerinde çalışılabilmiş bir konudur. (Kandel, 2013) Bu konu, görsel dikkat, görsel çalışma belleği, görsel uzun süreli bellek şeklinde üçe ayrılabilir. Kabul gören

kanı, görsel algılayışın bu üç yapının birlikte kullanıldığı, yapıcı bir süreç olduğudur. Bu konuda yapılan çalışmalar, bu algılayışın fizyolojik temelinin görsel yol (akış) adının verildiği, *ne (ventral)* ile *nerede (dorsal)* şekilde dallanmış yollara sahip olan bir yapı olduğunu göstermektedir. Detayları göz ardı edilirse görsel alan, öncelikle seri bir alçak katman işleme tabi tutulup, yön, renk, karşıtlık, disparite ¹, hareket yönü gibi temellerine ayrılmaktadır. Ancak bu ilk katmanda gerçekleşen işlemlerden sonra şekli farketme gibi daha bileşik işlemleri, benzer bileşik işlemler vasıtası ile de obje tanıma gibi en üst düzey işlevleri yerine getirebilmektedir. (Kandel, 2013) Bu çalışmada, görsel bilişin, ağırlıklı olarak görsel dikkati oluşturan yapısının hesaplamalı sinirbilimsel temellerini modellemeye gayret edeceğiz. Bilişin tamamlayıcı işlemlerine atıfta bulunarak konuyu önce yapıtaşları olan nöronlar üzerinden, sonra da üst düzey bilişsel süreçler üzerinden, yani sırası ile tüme varımcı ve tümenden gelimci tutumlar ile ele alalım.

Çalışmanın çıktılarından biri olan bu tez çalışmasında, Bölüm 2’de modelimizin yapı taşı olan sinir hücresinin konumuz ile ilgili özelliklerini, ve dinamik davranışını açıklayarak başlayacağız. Bölümün devamında, görsel algının işleyişi ve sinirbilimde algı ile ilgili teoriler ve tanımlar vereceğiz. Bunlara, her ne kadar benzetim ile test edilmesi güç bilişsel süreçler ile alakalı da olsa, ilerleyen zamanlarda yapılabilecek deneylerin kurgulanmasına ve konu ile ilgili düşünce deneylerine yön verebileceği için yer verdik.

Bölüm 3’de benzetimlerde kullanılan matematiksel modeli detaylı bir şekilde tanıttık. Bu tip modellerde gerçeğe uygunluğunun kontrolü zor olduğundan bir referansta detaylı bir şekilde tanımlanmış bir deneyi tekrarladık, ve modelin doğru çalıştığını gösterdik.

Bölüm 4’de, elimizdeki sinir hücresi modeli ile oluşturduğumuz bir sinir grubu ile, daha karmaşık bir süreç olan dikkat verme sürecinin çok basitleştirilmiş bir şeklini gözlemlemeye çalıştık. Kullandığımız modeller ve varsayımlar ile bu süreci gözlemleyemedik.

¹iki gözün görsel alanındaki görece uyumsuzluk, bir nevi derinlik algısının bir öncülü.

2. NÖRONLAR VE DEVRELERİNE İLİŞKİN TEMEL BİLGİLER

2.1 Nöron

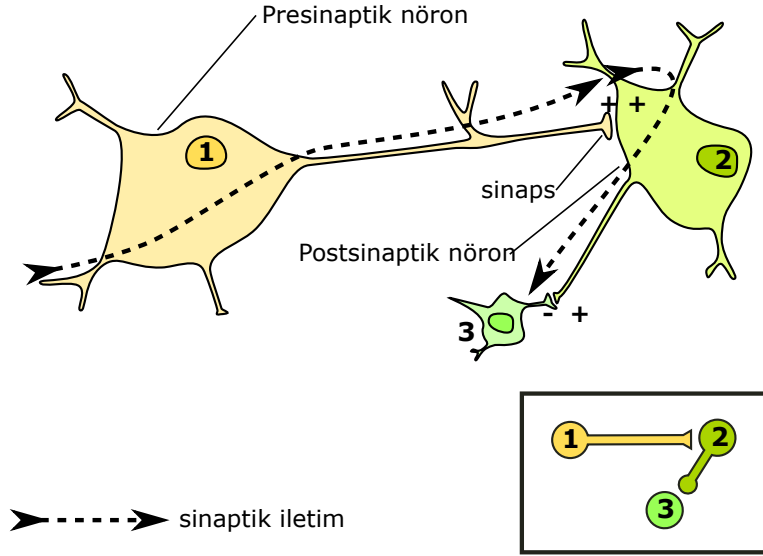
Nöron hücresinin, birbirinden fizyolojik farklar gösteren tipleri olsa da, davranışının belli başlı taraflarını yakalayan, farklı doğruluklarda modelleri mevcuttur. Çalışmada kullanılan model, fizyolojik temellere yakın olduğundan sinir hücresinin bazı temel özellikleri tanıtılacaktır. Burada kapsamlı olmaktansa çalışmada bahsi geçen kavramların tanıtımı amaçlanmıştır. Daha biyoloji temelli bir anlatım için Kandel (2013) de ifade edilen sayfalara göz atılabilir.

2.1.1 Membran potansiyeli

Sinir hücreleri içinde bulundukları ortamdaki iyonları, hücre zarı vasıtası ile çözeltideki iyonların akışını sınırlayacak şekilde ayırıştırır. (Kandel, 2013, p. 127) Fiziksel olarak yalıtılmış yüklerin denge durumu, elektrik devrelerindeki kapasiteler gibi bir potansiyele neden olur. Geçirgen bir zarda, iyonlar derişimin yüksek olduğu taraftan düşük olduğu tarafa doğru akacaktır (difüzyon vasıtası ile). Bu akışın neden olduğu yük farkı ise bir elektrik alanına oluşturup tersi yönde bir iyon akışına neden olur (elektrostatik kuvvetten dolayı). Difüzyon ve elektrostatik kuvvetin neden olduğu yük akımlarının eşit olduğu durumdaki potansiyele **membran potansiyeli** denir. Tek iyonun bulunduğu denge durumundaki membran potansiyeli, Nernst denklemi Denklem 2.1 ile hesaplanır.

$$E_k = \frac{RT}{zF} \ln \frac{[K^+]_{out}}{[K^+]_{in}} \quad (2.1)$$

Birçok iyonun olduğu durum için, her bir iyonun birbirine göre zardan ne kadar kolay geçtiğinin ölçütü olarak geçirgenlik ($P_x, [cm/s]$) tanımlanır. Dolayısı ile bu durumda membran potansiyeli Nernst denkleminin genelleştirilmiş bir hali olan Goldman denklemi Denklem 2.2 ile hesaplanabilir. (Kandel, 2013, p. 135)



Şekil 2.1 : Konu itibarı ile, nöronlardan bahsederken kullanılan dil iki nöronun arasındaki iletişimi sağlayan sinapsa göre tanımlanmıştır. Şekilde işaretlenen sinapsa göre **1** numaralı nöron sinaps öncesi, yani presinaptik **2** numaralı nöron ise sinaps sonrası, yani postsinaptik nöronudur. Sinaptik iletim, dolayısı ile ok ile gösterildiği gibi gerçekleşir. **1** numaralı nöron, sinaptik boşluğa NMDA ya da AMPA agonistleri (tetikleyicileri) bıraktığından **2** numaralı nöron uyarılacaktır. Benzer şekilde, **2** numaralı nöron, **3** numaralı nörona temas ettiği sinaptik boşluğa GABA agonisti bıraktığından **3** numaralı nöronu baskılayacaktır. Çerçevenilmiş resimde bu nöronların bu çalışmada kullanılacak olan gösterimlerini görebilirsiniz. Uyarıcı sinapslar üçgen tabanı, baskılayıcı sinapslar yuvarlak uçlar ile gösterilmiştir.

$$E = \frac{RT}{F} \ln \frac{P_K[K^+]_{out} + P_{Na}[Na^+]_{out} + P_{Cl}[Cl^-]_{in}}{P_K[K^+]_{in} + P_{Na}[Na^+]_{in} + P_{Cl}[Cl^-]_{out}} \quad (2.2)$$

2.2 Aksiyon Potansiyeli ve Sinaptik Dinamik

Bu çalışmada, sinir hücresinin sinapslarındaki geçitlerin geçirgenliğinin değişken olduğu bir model kullanılmıştır. Dolayısı ile, merkezi sinir sistemi nöronlarında sinaptik iletimin nasıl olduğunu anlatmalıyız. Sinapstan geçen sinyalin kaynağı nörona presinaptik nöron, hedefi olan nörona postsinaptik nöron diyelim. (Şekil 2.1)

Nöronlar, membran potansiyelleri denge durumunda iken, membranda üretilen küçük bir akım ile küçük bir pozitif membran potansiyeli artışı yaşarlar. Bu artış belli bir değerin altında kaldığında, iyonlar membran potansiyelini denge durumuna geri getirecek şekilde akarlar, ve *repolarizasyon* meydana gelir. Bu akım belli bir değeri aştığında (-50mV), ¹ membranın iyonlara olan geçirgenliği gerilime bağlılığından

¹Bu çalışmada, membran potansiyelinin azalması ya da artmasından bahsedildiğinde potansiyelin mutlak değerinden bahsedilmektedir.

dolay, *aksiyon potansiyeli* adı verilen bir potansiyel sıçraması yaşarlar (Izhikevich, 2007, Bölüm 2.3). Zaman içerisinde tekrarlı, bir döngü içerisinde ya da bir kere olabilecek bu olay, nöronların topluluk halinde ve bir mantığın temelini oluşturabilecek sahip dinamikler gösterebilmesini sağlar ². Komşu nöronların hücre içi ve dışındaki iyon derişimleri ne olursa olsun bu iki nöron arasındaki iletişim sadece aksiyon potansiyeli vasıtası ile gerçekleştiğinden, ve bir nörondan diğerine bilgi aktarımını sağladığından analizimizde *sinyal* olarak ele alınacaktır.

Bir nöronun diğer bir nöronu depolarize edip bu nöronu aksiyon potansiyeli yaşayacak duruma yaklaştırmasını uyarıcı bir etki, hiperpolarize edip bu olası durumdan uzaklaştırmasını ise baskılayıcı bir etki olarak ele alalım. Bir kası uyaran bir nöronda sadece bir uyarıcı ateşlenme kasın kasılmasına neden olacak kadar etkilidir. Bunun yanı sıra merkezi sinir sisteminde bir potansiyel sıçramasına neden olmak için bağlı olduğu uyarıcı nöronlardan 50 ila 100 tanesinin ateşlenmesi gerekebilir.

Merkezi sinir sistemideki nöron çiftleri, uyarıcı ve baskılayıcı olacak şekilde farklı süreçler ile birbirlerini etkileyebilirler. Her ne kadar iki nöron arasındaki bazı kimyasallar (nörotransmitterler) sinapslarda hem uyarıcı hem de baskılayıcı etkiye neden olabilse de çoğu nörotransmitter bir etki için özelleşmiştir (Kandel, 2013). Buradan yola çıkılarak, çalışmada kaba bir uyarıcı ve baskılayıcı nöron ayrımı yapacağız. Merkezi sinir sisteminde dominant olarak γ -aminobutirik asit (GABA) baskılayıcı reseptörleri, Glutamat ise uyarıcı reseptörleri (AMPA ve NMDA reseptörleri) tetiklediğinden, dominant olarak bu reseptörlere sahip nöronları uyarıcı ve baskılayıcı diye iki sınıfa ayıracağız. Şekil 2.1’de iki tip sinir hücresi de görülebilir.

2.3 Görsel Algı ve Dikkat ile İlgili Kavramsal Sinirbilim Tanımları

2.3.1 Görsel algı

Karmaşık görsel ortamlardaki nesneleri tanımlamak, yapay görme sistemlerinin henüz kopyalayamadığı olağanüstü bir hesaplama başarısıdır. Nasıl gördüğümüzün altında yatan mekanizmalar henüz tam açıklığa kavuşmamış durumdadır. Görme yalnızca nesne tanıma için değil, aynı zamanda hareketlerimize rehberlik etmek için de

²Sinir hücrelerinin biyolojilerinden kaynaklanan dinamiği kullanarak mantık operasyonlarının elde edilebileceği (Kandel, 2013, Ek E) de ele alınmıştır.

kullanılır. Bu ayrı işlevlere en az iki paralel ve birbirleriyle etkileşimli yol aracılık etmektedir. Görsel sistemdeki bu paralel yolların varlığı, kavrayış konusunun temel sorularından birini oluşturmaktadır. Ayrık yollar tarafından taşınan farklı bilgi türleri nasıl tutarlı bir görsel imge haline getirilir?

Genellikle kamera işleyişi ile karşılaştırılmasına rağmen görme retinadaki iki boyutlu görüntülerden yola çıkarak dünyanın üç boyutlu bir temsilini oluşturabilmesi ve farklı görsel koşullarda bile aynı algıyı üretebilmesi nedeniyle kamera analojisi ile açıklanamayacak kadar karmaşık bir süreçtir. Kameranın tersine, beyin, sahneleri farklı bileşenlere ayırır, ön planı arka plandan ayırır, hangi ışık uyaranlarının bir nesneye ait olduğunu ve diğerlerinin ne olduğunu belirler. Bunu yaparken, dünyanın yapısı hakkında önceden öğrenilmiş kuralları kullanır. Gelen görsel sinyal akışını analiz ederken, geçmiş deneyime dayanarak gözlere sunulan sahneyi tahmin eder. Bu nedenle, görme algısını sağlayan işlemler birleştirici bir doğaya sahiptir.

Görsel algı, retina, talamik çekirdekler ve serebral korteksin çoklu alanları arasındaki etkileşimi içerir. Retina, ince ayrıntıları çözme yeteneği, küçük hareketlerin ayrımı ve ışığın frekansındaki, karşıtlığındaki farklılıkları tespit etme kapasitesi gibi yetenekleri sayesinde görüşün sınırlarını tanımlar. Görsel korteks, karmaşık sahnelerden gelen bilgileri edinir, tek tek nesnelere ait yüzeylere ve konturlara ayırıştırır ve bu nesneleri arka planlarından parçalar. Bu süreç, oryantasyon, hareket yönü ve renk gibi yerel özelliklerin eş zamanlı bir analizini ve bu özelliklerin uzayda bütünleşmesini de içerir.

Görsel korteks, ventral ve dorsal akışlar ile birbirine bağlıdır. Ventral akış, temsil edilen nesnelerin doğası hakkında bilgiler sağlarken dorsal akış, okülomotor ve iskelet motor sisteminin hareket için kullanabileceği bilgileri sağlar. Dorsal akışın kritik bir fonksiyonu, daha ileri analiz için görsel alanda bulunan nesneleri seçmek için gereken dikkati yönlendirmektir (Ungerleider & Pessoa, 2008).

2.3.2 Dipten yukarı ve üstten aşağı süreçler

Duyum ve algı ile ilgili iki genel süreç vardır: dipten yukarı süreçler (*bottom-up processes*) ve üstten aşağı süreçler (*top-down processes*). Dipten yukarı süreçler, duyuusal bilgilerin içeri girişi sırasında işlenmesini ifade eder. Örneğin, bir görsel örüntü gösterildiğinde gözler özellikleri algılar, beyin bunları birleştirir ve örüntü algılanır. Görülen, yalnızca içeri giren duyuusal bilgiye dayalıdır. Dolayısıyla dipten yukarı

süreçlerde algı duyuşal bilginin en bayağı ve temel özelliklerinin algılanmasına neden olur. Üstten aşağıya süreçler ise biliş tarafından yönlendirilen algıyı ifade eder. Beyin, bildiğı ve algılamayı beklediğini uygulayarak, söz gelimi, boşlukları doldurur.

Üstten aşağıya süreçler, beynin frontal loblarında yer alan yönetici işlevleri tarafından beslenir. Öte yandan dipten yukarı süreçler algının temel bileşenlerine aracılık eden posterior korteks modüllerinden kaynaklanır.

Sinirbilimde, nedensellik izlenimlerinin altında yatan nöral süreçler konusunda bir tartışma vardır. Bazı görüşler algısal (dipten yukarı), diğerleri ise bilişsel/çıkarımsal (üstten aşağıya) yaklaşımları destekler. Dipten yukarı nedensellik algısının nöral temelleri ile ilgili olarak literatürde, nedenselliğin ilişkisel bilgi tipi olarak hedeflendiğı sadece birkaç nörobilimsel çalışma bulunmaktadır (Wende & Jansen, 2015).

2.3.3 Yanlı rekabet teorisi

Yanlı rekabet teorisi, *İng: biased competition theory* görsel dikkat ve görsel temsil ile ilgili sayısız hayvan ve insan çalışmalarını motive eden, etkili bir nöral sinir kuramı olmuştur. Yanlı rekabet teorisi, görsel alandaki her nesnenin kortikal temsil ve bilişsel işlem için rekabet ettiğı fikrini savunur. (Desimone & Duncan, 1995; Koch & Ullman, 1987) Bu teori, görsel işlem sürecinin, dikkat ve ileri işleme için bir nesnenin veya tüm öğelerin belirli özelliklerini önceliklendiren dipten yukarı ve üstten aşağıya sistemler gibi diğer zihinsel süreçler tarafından yanılabileceğini öne sürmektedir. Yanlı rekabet teorisi ile basitçe ifade edilen, nesnelerin işlenmek için birbirleriyle rekabetidir. Bu yarışma, genellikle görsel alanda devam eden nesneye veya alternatif olarak davranışla en ilgili nesneye doğru yönlendirilebilir. Üç en temel ilkenin lehine nöral kanıt bulunmaktadır: görsel sistemdeki temsilin rekabetçi olduğı, üstten aşağı ve dipten yukarıya eğilme mekanizmalarının rekabeti etkilediğı ve bu rekabet beyin sistemleri arasında bütünleşik olduğı (Beck & Kastner, 2009).

Seçici dikkatin yanılabet teorisi, görsel dikkat ve görsel temsil ile ilgili sayısız hayvan ve insan çalışmalarını motive eden, etkili bir sinir kuramı olmuştur. Şimdi en temel ilkelerinin üçünün lehine nöral kanıtlar var: görsel sistemdeki temsilin rekabetçi olduğı; yukarıdan aşağıya ve aşağıdan yukarıya eğilme mekanizmalarının devam eden rekabeti etkilediğı; ve bu rekabet beyin sistemleri arasında bütünleştirilmiş olduğı. Bu

üç ilke ve özellikle de bu özgün ilkelerden türetilen nöral tahminle ilgili altı daha belirli bulgulara işaret eden kanıtları gözden geçiririz (Beck & Kastner, 2009).

2.3.4 Özellik entegrasyon teorisi

İnsan görsel dikkatine dair en etkili psikolojik modellerden biri olan özellik entegrasyon teorisi *İng: biased competition theory* dikkatin görünürdeki her bir uyaran için seri olarak yönlendirilmesi gerektiğini önermektedir. Bu gereklilik olası nesneleri karakterize etmek veya ayırt etmek için birden fazla ayrılabilir özelliğin birleşimine ihtiyaç duyduğunda ortaya çıkar (Treisman & Gelade, 1980).

Özellik entegrasyon teorisinin ilk aşaması, dikkat öncesi aşamasıdır. Bu aşamada, beynin farklı bölümleri görsel alanda bulunan temel özellikler (renkler, şekil, hareket) hakkında bilgi toplar. Özelliklerin otomatik olarak ayrıldığı fikri mantıksız görünebilir; ancak, insan bu sürecin farkında değildir çünkü bu süreç, insan nesnenin bilincine varmadan önce algısal işlemenin başlangıcında gerçekleşir.

Özellik entegrasyon teorisinin ikinci aşaması odaklanmış dikkat aşamasıdır. Bu aşamada bir nesnenin bireysel özellikleri tüm nesneyi algılamak için birleşir. Bir nesnenin bireysel özelliklerini birleştirmek için dikkat gereklidir ve bu nesnenin seçimi, konumların bir "ana haritası" içinde gerçekleşir. Konumların ana haritası, özelliklerin tespit edildiği konumların tümünü içerir; ana haritadaki her konum, çoklu özellik haritalarına erişebilir (Deco & Rolls, 2005). Dikkat, harita üzerinde belirli bir yere odaklandığında, o konumda bulunan özellikler "nesne dosyalarına" kaydedilir ve depolanır. Nesne tanıdık ise, nesne ile önceki bilgi arasında bir ilişki kurulur, bu da bu nesnenin tanımlanmasıyla sonuçlanır (Geng & Behrmann, 2003). Bu aşamayı desteklemek için araştırmacılar genellikle Balint sendromundan şikayetçi olan hastalara baş vururlar. Parietal lobdaki hasar nedeniyle, bu insanlar bireysel nesneler üzerinde odaklanamazlar. Özellikleri birleştirmeyi gerektiren bir uyaran göz önüne alındığında, Balint sendromundan muzdarip insanlar, özellikleri birleştirmek için yeterince uzun süre odaklanamazlar (Robertson, Treisman, Friedman-Hill, & Grabowecky, 1997).

2.3.5 Görsel maskeleme

Geriye doğru maskeleme kavramı ilkin yüksek seslerden hemen önceki anlarda meydana gelen sessiz seslerin maskelenmesi olarak gözlemlenmiştir. 19. yüzyılın sonunda ve 20. yüzyılın başında yapılan deneyler sonucunda maskelemenin görsel algıda da mevcut olduğu ortaya çıkmıştır. Geriye doğru maskeleme, bir uyaran görüntünün sonrasında sunulan bir maske görüntünün, uyaran görüntünün algılanışını bastırması durumudur. Maskeleme, uyaran görüntü maske görüntüden daha önce sunulduğunda meydana geliyor ise ileriye doğru maskeleme, görüntüler aynı anda sunulduğunda meydana geliyor ise eşzamanlı maskeleme adını almaktadır. Bu fenomenin en belirginleştirildiği deneylerin birinde gözlemciye uyarıcı olarak önce gri bir daire, 50 milisaniye sonrasında da uyarıcının etrafını kaplayan kalın bir siyah çember gösterilir. Uyaran tek başına, maskeden sonra, maske ile aynı anda ya da 50 milisaniyeye görece uzun bir zaman önce gösterildiğinde algılanabilir iken, bahsi geçen durumda bastırılmaktadır (B. G. Breitmeyer & Ogmen, 2007; B. Breitmeyer & Ogmen, 2006).

Görsel maskeleme, uyaran maske çiftinin şekillerine göre geriye doğru, ileriye doğru, ya da eşzamanlı maskeleme fenomenini gösterebilir. Maske, uyaranın şekline sıkı sıkıya yaklaşan bir iç kontöre sahip ama uyaran ile üst üste gelmeyen bir şekle sahip ise sadece geriye doğru maskeleme yapabilirken, uyaranın üstünü de örten desenler hem ileri hem geriye doğru maskeleme yapabilmektedir (Enns & Di Lollo, 2000).

Örüntülerin şekillerindeki ve sunuşlarındaki zamanlama gibi farklılıkların maskelemeye etkisi, görsel devre ile ilgili teorilerin test edilmesini mümkün kılan bir ortam hazırlamıştır (B. Breitmeyer & Ogmen, 2006) .

Psikolojik ve üst seviye nörolojik deneyler sonucu gözlemlenmiş bu olguların, nöron düzeyinde (birimsel) yapılan benzetimler ile gösterilebilmesi, bu olguların altında yatan mekanizmaları, ve hepsinin ilintiye sahip olduğu görme algısının mekanizmasını meydana çıkarılmasında katkıda bulunacaktır. Gözlemlenebilecek üst seviye fenomenleri tanıttıktan sonra modelimizi açıklayıp, bu model ile benzetimimizde bu fenomenleri gözlemeye çalışacağız.



3. MODELLEME VE DOĞRULAMA

Bu çalışmada, tek bir sinir hücrelerini, farklı iyon kanallarının doğrusal olmayan geçirgenliğini, bu kanallar vasıtası ile oluşan akımları ve dolayısı ile oluşan membran potansiyelini hesaplayan bir model ile betimleyeceğiz. Nöronlardaki sinapslar ise, vurular sonucunda açılan iyon kanalları ile modellenmiştir. Bu nöron modelleri, birbirlerine yüksek bağlantılığa sahiptir. Hatta bu çalışmadaki kullanılan nöron grubunda 1000 adet nöron bulunmakta, bu nöronların her biri diğerine bağlı olduğundan, grup 1000000'e yakın bağlantıya sahiptir. Tüm bu bağlantılarda bahsi geçen iyon kanalları ayrık ve sürekli kısımlarının olduğu bir modele sahip olduğundan, ve sayılarından dolayı bu bağlantıların idaresi zorlaştığından, ve bu çalışmanın devamının getirilmesinin kolaylaşması için BRIAN2 kütüphanesi kullanılmıştır. Burada her bir sinapsı bağlantının hangi nöronlar arasında olduğuna göre bir bağlantı ağırlığı ile karakterize edip, sonucunda 1000 nüfuslu bir nöron grubunun toplu davranışını inceleyeceğiz.

3.1 Brunel Wang Modeli

Çalışmada benzetimini yaptığımız nöronların fizyolojik çalışma prensibini kullanılan modelin denklemleri ile açıklamanın yapıcı olacağını düşünüyorum. Öncelikle kolaylık olması için herhangi bir açıklama yapmadan modeli oluşturan denklemlerin tümünü yazalım (Brunel & Wang, 2001; Deco & Rolls, 2004).

$$C_m \frac{dV(t)}{dt} = -g_m (V(t) - V_L) - I_{syn}(t) \quad (3.1a)$$

$$I_{syn} = I_{AMPA_{rec}} + I_{AMPA_{ext}} + I_{NMDA} + I_{GABA} \quad (3.1b)$$

$$I_{AMPA_{ext}} = -g_{AMPA_{ext}} \cdot (V(t) - V_E) \sum_{j \in C_{ext}} s_j^{AMPA_{ext}}(t) \quad (3.1c)$$

$$I_{AMPA_{rec}} = -g_{AMPA_{rec}} \cdot (V(t) - V_E) \sum_{j \in C_E} w_j \cdot s_j^{AMPA_{rec}}(t) \quad (3.1d)$$

$$I_{NMDA} = -\frac{g_{NMDA} \cdot (V(t) - V_E)}{1 + \gamma \exp(-\beta V(t))} \sum_{j \in C_E} w_j \cdot s_j^{NMDA}(t) \quad (3.1e)$$

$$I_{GABA} = -g_{GABA} \cdot (V(t) - V_I) \sum_{j \in C_I} s_j^{GABA}(t) \quad (3.1f)$$

$$\frac{ds_j^{AMPA_{ext,rec}}(t)}{dt} = -\frac{s_j^{AMPA_{ext,rec}}(t)}{\tau_{AMPA}} + \sum_{k \in C_E} \delta(t - t_k) \quad (3.1g)$$

$$\frac{ds_j^{GABA}(t)}{dt} = -\frac{s_j^{GABA}(t)}{\tau_{NMDA,decay}} + \sum_{k \in C_E} \delta(t - t_k) \quad (3.1h)$$

$$\frac{ds_j^{NMDA}(t)}{dt} = -\frac{s_j^{NMDA}(t)}{\tau_{NMDA,decay}} + \alpha x_j(t) (1 - s_j^{NMDA}(t)) \quad (3.1i)$$

$$\frac{dx_j(t)}{dt} = -\frac{x_j(t)}{\tau_{NMDA,rise}} + \sum_{k \in C_E} \delta(t - t_k) \quad (3.1j)$$

Nöron dinamiğinin temelini iyon yüklerinin bir zar üzerindeki kontrollü transferine bağlı olması, zarın kayıplı bir kapasite ve bir akım kaynağı olarak modellenenebilmesini sağlar. Modelde, denklem 3.1 deki a-f aralığı elektriksel modeli, g-j aralığındaki denklemler ise sinaptik modeli oluşturmaktadır. Şekil 3.3’de görünen elektriksel model her bir nöron için, sinaptik model ise her bir sinaps için, sadece sinapsın tipine göre sadece ilgili denklemler kullanılarak oluşturulur.

$$C_m \frac{dV(t)}{dt} = -g_m (V(t) - V_L) - I_{syn}(t) \quad (3.1a \text{ yeniden})$$

Burada, zarın iki tarafındaki iyon yoğunluğu farkının sonucu olan zar potansiyeli $V(t)$ ile, bu iyon yoğunluklarının iyon kanallarının geçirgenliğiden doğan akım I_{syn} ile, iyonların iki tarafına dizilerek potansiyel farka neden oldukları membranın kapasitesi C_m , uyarının olmadığı zaman membran potansiyelinin sabit $V_L = -70mV$ a düşüren kaybın iletkenliği ise g_m ile gösterilmektedir.

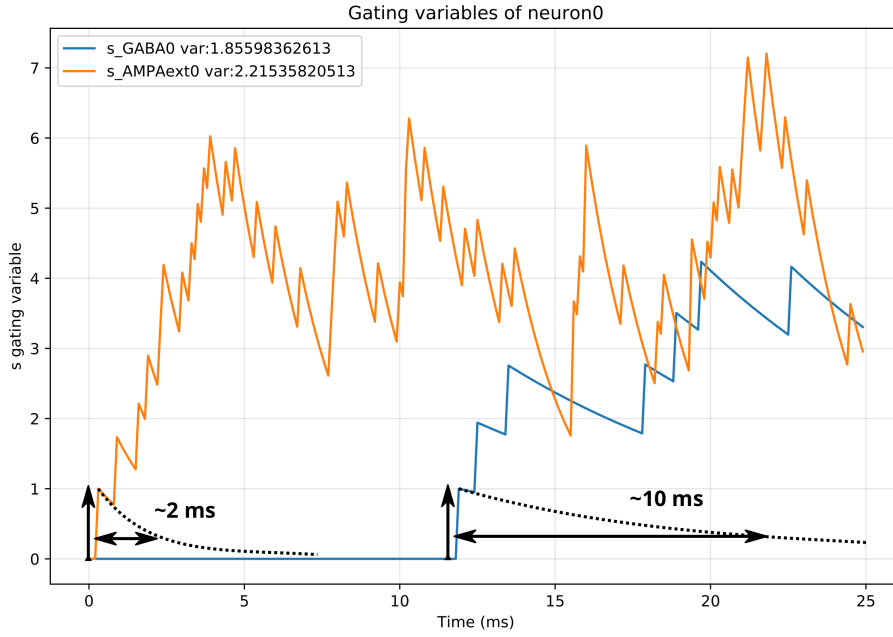
Burada, sinaptik akım, sinapsların üzerindeki çeşitli iyon kanallarından, AMPA, NMDA ve GABA reseptörlerinin ne kadarının açık olduğunu ifade eden s_j geçit katsayıları (gating variables) üzerinden hesaplanır. Burada, $I_{syn} = I_{AMPA_{ext}} + I_{AMPA_{rec}} + I_{NMDA} + I_{GABA}$ akımını oluşturan parçalar, sırası ile dış uyartının modellendiği AMPA reseptörlerinin, nöronlar arası bağlantıların AMPA reseptörlerinin, NMDA reseptörlerinin ve GABA reseptörlerinin neden olduğu akımlardır. AMPA ve NMDA reseptörleri membranın depolarizasyonunu ve sonucu olarak ateşlenmeyi, GABA ise hiperpolarizasyonunu ve sonuç olarak sönümlenmeyi sağlayacak şekilde çalışmaktadır.

Bu akımın AMPA ve GABA bileşenlerini sağlayan AMPA ve GABA reseptörlerinin geçit katsayıları, sinaps öncesi uyartı sonucunda Şekil 3.1 deki gibi bir birim yükselip sönümlenme zaman katsayıları olan $\tau_{AMPA} \approx 2ms$ ve $\tau_{GABA} \approx 10ms$ ile eksponansiyel azalırlar. Modelde AMPA reseptörlerinin dinamiğe etkisi *external* için *ext* ve *recurrent* için *rec* alt indisleri ile ayrılmıştır. Bu iki katkıdan ilki olan $AMPA_{ext}$, dış etkenlerin nöronlar üzerindeki etkisinin modellenmesi için kullanılmaktadır. Nöronlar, çalışılan bölgenin dışındaki diğer nöronların anlık aktivitesinin yarattığı arka plan gürültüsüne, ve dış dünyadan gelen uyarının bilgisine bu katkı ile maruz kalmaktadır. $AMPA_{rec}$ ise, AMPA reseptörlerinin nöronun öz dinamiğine olan katkısını temsil etmektedir. (Destexhe, Mainen, & Sejnowski, 1998)

$$\frac{ds_j^{AMPA_{ext,rec}}(t)}{dt} = -\frac{s_j^{AMPA_{ext,rec}}(t)}{\tau_{AMPA}} + \sum_{k \in C_E} \delta(t - t_k) \quad (3.1g \text{ yeniden})$$

$$\frac{ds_j^{GABA}(t)}{dt} = -\frac{s_j^{GABA}(t)}{\tau_{NMDA,decay}} + \sum_{k \in C_E} \delta(t - t_k) \quad (3.1f \text{ yeniden})$$

Reseptörlerin dinamiğini açıklayabilmek için yapılmış benzetimlerden örnek sonuçlar Şekil 3.2 ve Şekil 3.1’de sunulmuştur. NMDA reseptörlerinin aktivasyon zaman Şekil 3.2 de görüldüğü gibi yoksanamayacak kadar uzun olduğundan dinamikleri iki denklem ile modellenmiştir. NMDA reseptörünün açılması için iki agonist molekülün bağlanmak zorunda olması, kapanmasının ise reseptörden glutamatın yavaş ayrılması olduğu düşünülmektedir (Koch & Segev, 1998, 1. Bölüm). Denklemde birisi yükselme $\tau_{NMDA,rise}$ diğeri düşme $\tau_{NMDA,decay}$ için olacak şekilde iki tane zaman sabiti vardır. Bu



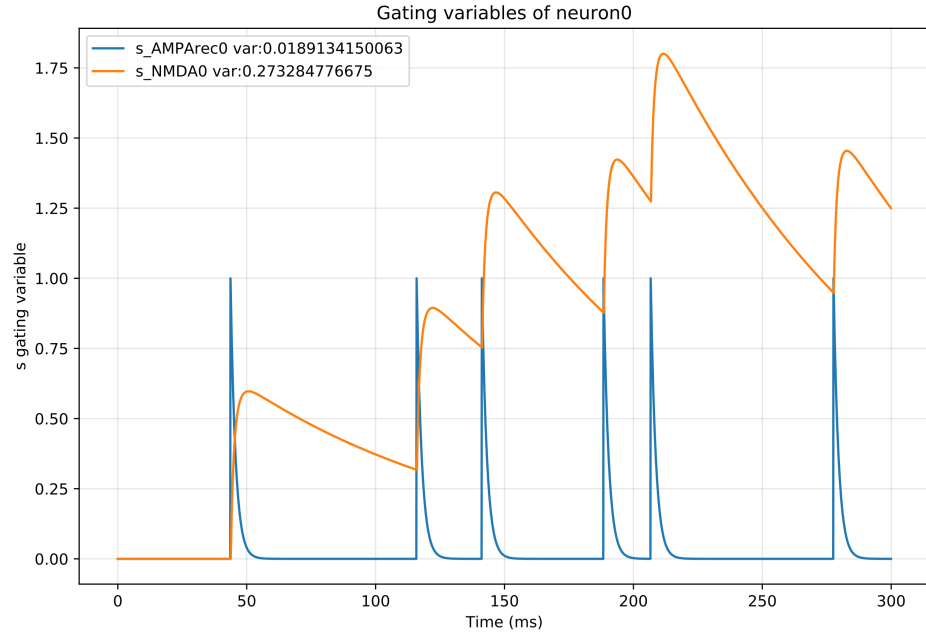
Şekil 3.1 : AMPA ve GABA reseptörlerinin geçit katsayılarının zamana göre değişimi

katsayılar deneysel olarak 32 deg C de $\tau_{NMDA,rise} = 20ms$ ve $\tau_{NMDA,decay} = 25 - 125ms$ olarak ölçülmüştür (Hestrin, Sah, & Nicoll, 1990).

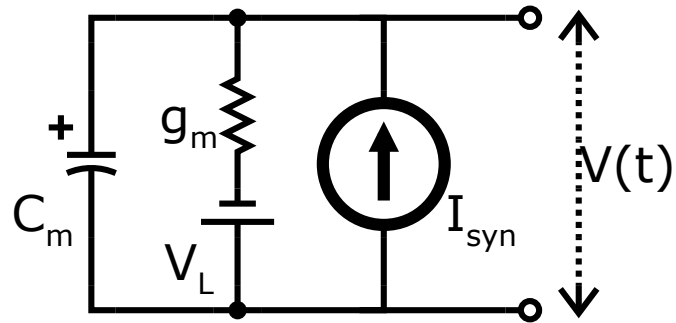
$$\frac{ds_j^{NMDA}(t)}{dt} = -\frac{s_j^{NMDA}(t)}{\tau_{NMDA,decay}} + \alpha x_j(t) (1 - s_j^{NMDA}(t)) \quad (3.1i \text{ yeniden})$$

$$\frac{dx_j(t)}{dt} = -\frac{x_j(t)}{\tau_{NMDA,rise}} + \sum_{k \in C_E} \delta(t - t_k) \quad (3.1j \text{ yeniden})$$

Presinaptik nöronun atım olayının post sinaptik nörondaki kanalları açması ve depolarizasyona neden olması sonucunda membranda $I_{syn} = I_{AMPA_{rec}} + I_{AMPA_{ext}} + I_{NMDA} + I_{GABA}$ akımı oluşur. Bu akım, Denklem 3.1a ile ifade edilen, Şekil 3.3’de şematize edilmiş elektriksel modeli süren akımdır. Nöronun tüm akım bileşenleri öncelikle ilgili sinapsların açık olma oranlarını sinapsın ağırlığı ile çarpıp nöronun tüm sinapsları üzerinden toplar, ve her bir reseptör için bir geçirgenlik hesaplanmış olur. (C_I baskılayıcı bağlantıların, C_E uyarıcı bağlantıların, C_{ext} ise uyarıcı dış bağlantıların kümesidir.) AMPA ve GABA reseptörleri için potansiyele göre doğrusal bir davranış kabul edilmektedir, dolayısı ile deneysel çalışmalardan elde edilen, Tablo 3.1 deki g_x iletkenlikleri kullanılır. NMDA reseptörlerinin davranışı kapıyı kapatabilen bir Magnezyum iyonu (Mg^{2+}) ile baskılanmaktadır, dolayısı ile NMDA geçirgenliği



Şekil 3.2 : AMPA ve NMDA reseptörlerinin geçit katsayılarının zamana göre değişimi



Şekil 3.3 : Nöronun elektriksel modeli

membran gerilimine göre nonlinear bir Fermi fonksiyonu $1 + \gamma \exp(-\beta V(t))$ ile çarpılır.

$$I_{syn} = I_{AMPA_{rec}} + I_{AMPA_{ext}} + I_{NMDA} + I_{GABA} \quad (3.1b \text{ yeniden})$$

$$I_{AMPA_{ext}} = -g_{AMPA_{ext}} \cdot (V(t) - V_E) \sum_{j \in C_{ext}} s_j^{AMPA_{ext}}(t) \quad (3.1c \text{ yeniden})$$

$$I_{AMPA_{rec}} = -g_{AMPA_{rec}} \cdot (V(t) - V_E) \sum_{j \in C_E} w_j \cdot s_j^{AMPA_{rec}}(t) \quad (3.1d \text{ yeniden})$$

$$I_{NMDA} = -\frac{g_{NMDA} \cdot (V(t) - V_E)}{1 + \gamma \exp(-\beta V(t))} \sum_{j \in C_E} w_j \cdot s_j^{NMDA}(t) \quad (3.1e \text{ yeniden})$$

$$I_{GABA} = -g_{GABA} \cdot (V(t) - V_I) \sum_{j \in C_I} s_j^{GABA}(t) \quad (3.1f \text{ yeniden})$$

3.1.1 Modelin ayrıklaştırılması

Benzetimde modelin tanımlanışı birçok farklı şekilde yapılabileceğinden, bu çalışmada kullanılan kütüphane Brian2 ile modelin nasıl tanımlandığının açıklanması yararlı olacaktır. Binler mertebesinde nöronun topyekün davranışının modelleniyor olması bu tanımlarla alakalı kuralları, tanımların nasıl yorumlandığını önemli kılmaktadır.

Şekil 3.4 de Brian2 de bir dinamiği tanımlamak için gereken tanımlar, yani nöron modelleri ve sinaps modelleri, birlikte gösterilmiştir. Nöron modelinde, her bir nöronun değişkenleri ve dinamiği ifade edilir. Şekil 3.4 de kırmızı kutu içerisine alınmış değişkenler, bu değişkenlerin dinamiği, ve nöronun ateşlenmesinin koşulu bu modelde ifade edilir. Binlerce nöron için bu modelin tanımlanmasının yanı sıra, binlerce nöronun milyonlara yaklaşabilecek sayıdaki her bir sinapsı için ise sinaps modeli tanımlanmalıdır. Bu modelde ise, presinaptik nöronun ateşlenmesinin sonucunda vuku bulan değişken değişimleri, ve her bir sinaps için ayrı tutulan değişkenlerin zaman içerisindeki dinamiği ifade edilir.

Şekil 3.4 üzerinden bir senaryo ele alalım. Presinaptik nöron uyarıcı ve sistem içinden bir nöron olsun. İncelediğimiz nöron sistemin içinden olduğundan sadece NMDA ve AMPArec reseptörleri dinamiğe katkıda bulunacaktır. Bu nöron ateşlendiğinde, sinaps modelinde ifade edildiği gibi, $s_{AMPA_{rec}}$ her bir sinapsa özgü w_i ağırlığı kadar artırılmaktadır. Bu, postsinaptik nörondaki AMPA reseptörlerinin açık olan oranının aniden artmasını modeller. Aniden artan $s_{AMPA_{rec}}$, zamanla $\delta_t s_{AMPA_{postrec}}$ ile ifade edilen diferansiyel denklem sonucunda Şekil 3.1 deki gibi azalacaktır. Geçici olarak

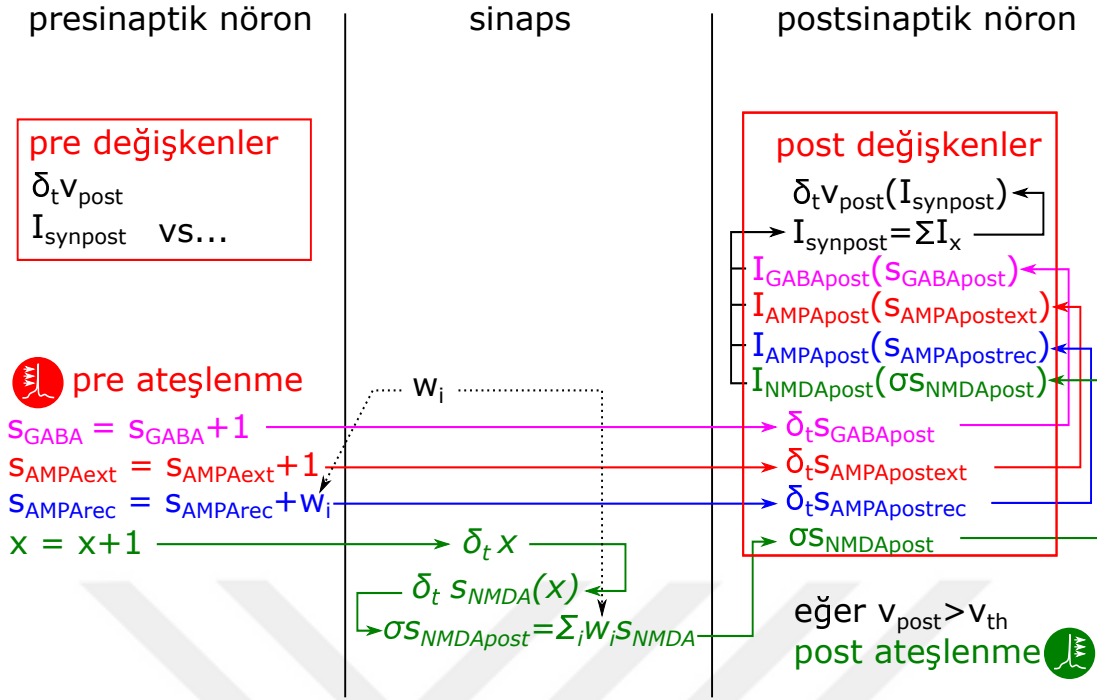
yükselmiş $s_{AMPA_{post_{rec}}}$, $I_{AMPA_{post_{rec}}}$ 'in artmasına neden olacak, bu da toplanan $I_{synpost}$ 'a katkıda bulunacaktır. $I_{synpost}$, ise membranı depolarize edecek, ve membran potansiyeli $|v_{post}|$ 'u düşürecektir. Dolayısı ile tüm bu süreç, postsinaptik nöronu ateşlenmeye yaklaştıracak, yani uyarmış olacaktır. Aynı şekilde x değişkeni de 1 artarak $\delta_t x$ ve $\delta_t s_{NMDA_{post}}$ dinamik sisteminin vasıtası ile NMDA reseptörlerinin açık olma oranı $s_{NMDA_{post}}$ 'un Şekil 3.1 deki gibi artmasına neden olacaktır. Post sinaptik nöronun tüm sinapslarındaki $s_{NMDA_{post}}$ değişkenleri toplanıp $\sigma_{s_{NMDA_{post}}}$ 'un artmasına, dolayısı ile $I_{NMDA_{post}}$ 'un artmasına, bu da toplanan $I_{synpost}$ 'a katkıda bulunacak, depolarizasyona ve $|v_{post}|$ 'un azalmasına neden olacaktır. Yani bir uyarıcı ateşlemenin iki farklı reseptör vasıtası ile etkisi yakalanabilmektedir.

Tersine, presinaptik nöronun baskılayıcı bir internöron olması durumunda ise GABA reseptörleri dinamik davranışı etkileyecektir. Presinaptik nöron ateşlendiğinde, bu sefer GABA reseptörlerinin açık olma oranı s_{GABA} bir artacak, benzer şekilde $\delta_t s_{GABA}$ diferansiyel denkleminin bu etkiye cevabı Şekil 3.1'deki gibi olacaktır. 2 milisaniyeliliğine artmış olan s_{GABA} , $I_{GABA_{post}}$ akımının artmasına neden olacaktır. $I_{GABA_{post}}$ hiperpolarizasyona neden olacak şekilde aktığından $|v_{post}|$ 'un artmasına neden olup post sinaptik nöronun baskılanmasına neden olacaktır.

3.1.2 Doğrulama benzetimi

Öncelikle, kendi benzetim ortamımızda deney parametrelerini de aldığımız, Brunel ve Wang (2001)'daki sonuçların yinelendiğini göstermemiz gerekmektedir. Makale, kortekste çalışan bellek yapısında etkin rol oynayan sürdürülen aktivitenin (persistent activity) tekrarlanan baskılayıcı ağlar ile (recurrent inhibition) oluştuğunu, ve bu aktivitenin bu yapıda bellek niteliğini açıklayabildiğini göstermektedir. Burada Brunel ve Wang (2001) makalesini, ele aldığımız Deco ve Rolls (2003) makalesinde kullanılan nöron modelinin kaynağı olduğundan, dolayısı ile modelin görsel yolda odaklanma ile ilgili de kullanılabileceğinin onaylanmış olmasından seçtik. Bu bölümde, kurduğumuz modeli ve model için yazdığımız kodu test edeceğiz.

Öncelikle, denklemlerini Alt Bölüm 3.1'de ifade etmiş olduğumuz modelin kısıtlarından ve benzetimde nasıl kullanıldığından bahsetmeliyiz. Sinir hücresini temel hesaplama ünitesi olarak alan çalışmalardaki temel zorluklardan biri, hücrenin konu-



Şekil 3.4 : Kullanılan nöron modelinin Brian2 deki sinaplara tanımlanışı. δ_t yanındaki değişkenin bir diferansiyel denklem ile tanımlandığını, $\sum_i$ tüm sinapslar üzerinden alınan toplamı, post ve pre son ekleri değişkenlerin hangi nöronun (postsinaptik ya da presinaptik) değişkenleri olduğunu ifade eder. Pre ateşlenme yazısının altında listelenmiş pre sinaptik ateşleme sonucunda yapılan değişken atamalarından çıkan oklar, bu atamaların dinamiğe etkisinin kanallarını göstermektedir.

muna ve işlevine ¹ göre yüksek derecede farklılaşabilmesidir. Bundan, benzetimler incelenecek olan fenomenin açıklayıcılığına katkıda bulunmayacak farklılaşmalardan arındırılır. Bunlardan birkaçından bahsetmek benzetimin kapsamını belirlemek için yardımcı olacaktır:

- Nöronların fizyolojik parametreleri eşit alınmıştır. Yani büyüklük, şekil ya da uzunluk göz ardı edildiğinden zar kapasitesi, zar kaybı, sinyal gecikmeleri gibi parametreler sabittir.
- Sadece GABA, NMDA ve AMPA reseptörleri ve ilgili antagonistleri modellenmiştir. Çalışılan bölge korteks olduğundan glisin ya da benzerleri ele alınmamıştır.
- Kullanılan kaynakların halen nöronlar arası bağlantılarda rastlantısallığı ele almamalarından yola çıkarak, tüm nöronlar birbirine bağlanmıştır.

¹ya da incelerken işlevi tayin ettiğimiz davranışına

Bu çalışmadaki de dahil olmak üzere, biriktir ve ateşle nöron modellerinin kullanıldığı benzetimler bini aşkın nöron ile yapılmaktadır. Aynı zamanda bu benzetimlerde baskılayıcı ve uyarıcı olmak üzere iki nöron tipi de kullanılmaktadır. Bu iki nöron tipi genelde Gray I tipi ve Gray II tipi adı verilen fizyolojik farklara sahip sinapslara sahip olduğundan modelimizde farklı katsayılara sahip olacaktır.

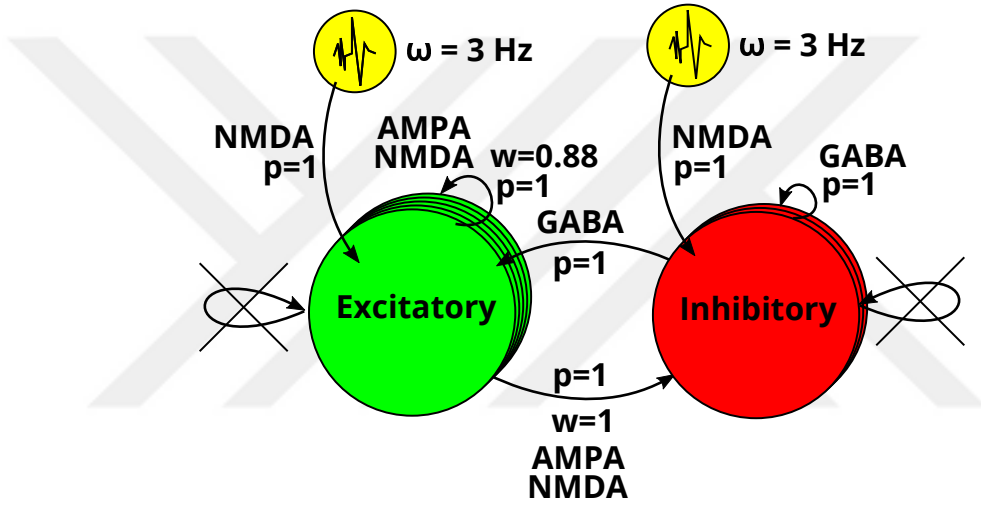
Bu nöronlar, uyarıldıklarında farklı antagonistler (tetikleyiciler) serbest bırakmakta, ve aynı zamanda farklı antagonistlerin kendilerinde yarattığı etki (akıma katkısı) da farklı olmaktadır. Tablo 3.1 de, NMDA , AMPA ve GABA antagonistlerinin akıma katkısının modellenmesinde kullanılan katsayıların iki nöron tipi için farklı olduğu görülebilir. (sinaps modelleri farklıdır) Bunun yanı sıra, hücre zarı kapasitesi C_m , hücre zarı kaybı g_m ve uyartının sonucunda oluşan gecikmenin modellendiği tepkisizlik zamanı τ_{rp} gibi nöronda harekete geçirilen akımın zar gerilimini nasıl değiştireceğinin modelinin katsayıları da hücre tipine göre fark gösterdiği tablodan görülebilir. Tablodaki değerler referans Brunel ve Wang (2001) dan alınmıştır.

Çizelge 3.1 : Benzetimde kullanılan katsayılar

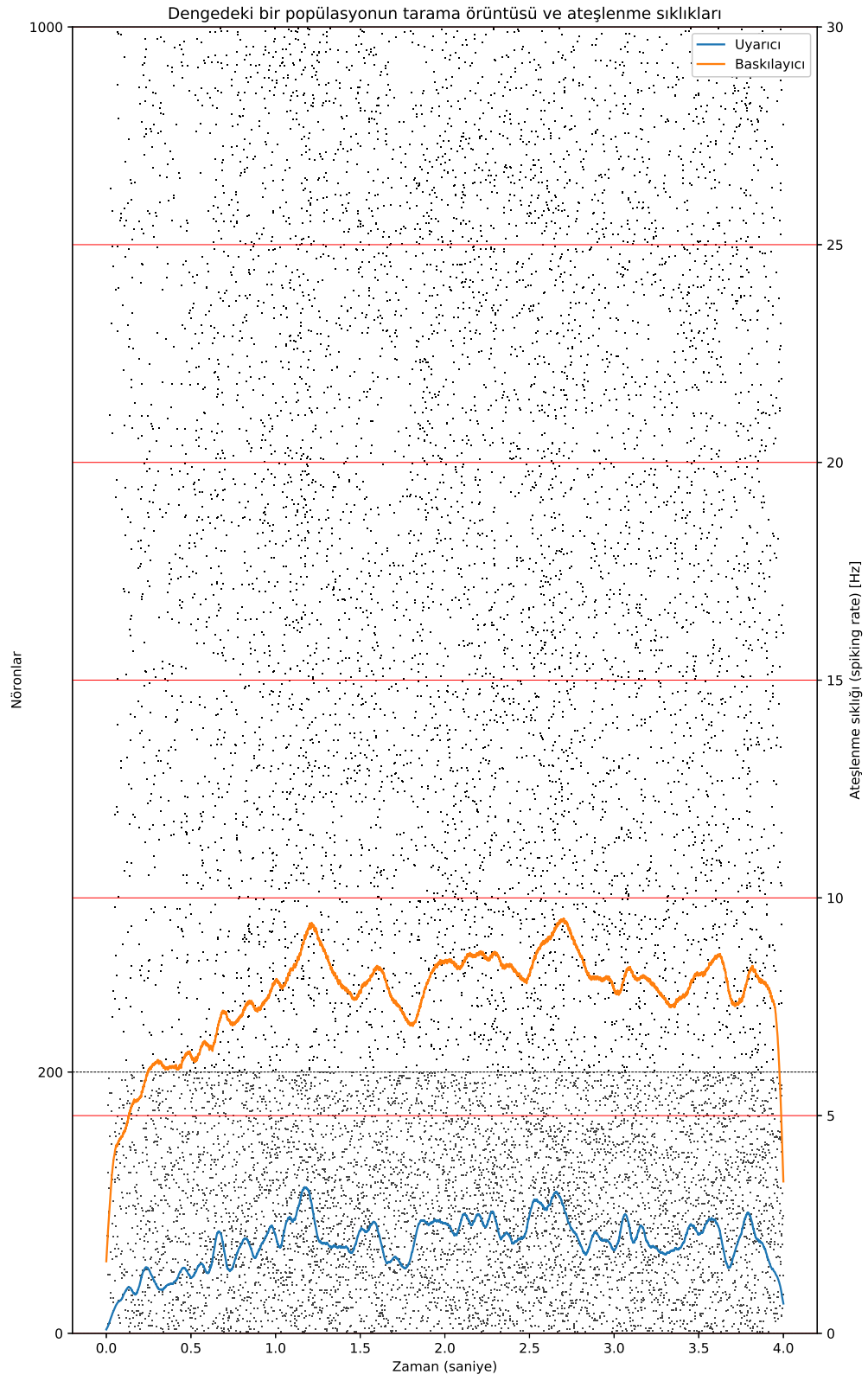
	Uyarıcı	Baskılayıcı	Açıklama
N	800	200	Nöron sayısı
C_m [nF]	0,5	0,2	Zar Kapasitansı
g_m [nS]	25	20	Zar kaybı
τ_{rp} [ms]	2	1	Tepkisizlik zamanı
AMPA geçit değişkeni			
$g_{AMPA_{ext}}$ [nS]	2,08	1,62	dış bağlantılar için kaybı
$g_{AMPA_{rec}}$ [nS]	0,104	0,082	iç bağlantılar için kaybı
τ_{rp} [ms]		2	düşüş zaman sabiti
NMDA geçit değişkeni			
g_{NMDA} [nS]	0.327	0.258	tüm bağlantılar için kaybı
α [1/ms]	0.327	0.258	
$\tau_{NMDA_{rise}}$ [ms]		2	yükselme zaman sabiti
$\tau_{NMDA_{decay}}$ [ms]		100	düşüş zaman sabiti
GABA geçit değişkeni			
g_{GABA} [nS]	1,25	0,973	tüm bağlantılar için kaybı
τ_{GABA} [ms]		10	düşüş zaman sabiti

Modeli doğrulamak için kullandığımız deneyde, toplam 1000 tane nöron tanımlanmıştır. Kaynaktaki gibi, her bir nöron, kendisi haricindeki tüm diğer nöronlara bağlıdır. Bu nöronların 200 tanesi baskılayıcı internöronlar, diğer 800 tanesi ise uyarıcı nöronlardır. Makaledeki gibi 800 uyarıcı nöronun 80'erli gruplar halindeki özelleşmiş 5 grubu, daha önceden belirli uyaranlara kodlandığı varsayılmış, dolayısı ile grup içi bağlantıları güçlü $w_i = w^+ = 2,1$ ama gruplar arası bağlantıları zayıf $w_i = w^- = 0,88$ kılınmıştır. Bir koda özelleşmemiş, uyaran ayırt etmeyen kalan uyarıcı 400 nöron, kodlara özelleşmiş bu beş alt popülasyona olan bağlantıları zayıf, ters yöndeki bağlantıların ağırlığı birim $w_i = 1$ alınmıştır.

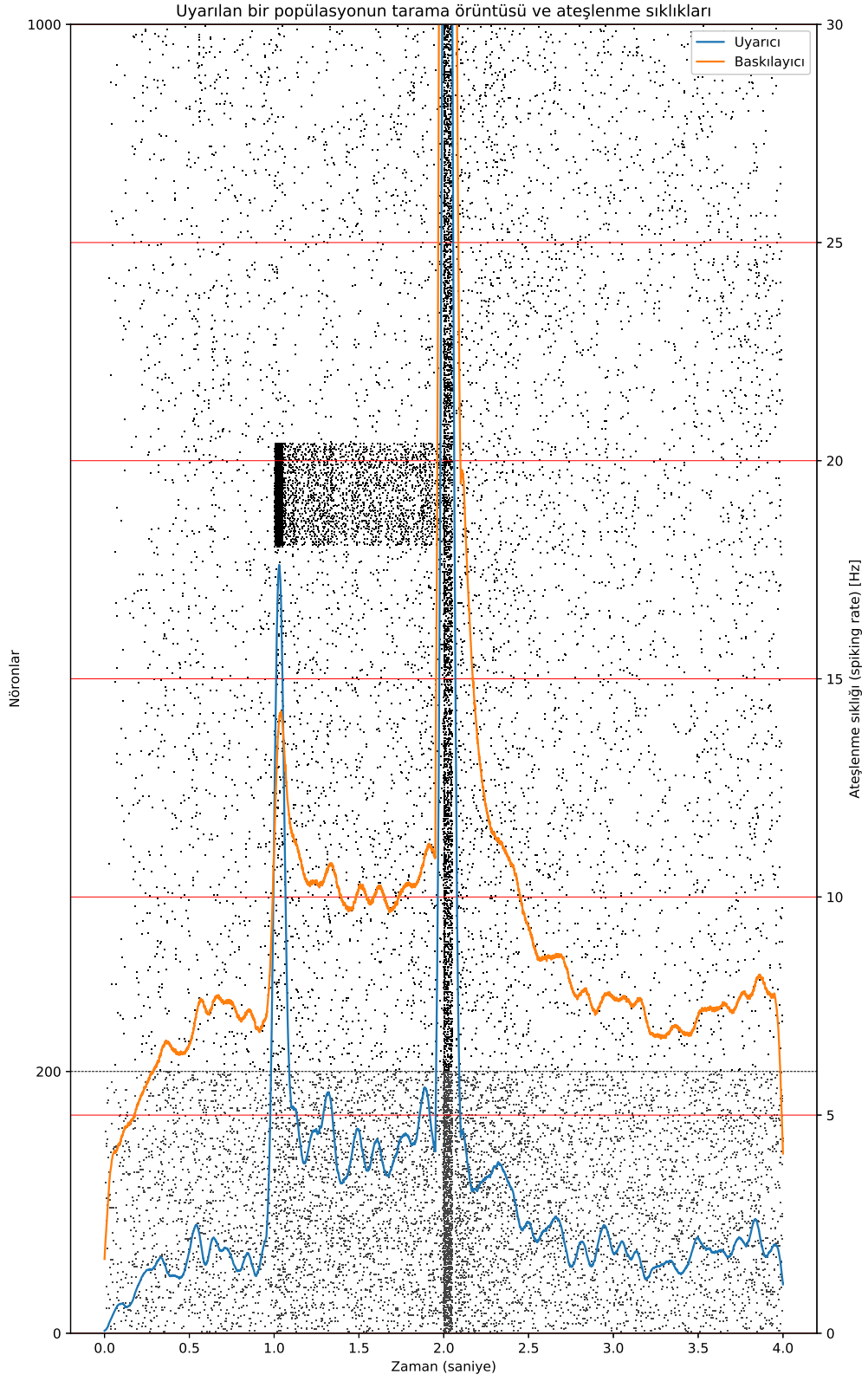
Baskılayıcı 200 internöron, uyarıcı nöronların sürekli birbirlerini uyardıkları bir pozitif geri besleme durumunda kalmamalarını sağlamaktadır. Dolayısı ile bu nöronlar, popülasyondaki tüm nöronlara bağlı olup, tüm popülasyonun uyarıl birbirlerine yaptıkları, her bir 800 uyarıcı nörona yaptıkları, ve her bir uyarıcı nöronun onlara yaptığı bağlantı ağırlıkları da birimdir. Dışarıdan modelin koşabilmesi için gereken aktivasyon enerjisini veren gürültüyü modellediğimiz 800 nöron ise tüm nöronlara bire bir, birim ağırlık ile bağlıdır. Benzetimdeki her bir nörona Poisson dağılımına göre ortalama $3Hz$ sıklıkta ateşlenen bir nöron grubu bağlıdır. (Aslında bu tip bir bağlantının değişkenler üzerindeki etkisi kullanılmaktadır.) Bağlantıların bir diğer gösterimi de Brunel ve Wang (2001, Figure 1) de görünebilir.



Şekil 3.5 : Doğrulama deneyinde kullanılan bağlantılandırılmış nöronlar 800 uyarıcı ve 200 baskılayıcı nörona sahip iki havuza ayrılmış, ve bunlar kendi kendileri haricinde bire bir bağlantılandırılmıştır. İncelenen bölgenin dışındaki bölgelerden gelen gürültü, 3Hz eşik frekansındaki Poisson dağılımına sahip bir gürültü ile benzetilmiştir.



Şekil 3.6 : Denge durumundaki nöronların tarama örüntüsü



Şekil 3.7 : Uyarılmış durumundaki nöronların tarama örüntüsü

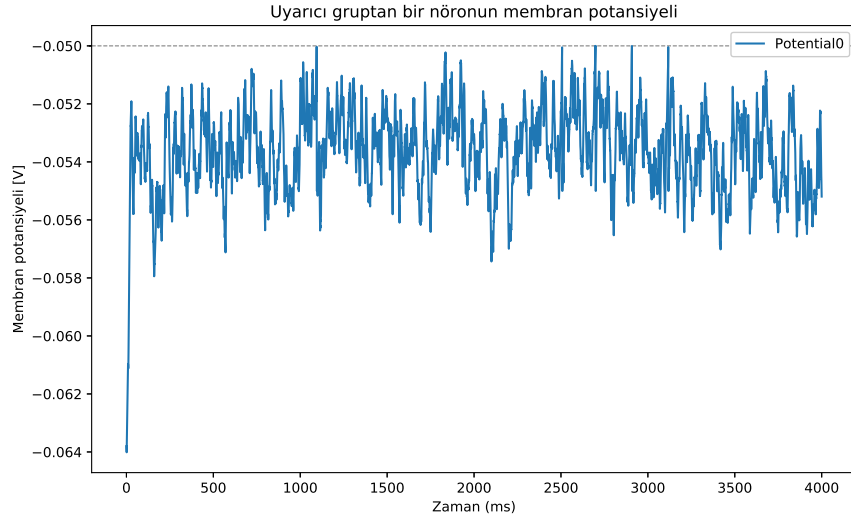
Şekil 3.6 ve Şekil 3.7’de yapılan deneyin sonucunu görebiliriz.

İlk deneyde Şekil 3.6’de sadece arka fon gürültüsüne maruz bırakılmış 1000 nöronun 500 ünün 4 saniye boyunca topyekün davranışı görülebilir. Aşağıda kalan ilk 100 örüntü, baskılayıcı interneuronların, geriye kalan 400 tane örüntü uyarıcı nöronların ateşlenmelerini göstermektedir. Burada baskılayıcı nöronlar yüksek bir sıklıkta ve uyarıcı nöronlar nazaran daha düşük bir sıklıkta, denge durumunun civarında bir spontane ateşlenme göstermektedir. Üstüste bindirilmiş olan grafikte ise tüm popülasyondaki uyarıcı nöronların ve baskılayıcı internöronların üzerinden ortalaması alınan ateşleme frekansları gösterilmektedir. Grafiklerde ortalama, 25 milisaniyelik bir yürüyen pencere ile alınmıştır. Bu ateşlenmenin sıklığı, tekrarladığımız makaledeki ortalama alan teorisi sonucunda ortaya çıkan sıklıklara ($\omega_E = 3Hz$, $\omega_I = 9Hz$) ve serebral korteksteki spontan aktiviteye yakın çıkmıştır. (Kaynak olarak aldığımız (Brunel & Wang, 2001; Deco & Rolls, 2004) da frekans ω ile gösterildiğinden bu gösterim adet edinilmiştir.)

Makalenin kendisinde de simülasyonlarda elde edilen değerlerin hesaplamadan daha düşük olması durumu yapılan basitleştirmelere dayandırılmıştır.

Karşılaştırma için denge durumundaki sistemdeki nöronların değişkenlerinden örnek vermekte de yarar vardır. Şekil 3.8 da benzetimdeki bir nöronun membran potansiyelini görülebilir. Örnek olarak $t = 1,1sn.$ ’de membran potansiyeli $v = -0,050V$ ’a düşen nöron ateşlenme yaşamış, ve membran potansiyeli $v = -0,055V$ ’a düşmüştür.

İkinci deneyde, özelleşmiş gruplardan birisi arka plan gürültüsünden daha yüksek sıklıkta bir gürültü ($25Hz \cdot 80$) ile uyarılsın. Bu uyarı, arka plan gürültüsünün nazaran az olduğu denge durumundaki popülasyonu, birbirine bağlı nöron popülasyonun oluşturduğu toplam dinamik sistemini yinelenen ateşlemelerin olduğu yeni bir denge noktasına taşımaktadır. Bu deneyin sonucunu Şekil 3.7’de görebilirsiniz. Bu tekrarlılığa sahip durum, aslında Brunel ve Wang (2001, Figure 4) deki dallanma diyagramında görülebilecek şekilde, sadece belli bir w_i aralığında gerçekleşmektedir. Makaledeki 3.1 numaralı sonucu tekrarlayabildiğimizden dolayı kurguladığımız ve denklemlerini yazdığımız bu model, başka deneyler için de kullanılabilir.



Şekil 3.8 : Benzetimin denge durumundan alınan bir nöronun membran potansiyeli. Örnek olarak, 1,1 inci saniyede membranın potansiyeli eşik değeri ilk kez geçmiş ve bu durum bir ateşlenmeye neden olmuştur.

Bu çalışmanın devamı niteliğindeki çalışmada, görsel dikkatin modellenmesi ele alınacaktır (Deco & Rolls, 2003, 2005).



4. DENEY

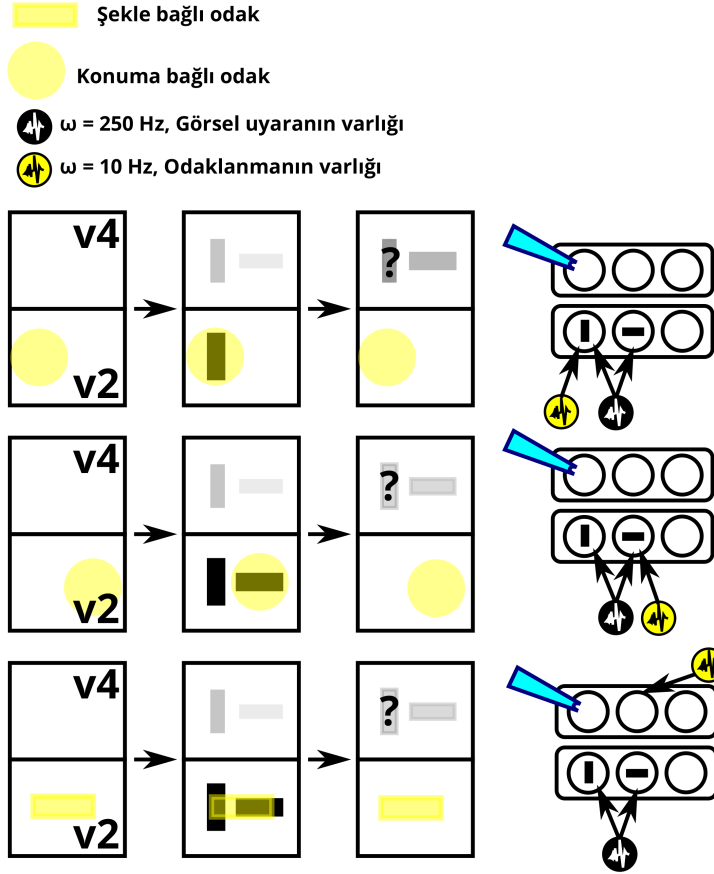
4.1 Görsel Dikkat Deneyi

Bölüm 3’de Brunel ve Wang (2001) makalesindeki topla ve ateşle sinir modelinin benzetimi sınıdıldığından, bu modeli görsel dikkat fenomenini açıklayabilmek için kullanabiliriz. Bir önceki bölümde, uyarıcı ve baskılayıcı popülasyonlara sahip bir sinir grubunda spontane ve sürdürülebilir vuruların (*Ing:spike*), ikili stabil bir yapının gözlemlenebildiğini gördük. Bu, en azından biyolojik olarak bir bellek davranışının sergilenebildiği bir yapı taşıyı elde ettiğimizi göstermektedir. Bunun yanı sıra, yanlış rekabet (*biased competition*) özelliğinin de olup olmadığını test etmemizi sağlayacak bir ortam oluşturmuş oluyoruz.

Bundan sonraki adım, bu model ile görsel devrenin indirgenmiş bir halinin gösterimi olabilir. Nöronların atmalarını teker teker ele almayan, atma sıklıklarını sürekli bir dinamik sistem olarak tanımlayarak bu modellemeyi yapan modeller mevcuttur (Deco & Rolls, 2005). Aynı zamanda bu deneyde, görsel olarak kullanılan imgelerin değişmez gösterimlerini (*invariant representation*) elde edebilmek için çoğu durumda bir Hebbian eğitim kurgulanmaktadır. Bu gösterimleri biyolojik olarak gerçeğe yakın bir şekilde elde edebilmek için ise 2 boyutlu görsellerin Gabor fonksiyonlarına ¹ olan açılımları kullanılmalıdır.

Elimizdeki modelin karmaşıklığı göz önünde tutularak, dikkat problemini Reynolds, Chelazzi, ve Desimone (1999) makalesindeki biyolojik sonuçları yinelemeye çalışan Deco ve Rolls (2005) makalesine dayandıracağız. Burada, görsel sistem V4 ve V2 devrelerine indirgenmiş, ve iki adet seçilen örüntünün V2 deki popülasyona etkisinin V4 popülasyonuna nasıl yansıdığı incelenecektir. Fizyolojik deneylerde, çeldirici örüntülerin karşıtlığının diğer cismin tanımlanmasındaki başarımı etkilediği gözlemlenmektedir.

¹Gabor fonksiyonları ile bir görüntünün öteleme, döndürme ve ölçeklenmeye karşı değişmez gösterimleri elde edilebilmektedir.



Şekil 4.1 : Deco ve Rolls (2005)'de geçen deneylerde, belli bir tercih edilen görsel uyarana özelleştiği bilinen bir V4 katmanı sinirinin, 1. Odak tercih edilen uyarının konumunda iken hem tercih edilen, hem de çeldirici uyarın geldiğindeki davranışı, 2. Odak çeldirici uyarının konumunda iken hem tercih edilen, hem de çeldirici uyarın geldiğindeki davranışı, 3. Odak tercih edilen uyarının şeklinde (V4 teki gösterimine) iken hem tercih edilen, hem de çeldirici uyarın geldiğindeki davranışı incelenmiştir. Bu çalışmada, sadece ilk iki deney yapılacaktır.

Simülasyon, iki farklı odak deneyini mümkün kılmaktadır. Bunlardan ilkinde, örüntüdeki bir konumun temsili uyarıldığında o konumda bulunan cismin temsiliinde fazladan bir uyarılma görülmesi, ikincisinde ise örüntüde bulunan bir cismin şekil temsili uyarıldığında o şeklin örüntüde bulunduğu konumun temsiliinin uyarıldığının görülmesi amaçlanmaktadır. Yardım için Deco ve Lee (2004)'deki Fig. 4'e göz atılabilir. Şekilde, V1 modülünden ikiye ayrılıp biri dorsal "nerede" yolunu diğeri ise ventral "ne" yolunu oluşturan, sırası ile DM ve VM modülleri gösterilmektedir. Burada, Hebbian kural ile eğitilmiş bir ağda, "nerede" modülünde belirli bir konumu üstten aşağı kutuplayarak, o konumdaki cismin "ne" modülündeki gösteriminde sinirsel aktivitenin arttığı gözlemlenebilmektedir. Aynı şekilde "ne" modülünde bir cismin gösterimi üstten aşağı bir süreç ile kutuplayarak, o cismin örüntüde öğrenilmiş konumunu "nerede" modülündeki aktivite artışından bulunabilir. Elimizdeki modelde üç ayrı sinir grubunu betimlemek nazaran karmaşık olacağından bu dinamiği V2 ve V4 bölgeleri ile betimleyen, Deco ve Rolls (2005)'deki deney tekrarlanmıştır.

Burada, Deco ve Rolls (2005)'daki gibi, daha önceden bir öğrenme döneminden geçmiş olduğunu varsaydığımız bir ağ kurgulayacağız. Bir önceki bölümde kullanılan 1000 nörona sahip modelden iki adet kullanıp, makaledeki gibi birini V2 diğeri V4 görsel modüllerini modellemek için kullanacağız. Bu deneyde V2 modülü, 200 adet baskılayıcı, 80i referans cismi (S1), 80i çeldirici cismi (S2) temsilen, 800 tane uyarıcı nörondan oluşmaktadır. S1 ve S2 grubundaki nöronlar, bu cisim görsel alan içerisinde ise $\lambda = 250Hz^2$ karakteristik frekansına sahip bir Poisson gürültüsüne tabii tutularak, eğer dikkat üzerlerinde ise $\lambda = 10Hz$ lik fazladan bir gürültüye tabi tutularak uyarılır. Dikkati temsil eden gürültü, aslında üstten aşağı etkiyi temsil eden anlık aktivitenin artışıdır. Deco ve Rolls (2005) da aslında bu uyarı ventral yoldan V1 modülünü, dolayısı ile dorsal modülü tetikleyebilmektedir. Lakin, burada bir basitleştirmeye gidilmiş ve bu gürültünün kaynağı irdelenmemiştir.

V2 modülündeki S1 ve S2 gruplarındaki nöronlar, daha üstteki V4 te bulunan S1' ve S2' uyarıcı nöron grupları ile Jb, Kb, Jf ve Kf ağırlıkları ile bağlanmaktadır. Bahsi geçen makalede yapılan analiz sonucunda bu parametrelerin yanlı yarışı mümkün kılacak dar bir parametre uzayı çıkarılmıştır. Jf ve Kf ağırlıkları sırası ile ($S1 \rightarrow$

²Poisson dağılımının alışlagelmiş gösteriminde frekans biriminde olan sıklık göstergesi λ ile gösterildiği için aynı gösterim benimsenmiştir.

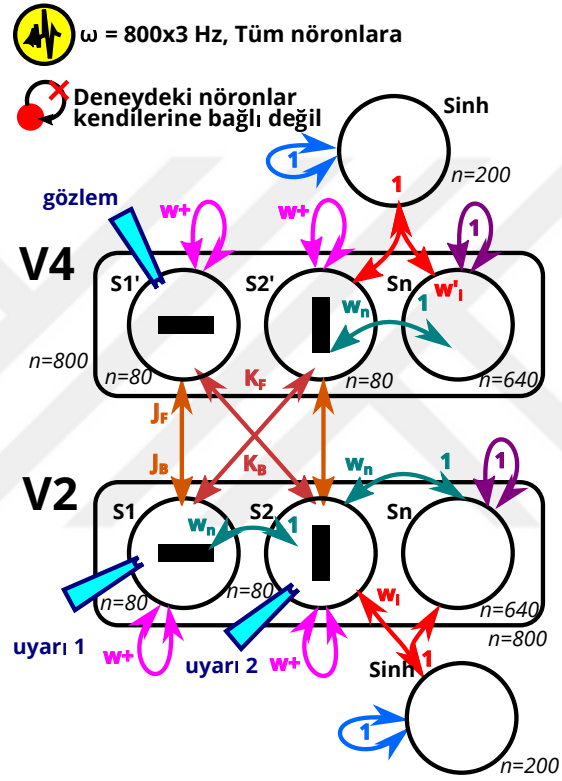
$S1'$), $(S2 \rightarrow S2')$ bağlantılarının ve $(S1' \rightarrow S1)$, $(S2' \rightarrow S2)$ bağlantılarının katsayısıdır. K_b ve K_f ise, aynı yönelimde, ama çapraz bağlantıların ileribesleme ve geribesleme katsayılarıdır. Bu deneyde, $J_f = 1.6$, $J_b = 0.5$, $K_f = 0.15$, $K_b = 0.06$ alınmıştır. Bölüm 3'den farklı olarak, V4 katmanının alıcı alanı daha büyük olduğundan, burada baskılayıcı bağlantılar da daha yüksek bir katsayı ile ağırlıklandırılmaktadır. V2'de baskılayıcı sinirlerin çıkışları $w_I = 1$, V4'te ise $w'_I = 1.25$ ile ağırlıklandırılır. Her bir modülde yinelenen atmaların sürdürülebilmesi şartının sağlanması için aynı zamanda, cisimlerden bağımsız *nonselective* grubundan cisimleri temsil eden gruplara sırası ile V2 ve V4'te,

$$w_n = (-f * J_b - f * K_b) / (1 - 2 * f) + w_- \quad (4.1)$$

ve

$$w'_n = (-f * J_f - f * K_f) / (1 - 2 * f) + w_- \quad (4.2)$$

katsayıları atanmalıdır. Şekil 4.2 de ağırlıkların nasıl alındığı gösterilmiştir. Aynı zamanda, çalışmada referans alınan çalışmadaki gibi Kalsiyum derişiminin sinaptik akıma olan katkısı da modellenmiştir, lakin, sonuçlarda bir değişikliğe neden olmadığından, Brunel ve Wang (2001) deki temel modele geri dönlmüştür.



Şekil 4.2 : Bu deneyde iki adet katman bulunduğundan Bölüm 3 teki popülasyonlardan iki adet kullanılmıştır. Bunun sonucunda, iki katmanı birbirine bağlantılandırılması, ve Bölüm 3 ten farklı olarak ağırlığı değiştirilen bağlantıların hangileri olduğu bu şekilde gösterilmiştir. Çift taraflı okların yanındaki katsayılar okun yönündeki bağlantının katsayısını gösterir.

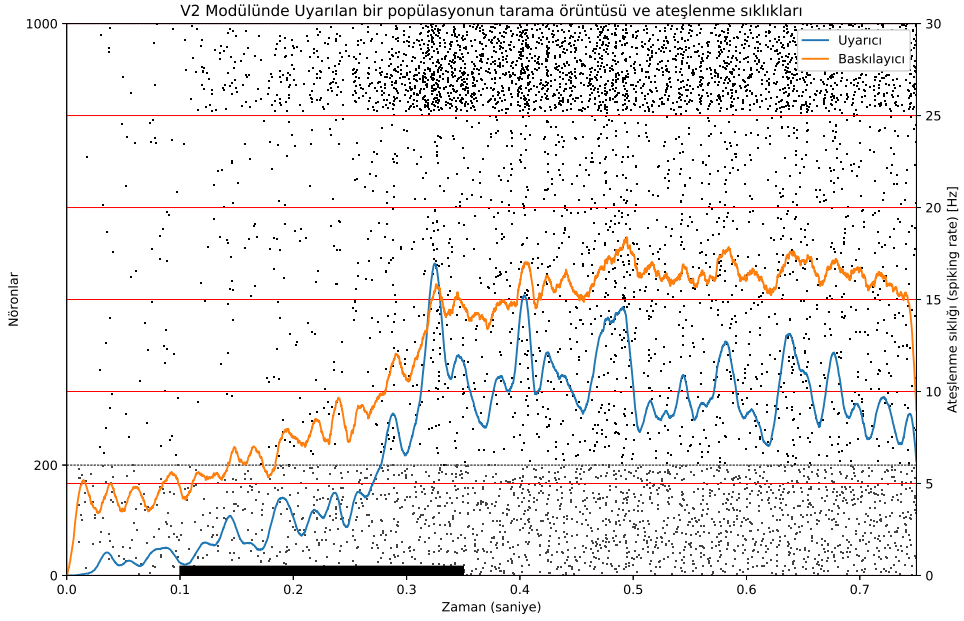


5. SONUÇLAR

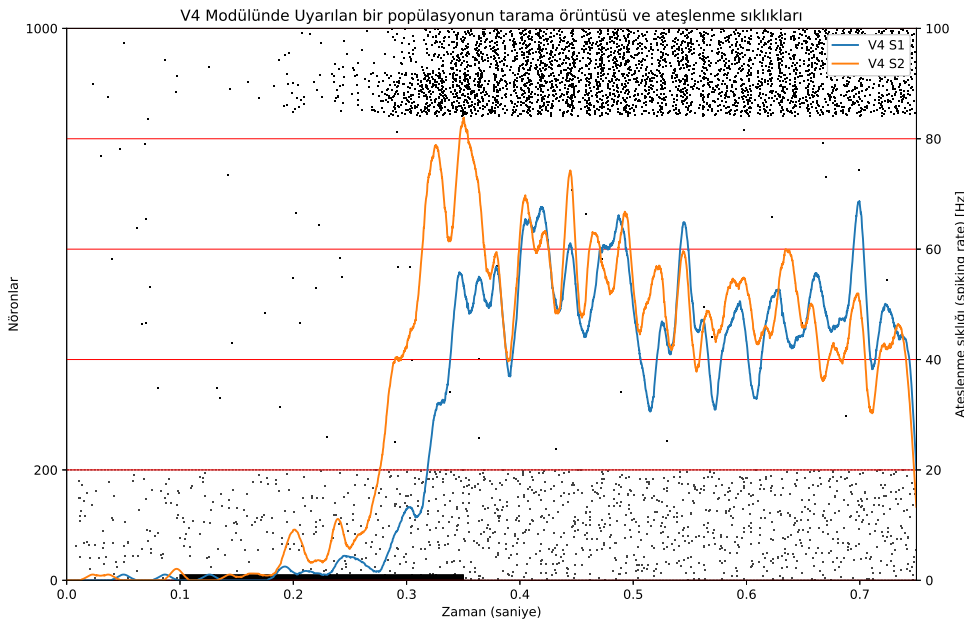
Bu çalışmanın sonucunda yapılan deneyde, görsel algıyı oluşturan modüllerden ikisini, V2 ve V4 modüllerini 1000'er sinire sahip modüller olarak aldık, ve bu modüllerde 80'er grupları değişmez gösterimleri temsil edecek şekilde ayırdık. Bu iki modüldeki farklı grupları kendi aralarında, ve modüller arasında birbirleri ile yanlı rekabete girecek şekilde bağlantılandırdık. Bunun sonucunda bir gösterime yapılan uyarının (denge durumundakinden daha yüksek bir gürültü olarak modellenmiştir) yanlı rekabette bir fark yaratıp yaratmadığına baktık. Yanlı rekabet söz konusu ise, V2 grubunda uyarının uygulandığı gruba bağlı olarak V4'deki grupların uyarılmalarında azımsanamayacak bir fark görülmelidir.

Bağlantıları açıklandığı gibi yapılmış olan sinir gruplarından, S1'e odağı temsil eden gürültü ve S1 ve S2'ye aynı anda cismin görüş alanına girdiğine dair gürültü $t = 100ms$ 'den $t = 350ms$ 'ye kadar tatbik edilmiştir. Bunun sonucunda Şekil 5.1 ve Şekil 5.2'de görüldüğü gibi, uyarının söz konusu olduğu zaman diliminde V2'deki S1 ve S2 grupları, ve V2'deki artış dolayısı ile V4'deki S1' ve S2' gruplarındaki aktivite artış göstermiş ve bu popülasyonu sürekli atma yapacak şekilde uyarmıştır.

Odaklanmanın Deco ve Rolls (2005) daki gibi bir etkisinin olduğunu gösterebilmek için, odağın S1 ile S2'ye tatbik edilmesinin gözlemlenebilir bir farka neden olduğunu gösterebilmemiz gerekmektedir. Şekil 5.2'deki gibi tekil deneyler, bu konuda yanıltıcı olabileceğinden, makaledeki gibi bir ortalama alıp karşılaştırmak gerekmektedir. Bu ortalamanın sonucunu Şekil 5.3'de görebilirsiniz. Bu ortalamadan da görülebileceği üzere, yaptığımız deneyin sonucunda Deco ve Rolls (2005) makalesinde gözlemlendiği gibi bir asimetri gözlemlenmedi. Bunun yanı sıra, Brunel ve Wang (2001) sinir modelinin önemli bir özelliği olan tekrarlı uyarıma sahip olma, *recurrent excitation*, elde ettiğimiz sonuçlarda gözlemlenmekte, uyarılan V2 ve V4 nöronları yükseltilmiş bir aktivite seviyesinde kalmaktadır. Buna karşın doğrulamaya çalıştığımız davranışta,



Şekil 5.1 : Bu şekilde uyarının uygulandığı zaman alttaki siyah bant ile gösterilmiştir. Mavi grafik tüm uyarıcı nöronların, turuncu grafik, tüm baskılayıcı nöronların aktivitesinin 10Hz lik bir pencere ile yumuşatılmış gösterimidir. Arka plandaki örüntüde en alttaki ilk 200 nöron baskılayıcı nöronları, En üstteki 920-1000 aralığındaki nöronlar S1, 840-920 aralığındaki nöronlar ise S2 grubundaki nöronlardır.



Şekil 5.2 : Şekil 5.1 deki gibi, V4 modülünde S1' ve S2' ye uyarıların etkisinin grafiğidir. Buradaki etki, V2 ye tatbik edilen uyarının dolaylı sonucudur.

aktivitede üssel bir düşüş gözlemlenmektedir. Burada bir şekilde Deco ve Rolls (2005) tekrarlı uyarılmanın gözlemlendiği parametre uzayından çıkmış olabilir.

Çalışmada her ne kadar dikkatin nörodinamik altyapısı ile ilgili bir kazanım elde edilememiş de olsa, bazı konularda çalışabilmek için bir ara basamak olabileceğini düşünüyorum. Modellemede, vuru sıklığını temel alan modeller ile topla ve ateşle sınırları gözlemlenen fenomenlere açıklık getirmede bir çekişme içerisindedir. Dünyadaki bilgi işlem gücü arttıkça topla ve ateşle nöron modelleri gibi detaylı modeller gittikçe daha büyük nöron gruplarını benzetebilmektedir. Bunun yanı sıra, vuru sıklığı dinamiğini temel olan modeller ise daha kesin ölçümlere dayandıkça ve daha açıklayıcı modeller bulundukça büyük nöron gruplarındaki davranışı gittikçe daha iyi açıklayabilmektedir. Dolayısı ile bu iki uçtaki modelleme tarzları, bir süre sonra aralarındaki ölçekten kaynaklanan boşluğu giderecek şekilde birbirlerini yaptıkları karakterizasyonlar ile iyileştirecektir. Farklı modelleme pratikleri ile ilgili bir tartışma Poggio ve Serre (2013)'de bulunabilir.

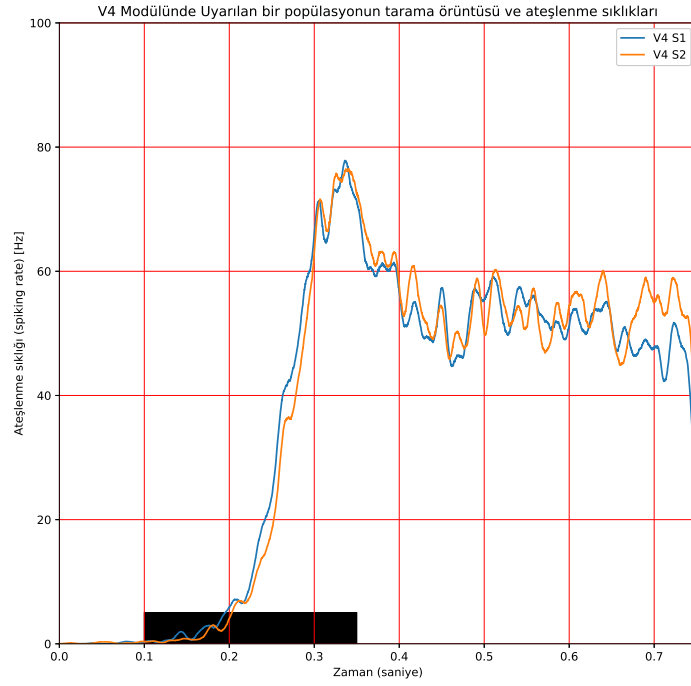
Özellikle topla ve ateşle nöronları, temel mantık bloklarını (Kandel, 2013, Appendix F), tekrarlanan bastırılmayı ya da uyarılmayı, doğrusal olmayan fenomenleri sergilediğinden karmaşık sistem davranışlarını açıklayabilecek bir aday olmayı sürdürebilir. Dolayısı ile, topla ve ateşle siniri gibi birimsel (*Ing:unitary*) bir model ile yapılmış bu benzetim, benzerlerine bir taban sağlayabilir.

Bunun yanı sıra, dorsal ve ventral yolların benzetimininde de ilerlemeye değer bazı noktalar bulunmaktadır. Öncelikle, Deco ve Rolls (2005) daki indirgenmiş modelin bu çalışmada kullanılan model ile gerçekleşmesi olası görünmektedir. V1-Dorsal ve V1-Ventral kolların V şeklinde bağlantılandırılmış modüllerinin nasıl etkileşeceği, bu topoloji ile daha iyi açıklayıcılığa sahip modeller oluşturulabilir mi bakılabilir. Aynı zamanda gözetimli öğrenme ile gerçek örüntülerin Gabor gösterimlerinin çıkarılması ile uygulamaya daha yakın araştırmalar da yapılabilir. Burada elde edilmesinde yarar bulunan becerilerden biri, ortalama alan teorisi ile bu modelin farklı fizyolojik ortamlarda da kullanılmasını sağlayacak parametre değişikliklerinin yapılabilmesi olabilir. Bu sayede hitap edilen fizyolojik hacim de görsel korteksin dışına çıkabilir, aynı zamanda farklı bağlantı koşullarında halen tekrarlanan uyarı ve baskılama rejiminde kaldığımızdan emin olarak yeni denemeler yapılabilir.

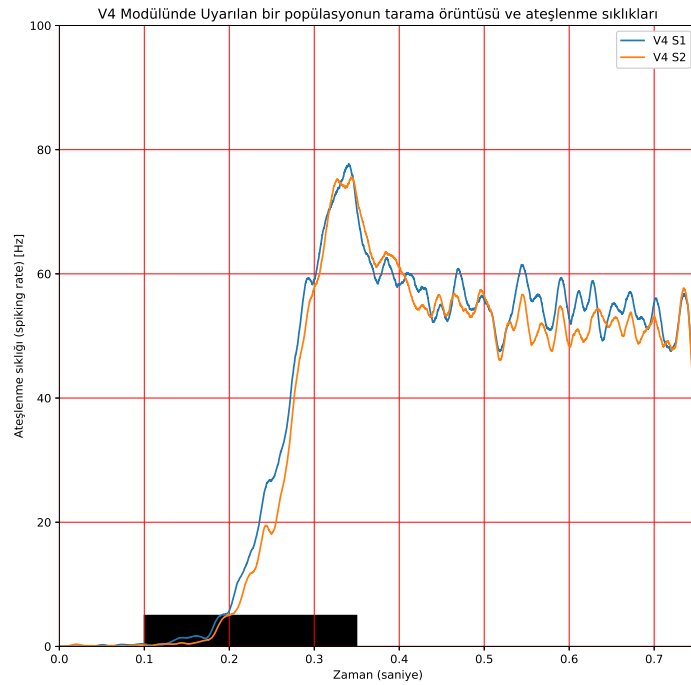
Sinirbilim makalelerinde bu tip deneylerin ifade biçimleri deneylerin yeniden üretilmesini güç ya da imkansız kılabilmektedir. Deney kurgusunu Deco ve Rolls (2005)'ü baz alarak oluşturduğumuzdan, buradaki sonuçlar aynı zamanda bahsi geçen makaledeki sonuçların açık bir yazılım kütüphanesi ile yeniden üretilmesine de yaramaktadır. Bu, birçok yinleme ve gerçeğe uygunluk testi ile mümkün olmuştur, ve çalışmanın BRIAN2 ile gerçekleştirilmesinden ve kodu açık tutulacağından bundan sonraki bir çok çalışmayı kolaylaştıracaktır. Bu çalışmanın çıktılarından birisi, açık kaynak olarak bulunabilecek bir çalışma defteridir ¹.

Büyük resimde ise üstten aşağı ve alttan yukarı süreçlerin birlikte oynayabildiği ve davranışsal makro fenomenlere ipucu uzatan mikroskopik koşulları bulmak önem teşkil etmektedir. Bunun için davranışsal literatüre hakimiyet temeldir, ve belki disiplinlerarası grup çalışmalarından destek görebilir.

¹Jupyter defteri olarak yapılan paylaşımlar ile bu çalışmada yapılan tüm denemeler ve çalışmanın tüm deneyleri başkası tarafından yeniden çalıştırılabilir, ve kolaylıkla ilerletilebilir. Github sitesindeki mehmatalianil kullanıcısının yazılım havuzunda tutulacaktır.



(a) Odak S2



(b) Odak S1

Şekil 5.3 : 8 tane deneyin sonucunda V4 grubundaki S1' ve S2' alt gruplarındaki aktivitenin ortalaması. Bu ortalamada Deco ve Rolls (2005) dan farklı olarak odağa bağlı bir fenomen gözlemlenmemiştir.



KAYNAKLAR

- Aviv, R. (2018, April 2). The Edge of Identity. *The New Yorker*.
- Beck, D. M., & Kastner, S. (2009). Top-down and bottom-up mechanisms in biasing competition in the human brain. *Vision research*, 49(10), 1154–1165.
- Breitmeyer, B. G., & Ogmen, H. (2007). Visual masking. *Scholarpedia*, 2(7), 3330. revision #91927. doi:10.4249/scholarpedia.3330
- Breitmeyer, B., & Ogmen, H. (2006). *Visual masking: time slices through conscious and unconscious vision*. Oxford Psychology Series. OUP Oxford. Retrieved from <https://books.google.com.tr/books?id=OGW2Lpu2p0sC>
- Brunel, N., & Wang, X.-J. (2001). Effects of neuromodulation in a cortical network model of object working memory dominated by recurrent inhibition. *Journal of computational neuroscience*, 11(1), 63–85.
- Deco, G., & Lee, T. S. (2004). The role of early visual cortex in visual integration: a neural model of recurrent interaction. *European Journal of Neuroscience*, 20(4), 1089–1100.
- Deco, G., & Rolls, E. T. [Edmund T]. (2005). Neurodynamics of biased competition and cooperation for attention: a model with spiking neurons. *Journal of neurophysiology*, 94(1), 295–313.
- Deco, G., & Rolls, E. T. [Edmund T.]. (2003). Attention and working memory: a dynamical model of neuronal activity in the prefrontal cortex. *European Journal of Neuroscience*, 18(8), 2374–2390.
- Deco, G., & Rolls, E. T. [Edmund T.]. (2004, March 1). A neurodynamical cortical model of visual attention and invariant object recognition. *Vision Research*, 44(6), 621–642. doi:10.1016/j.visres.2003.09.037
- Deco, G., & Rolls, E. T. [Edmund T.]. (2005). Attention, short-term memory, and action selection: a unifying theory. *Progress in neurobiology*, 76(4), 236–256.
- Desimone, R., & Duncan, J. (1995). Neural mechanisms of selective visual attention. *Annual review of neuroscience*, 18(1), 193–222.
- Destexhe, A., Mainen, Z. F., & Sejnowski, T. J. (1998). Kinetic models of synaptic transmission. *Methods in neuronal modeling*, 2, 1–25.
- Enns, J. T., & Di Lollo, V. (2000). What's new in visual masking? *Trends in cognitive sciences*, 4(9), 345–352.
- Geng, J. J., & Behrmann, M. (2003). Selective visual attention and visual search: behavioral and neural mechanisms. In *Psychology of learning and motivation* (Vol. 42, pp. 157–191). Elsevier.
- Goodman, D. F. M., & Brette, R. (2009). The brian simulator. *Frontiers in Neuroscience*, 3. doi:10.3389/neuro.01.026.2009
- Hestrin, S., Sah, P., & Nicoll, R. A. (1990). Mechanisms generating the time course of dual component excitatory synaptic currents recorded in hippocampal slices. *Neuron*, 5(3), 247–253. doi:[https://doi.org/10.1016/0896-6273\(90\)90162-9](https://doi.org/10.1016/0896-6273(90)90162-9)
- Izhikevich, E. M. (2007). *Dynamical systems in neuroscience*. MIT press.
- Kandel, E. R. (2013). *Principles of neural science, fifth edition*. McGraw Hill Professional.

- Koch, C., & Segev, I.** (1998). *Methods in neuronal modeling: from ions to networks*. MIT press.
- Koch, C., & Ullman, S.** (1987). Shifts in selective visual attention: towards the underlying neural circuitry. In *Matters of intelligence* (pp. 115–141). Springer.
- Poggio, T., & Serre, T.** (2013). Models of visual cortex. *Scholarpedia*, 8(4), 3516. revision #149958. doi:10.4249/scholarpedia.3516
- Reynolds, J. H., Chelazzi, L., & Desimone, R.** (1999). Competitive mechanisms subserve attention in macaque areas v2 and v4. *Journal of Neuroscience*, 19(5), 1736–1753.
- Robertson, L., Treisman, A., Friedman-Hill, S., & Grabowecky, M.** (1997). The interaction of spatial and object pathways: evidence from balint's syndrome. *Journal of Cognitive Neuroscience*, 9(3), 295–317.
- Rothman, J.** (2018, April 2). As Real as It Gets. *The New Yorker*.
- Sacks, O.** (2009). *The man who mistook his wife for a hat*. Pan Macmillan UK. Retrieved from <https://books.google.com/books?id=OmSV5tKzvboC>
- Taylor, J. B.** (2008). My stroke of insight. Retrieved May 1, 2018, from https://www.ted.com/talks/jill_bolte_taylor_s_powerful_stroke_of_insight
- Treisman, A. M., & Gelade, G.** (1980). A feature-integration theory of attention. *Cognitive psychology*, 12(1), 97–136.
- Ungerleider, L. G., & Pessoa, L.** (2008). What and where pathways. *Scholarpedia*, 3(11), 5342.
- Wende, K., & Jansen, A.** (2015). Top-down and/or bottom-up causality: the notion of relatedness in the human brain. In *Advances in cognitive neurodynamics (v): proceedings of the fifth international conference on cognitive neurodynamics - 2015* (p. 169).

EKLER

EK A : Brian ile İlgili Notlar

EK B : Modelin sınanmasında kullanılan kaynak kodu

EK C : Deneyde kullanılan kaynak kodu





EK A Brian ile İlgili Notlar

Bitirme çalışmasının zaman olarak hatırı sayılır bir kısmı Brian2 kütüphanesini kullanarak benzetimin kurgulanması, çalışır hale getirilmesi ve elde edilen sonuçların gerçekte örtüştüğünden emin olmakla geçmektedir. Dolayısı ile sonuçlardan bağımsız olarak sadece Brian2 ile ilgili bazı noktalara değinmekte yarar vardır. Brian2, hesaplamalı sinirbilimde kullanılan bazı modelleri, dinamik denklem çözümü için sık sık kullanılan bazı basitleştirmeleri bünyesinde tutan Python dilinde yazılmış bir kütüphanedir. (Goodman & Brette, 2009)

Dikkat Edilmesi Gereken Noktalar

- Python gibi yorumlanan (interpreted) bir dil için yazıldığından, sıklıkla alt seviye cebirsel işlem gerektiren simulasyonlarda, mesela her saat işaretinde çözülen (clock-driven) denklemlerde çok yavaşlayabilir.
- Bu yavaşlık, belli bir oranda problemi Python yorumlayıcı ile değil, C dilinde derlenip koşmasını sağlayarak azaltılabilir. Bunun için

```
set_device('cpp_standalone')
```

fonksiyonu kullanılabilir.

- C ile derlenmek üzere kodlanan bir benzetimde, Brian2 nin Python sözdizimindeki gibi kullandığı bazı liste dilimleme operasyonları kullanılamamaktadır, dolayısı ile bir

```
set_device('cpp_standalone')
```

```
Sexc_exc.w[:,400:480] = w_minus # yerine  
Sexc_exc.w['j < 480 and 400 < j'] = w_minus # kullanilmali
```

- Aynı zamanda hangi işlemlerin her saat işaretinde yapılmasını gerektiğini belirtmek benzetimi çok hızlandırabilir.
- Eğer problem çok daha büyük ise, doğru tanımlandığından emin olunmuş bir kod CUDA yardımı ile bir GPU ya derlenebilir. Bunun için bir bulut servisinden GPU kiralayıp çok küçük meblağlara benzetiminizi yapabilirsiniz.
- Brian2 nin paralel işlem yetenekleri halen beklenmedik problemlere neden olmaktadır. Bu çalışmanın büyük bir kesiminde, bir bulut hizmetinden çok çekirdekli sanal makineler kullanılmış, ve benzetimler gerçekçi görünmüştür. Ama simulasyonların çekirdek sayısı değiştirilip yeniden koşturulduğunda tekrarlanmadığı görüldüğü için vazgeçilmiştir. Bu konunun gerçekleşmesi yeterli düzeyde OpenMP bilgisi gerektirmektedir, dolayısı ile tavsiye edilmez. Bu konuda yazılımın ekrana verdiği uyarıyı dikkate alın.

- Parametrelere rastgele sayılar atandığında tüm sinapslar için atıldığı unutulmamalıdır. Burada mesela kabaca 640000 tane w ataması yapılmaktadır, rastgele sayıların tekrarını önleyebilmek için seed komutunun çağrılmasına özen gösterilmelidir.

```
Sexc_exc.w = '0.1*rand()+2.1'  
seed()
```



EK B Modelin sınanmasında kullanılan kaynak kodu

Aşağıda örnek olması için kullanılan kod verilmiştir. Lütfen kaynak kodunu buradan değil, tez ile birlikte teslim edilen veri paketinde, ya da e-posta vasıtası ile alın.

```
from brian2 import *
import shelve
from time import mktime, gmtime

clean=True
set_device('cpp_standalone')

# MULTITHREADING OFF
# prefs.devices.cpp_standalone.openmp_threads = 3

# (Deco and Rolls 2003)
# (Brunel and Wang 2001)
# Neuron resting potential V_L = -70mV
# Threshold Theta = -50mV
# Reset V_reset = -55mV
# Membrane Capacitance Cm = 0.5 nF , 0.2nF (pyramidal,interneuron)
# Membrane leak gm = 25nS , 20nS
# Refractory period tau_r = 2ms , 1 ms
# membrane time constant tau_m = Cm/gm = 20ms, 10ms

tau_m = 20.*ms
C_m_pym = 0.5*nF
C_m_inter = 0.2*nF
g_m_pym = 25.*nS
g_m_inter = 20.*nS
tau_r = 2.*ms
theta = -50.*mV
V_L=-70.*mV
V_r=-55.*mV

## Pyramidal Neuron Parameters for
g_GABA_pym = 1.25 * nS
g_AMPArec_pym= 0.104 * nS
g_AMPAext_pym = 2.08 * nS
g_NMDA_pym = 0.327 * nS
## Interneuron Parameters for (Brunel Wang)
g_GABA_inter = 0.973 * nS
g_AMPArec_inter = 0.081 * nS
g_AMPAext_inter = 1.62 * nS
g_NMDA_inter = 0.258 * nS

CMg = 1
V_E = 0 * mV
NMDA_exponent = (-0.062) / mV
```

```

alpha = 0.5 / ms
tau_AMPA = 2. *ms
tau_NMDArise = 2. *ms
tau_NMDAdecay = 100.* ms
tau_GABA = 10 *ms

Nexcitatory = 800
Ninhibitory = 200
Npool = Nexcitatory+Ninhibitory

# external stimuli
rate = 3 * Hz

# subpopulations
f = 0.1
p = 5
N_sub = int(Nexcitatory * f)
N_non = int(Ninhibitory * (1. - f * p))
w_plus = 2.1
w_minus = 1. - f * (w_plus - 1.) / (1. - f)

eqs_pym = '''
dv/dt = (-1) * g_m_pym * (v - V_L) /C_m_pym -
        I /C_m_pym : volt (unless refractory)
I = I_GABA + I_AMPArec + I_AMPAext + I_NMDA : amp

I_GABA = g_GABA_pym * (v-V_L)*S_GABA : amp
dS_GABA/dt =(-1) * S_GABA/tau_GABA : 1

I_AMPArec = g_AMPArec_pym * (v-V_E)*S_AMPArec : amp
dS_AMPArec/dt =(-1) * S_AMPArec/tau_AMPA : 1

I_AMPAext = g_AMPAext_pym * (v-V_E)*S_AMPAext : amp
dS_AMPAext/dt = (-1) * S_AMPAext/tau_AMPA : 1

I_NMDA = g_NMDA_pym * (v-V_E) / (1+CMg*exp(NMDA_exponent*v) / 3.57) *
        Sum_S_NMDA : amp
Sum_S_NMDA : 1
'''

eqs_inter = '''
dv/dt = (-1) * g_m_inter * (v - V_L) /C_m_inter -
        I /C_m_inter : volt (unless refractory)
I = I_GABA + I_AMPArec + I_AMPAext + I_NMDA : amp

I_GABA = g_GABA_inter * (v-V_L)*S_GABA : amp
dS_GABA/dt =(-1) * S_GABA/tau_GABA : 1

I_AMPArec = g_AMPArec_inter * (v-V_E)*S_AMPArec : amp
dS_AMPArec/dt =(-1) * S_AMPArec/tau_AMPA : 1

I_AMPAext = g_AMPAext_inter * (v-V_E)*S_AMPAext : amp
dS_AMPAext/dt = (-1) * S_AMPAext/tau_AMPA : 1

I_NMDA = g_NMDA_inter * (v-V_E) / (1+CMg*exp(NMDA_exponent*v) / 3.57)
        * Sum_S_NMDA : amp
Sum_S_NMDA : 1
'''

```

```

# Noronlari Yaratiyoruz
inhibitory = NeuronGroup(Ninhibitory, eqs_inter ,
                        threshold='v > theta', reset='v = V_r',
                        refractory=1*ms, method='euler', name = "InhibitoryGroup")
excitatory = NeuronGroup(Nexcitatory, eqs_pym ,
                        threshold='v > theta', reset='v = V_r',
                        refractory=2*ms, method='euler', name = "ExcitatoryGroup")

eqs_glut = '''
Sum_S_NMDA_post = w*S_NMDA : 1 (summed)
dS_NMDA/dt = (-1) * S_NMDA/tau_NMDAdecay +
              alpha*x*(1-S_NMDA) : 1 (clock-driven)
dx/dt = (-1) * x / tau_NMDArise : 1 (clock-driven)
w : 1
'''

eqs_pre_glut = '''
S_AMPAreac += w
x += 1
'''

eqs_pre_gaba = '''
S_GABA += 1
'''

eqs_pre_ext = '''
S_AMPAext += 1
'''

# Poisson Giris --- (0--- havuz noronlari sinapslari
Sinput_exc = PoissonInput(excitatory, 'S_AMPAext',
                          Nexcitatory, rate, '1')
Sinput_inh = PoissonInput(inhibitory, 'S_AMPAext',
                          Nexcitatory, rate, '1')
# Poisson Giris --- (0--- inhibiting havuz noronlari sinapslari

# havuz noronlari --- (0--- havuz noronlari sinapslari
Sexc_exc = Synapses(excitatory, excitatory, model=eqs_glut,
                    on_pre=eqs_pre_glut, method='euler',
                    name="ExcitatoryLoopback" )
# excit noronlari --- (0--- inhib noronlari sinapslari
Sexc_inh = Synapses(excitatory, inhibitory, model=eqs_glut,
                    on_pre=eqs_pre_glut, method='euler',
                    name = "InhibitorExciting")
# inhib noronlari --- (0--- excit noronlari sinapslari
Sinh_exc = Synapses(inhibitory, excitatory, on_pre=eqs_pre_gaba,
                    method='euler', name = "ExcitatoryInhibiting")
# inhib noronlari --- (0--- excit noronlari sinapslari
Sinh_inh = Synapses(inhibitory, inhibitory, on_pre=eqs_pre_gaba,
                    method='euler', name = "InhibitoryLoopback")

# exc noronlari --- (0--- exc noronlari sinapslari
# 0.2 olasilik ile baglantilandi
Sexc_exc.connect('i != j', p=1)
# Sexc_exc.w = '0.1*rand()+2.1'
Sexc_exc.w[:] = 1
# seed()

```

```

# internal other subpopulation to current nonselective
Sexc_exc.w['j < 480 and 400 < j'] = w_minus
Sexc_exc.w['j < 560 and 480 < j'] = w_minus
Sexc_exc.w['j < 640 and 560 < j'] = w_minus
Sexc_exc.w['j < 720 and 640 < j'] = w_minus
Sexc_exc.w['j < 800 and 720 < j'] = w_minus

# internal current subpopulation to current subpopulation
Sexc_exc.w['i < 480 and 400 < i and j < 480 and 400 < j'] = w_plus
Sexc_exc.w['i < 560 and 480 < i and j < 560 and 480 < j'] = w_plus
Sexc_exc.w['i < 640 and 560 < i and j < 640 and 560 < j'] = w_plus
Sexc_exc.w['i < 720 and 640 < i and j < 720 and 640 < j'] = w_plus
Sexc_exc.w['i < 800 and 720 < i and j < 800 and 720 < j'] = w_plus

#for pi in range(N_non, N_non + p * N_sub, N_sub):
#internal other subpopulation to current nonselective
#Sexc_exc.w[Sexc_exc.indices[:, pi:pi + N_sub]] = w_minus
#internal current subpopulation to current subpopulation
#Sexc_exc.w[Sexc_exc.indices[pi:pi + N_sub, pi:pi + N_sub]] = w_plus

# exc noronlari --- (0--- inh noronlari sinapslari
# 1 olasilik ile baglantilendirildi
Sexc_inh.connect(p=1)
#Sexc_inh.w = '0.1*rand()+2.1'
Sexc_inh.w[:] = 1
seed()

# inh noronlari --- (0--- exc noronlari bsinapslari
# 0.5 olasilik ile baglantilendirildi
Sinh_exc.connect(p=1)

# inh noronlari --- (0--- inh noronlari sinapslari
# 0.5 olasilik ile baglantilendirildi
Sinh_inh.connect('i != j', p=1)

net = Network([excitatory, inhibitory, Sinh_exc, Sexc_inh,
               Sexc_exc, Sinput_exc, Sinput_inh, Sinh_inh])

# baslangic degerleri gerilimleri -65mV +- 2mV
excitatory.v = ' (rand()*2-65)*mV'
inhibitory.v = ' (rand()*2-65)*mV'

# Durum ve Spike monitorleri
Mspike_exc = SpikeMonitor(excitatory[0:800])
Mspike_inh = SpikeMonitor(inhibitory[0:200])
Mexc = StateMonitor(excitatory, True, record=True)
Minh = StateMonitor(inhibitory, True, record=True)
net.add([Mspike_exc, Mspike_inh, Mexc, Minh])
Pop_inh = PopulationRateMonitor(inhibitory)
Pop_exc = PopulationRateMonitor(excitatory)
net.add([Pop_exc, Pop_inh])

net.run(4000*ms, report='stdout')

print ("Excitatory Spikes:" + str(Mspike_exc.num_spikes))
print ("Inhibitory Spikes:" + str(Mspike_inh.num_spikes))

```

EK C Deneyde kullanılan kaynak kodu

Aşağıda örnek olması için kullanılan kod verilmiştir. Lütfen kaynak kodunu buradan değil, tez ile birlikte teslim edilen veri paketinde, ya da e-posta vasıtası ile alın. Tüm deneylerin Jupyter defteri mevcuttur.

Genel parametreler

```
from brian2 import *
import shelve
from time import mktime, gmtime

clean=True
set_device('cpp_standalone')

# These are Deco Rolls 2005a for V2 and V4. Forward connections
# are dominant. Ratio ~ 2.5, in line with physio.

Jb = 0.5
Jf = 1.6
c = 0.1 # For Martinez 0.075

Kb = 0.06
Kf = 0.15

Nexcitatory = 800
Ninhibitory = 200
Npool = Nexcitatory+Ninhibitory

# external stimuli
rate = 3 * Hz
rate_attention = 10 * Hz
rate_input = 250 * Hz

# subpopulations
f = 0.1
p = 2 # We want 2 subpopulations here, S1 and S2.
N_sub = int(Nexcitatory * f)
N_non = int(Ninhibitory * (1. - f * p))
w_plus = 1.5 #This was 2.1
w_minus = 1. - f * (w_plus - 1.) / (1. - f)
# so that the overall recurrent excitatory
# synaptic drive in the spontaneous state
# remains constant as w+ is varied (Deco 2005a)
# inhibition
w_i = 1 # V2

# This is because the inhibition in both layers
```

```

# is local and therefore the stronger the neighborhood
# relationship, the stronger the inhibition. V4 has
# the e greater degree of overlapping
# of the receptive fields, the stronger the competition

# non selective to selective
w_n      = (-f*Jb-f*Kb)/(1-2*f)+w_minus   # V2

```

V4 için parametreler

```

# inhibition
w_i_prime = 1.25 # V4

# nonselective to selective
w_n_prime = (-f*Jf-f*Kf)/(1-2*f)+w_minus # V4

```

Tek modülün modeli

```

#prefs.devices.cpp_standalone.openmp_threads = 2

# (Deco and Rolls 2003)
# (Brunel and Wang 2001)
# Neuron resting potential V_L = -70mV
# Threshold Theta = -50mV
# Reset V_reset = -55mV
# Membrane Capacitance Cm = 0.5 nF , 0.2nF (pyramidal,interneuron)
# Membrane leak gm = 25nS , 20nS
# Refractory period tau_r = 2ms , 1 ms
# membrane time constant tau_m = Cm/gm = 20ms, 10ms

tau_m = 20.*ms
C_m_pym = 0.5*nF
C_m_inter = 0.2*nF
g_m_pym = 25.*nS
g_m_inter = 20.*nS
tau_r = 2.*ms
theta = -50.*mV
V_L=-70.*mV
V_r=-55.*mV

## Pyramidal Neuron Parameters for
g_GABA_pym = 1.287 * nS
g_AMPAreC_pym= 0.104 * nS
g_AMPAext_pym = 2.08 * nS
g_NMDA_pym = 0.327 * nS
g_AHP_pym = 7.5 * nS
## Interneuron Parameters for (Brunel Wang)
g_GABA_inter = 1.002 * nS
g_AMPAreC_inter = 0.081 * nS
g_AMPAext_inter = 1.62 * nS

```

```

g_NMDA_inter = 0.258 * nS
g_AHP_inter = g_AHP_pym

CMg = 1
V_E = 0 * mV
V_K = -80 * mV
NMDA_exponent = (-0.062) / mV
alpha = 0.5 / ms
alphaCa = 0.005
tau_AMPA = 2. *ms
tau_NMDArise = 2. *ms
tau_NMDAdecay = 100.* ms
tau_GABA = 10 *ms
tau_Ca = 600 * ms

eqs_pym = '''
dv/dt = (-1) * g_m_pym * (v - V_L) / C_m_pym -
        I / C_m_pym : volt (unless refractory)
I = I_GABA + I_AMPArec + I_AMPAext + I_NMDA : amp

I_GABA = g_GABA_pym * (v-V_L)*S_GABA : amp
dS_GABA/dt = (-1) * S_GABA/tau_GABA : 1

I_AMPArec = g_AMPArec_pym * (v-V_E)*S_AMPArec : amp
dS_AMPArec/dt = (-1) * S_AMPArec/tau_AMPA : 1

I_AMPAext = g_AMPAext_pym * (v-V_E)*S_AMPAext : amp
dS_AMPAext/dt = (-1) * S_AMPAext/tau_AMPA : 1

I_NMDA = g_NMDA_pym * (v-V_E) / (1+CMg*exp(NMDA_exponent*v)/3.57)
        * (Sum_S_NMDA+Sum_S_NMDA_v2) : amp
Sum_S_NMDA : 1
Sum_S_NMDA_v2 : 1
'''

eqs_inter = '''
dv/dt = (-1) * g_m_inter * (v - V_L) / C_m_inter -
        I / C_m_inter : volt (unless refractory)
I = I_GABA + I_AMPArec + I_AMPAext + I_NMDA : amp

I_GABA = g_GABA_inter * (v-V_L)*S_GABA : amp
dS_GABA/dt = (-1) * S_GABA/tau_GABA : 1

I_AMPArec = g_AMPArec_inter * (v-V_E)*S_AMPArec : amp
dS_AMPArec/dt = (-1) * S_AMPArec/tau_AMPA : 1

I_AMPAext = g_AMPAext_inter * (v-V_E)*S_AMPAext : amp
dS_AMPAext/dt = (-1) * S_AMPAext/tau_AMPA : 1

I_NMDA = g_NMDA_inter * (v-V_E) /
        (1+CMg*exp(NMDA_exponent*v)/3.57) * Sum_S_NMDA : amp
Sum_S_NMDA : 1
'''

reset_glut = '''
v = V_r
'''

```

```

# Noronlari Yaratiyoruz
inhibitory = NeuronGroup(Ninhibitory, eqs_inter ,
    threshold='v > theta', reset=reset_glut,
    refractory=1*ms, method='euler', name = "InhibitoryGroupv2")
excitatory = NeuronGroup(Nexcitatory, eqs_pym ,
    threshold='v > theta', reset=reset_glut,
    refractory=2*ms, method='euler', name = "ExcitatoryGroupv2")

# Noronlari Yaratiyoruz
inhibitory_prime = NeuronGroup(Ninhibitory, eqs_inter ,
    threshold='v > theta', reset=reset_glut,
    refractory=1*ms, method='euler', name = "InhibitoryGroupv4")
excitatory_prime = NeuronGroup(Nexcitatory, eqs_pym ,
    threshold='v > theta', reset=reset_glut,
    refractory=2*ms, method='euler', name = "ExcitatoryGroupv4")

eqs_glut = '''
Sum_S_NMDA_post = w*S_NMDA : 1 (summed)
dS_NMDA/dt = (-1) * S_NMDA/tau_NMDAdecay +
alpha*x*(1-S_NMDA) : 1 (clock-driven)
dx/dt = (-1) * x / tau_NMDArise : 1 (clock-driven)
w : 1
'''

eqs_glut_v2 = '''
Sum_S_NMDA_v2_post = w*S_NMDA : 1 (summed)
dS_NMDA/dt = (-1) * S_NMDA/tau_NMDAdecay +
alpha*x*(1-S_NMDA) : 1 (clock-driven)
dx/dt = (-1) * x / tau_NMDArise : 1 (clock-driven)
w : 1
'''

eqs_gaba = '''
w : 1
'''

eqs_pre_glut = '''
S_AMPAreac += w
x += 1
'''

# Here, GABA is changed to accept a weight.
eqs_pre_gaba = '''
S_GABA += w
'''

eqs_pre_ext = '''
S_AMPAAext += 1
'''

# Poisson Giris --- (0--- havuz noronlari sinapslari
Sinput_exc = PoissonInput(excitatory, 'S_AMPAAext',
    Nexcitatory, rate, '1')
Sinput_inh = PoissonInput(inhibitory, 'S_AMPAAext',
    Nexcitatory, rate, '1')
# Poisson Giris --- (0--- havuz noronlari sinapslari
Sinput_exc_prime = PoissonInput(excitatory_prime,
    'S_AMPAAext', Nexcitatory, rate, '1')
Sinput_inh_prime = PoissonInput(inhibitory_prime,

```



```
'S_AMPAext', Nexcitatory, rate, '1')
```

V2 Bağlantıları

```
# havuz noronlari --- (0--- havuz noronlari sinapslari
Sexc_exc = Synapses(excitatory, excitatory, model=eqs_glut,
    on_pre=eqs_pre_glut, method='euler', name="ExcitatoryLoopbackV2" )
# excit noronlari --- (0--- inhib noronlari sinapslari
Sexc_inh = Synapses(excitatory, inhibitory, model=eqs_glut,
    on_pre=eqs_pre_glut, method='euler', name = "InhibitorExcitingV2")
# inhib noronlari --- (0--- excit noronlari sinapslari
Sinh_exc = Synapses(inhibitory, excitatory, model=eqs_gaba,
    on_pre=eqs_pre_gaba, method='euler', name = "ExcitatoryInhibitingV2")
# inhib noronlari --- (0--- excit noronlari sinapslari
Sinh_inh = Synapses(inhibitory, inhibitory, model=eqs_gaba,
    on_pre=eqs_pre_gaba, method='euler', name = "InhibitoryLoopbackV2")

# exc noronlari --- (0--- exc noronlari sinapslari
Sexc_exc.connect('i != j', p=1)
# Sexc_exc.w = '0.1*rand()+2.1'
Sexc_exc.w[:] = 1
# seed()

# all connections to selective
Sexc_exc.w['j < 800 and 640 < j'] = w_n

# selective connections to selective
Sexc_exc.w['i < 800 and 640 < i and j < 800 and 640 < j'] = w_minus

# internal current subpopulation to current subpopulation
Sexc_exc.w['i < 720 and 640 < i and j < 720 and 640 < j'] = w_plus
Sexc_exc.w['i < 800 and 720 < i and j < 800 and 720 < j'] = w_plus

# exc noronlari --- (0--- inh noronlari sinapslari
Sexc_inh.connect(p=1)
# Sexc_inh.w = '0.1*rand()+2.1'
Sexc_inh.w[:] = 1
seed()

# inh noronlari --- (0--- exc noronlari sinapslari
Sinh_exc.connect(p=1)
Sinh_exc.w[:] = w_i

# inh noronlari --- (0--- inh noronlari sinapslari
Sinh_inh.connect('i != j', p=1)
Sinh_inh.w[:] = 1
```

V4 bağlantıları

```
# havuz noronlari --- (0--- havuz noronlari sinapslari
```

```

Sexc_exc_prime = Synapses(excitatory_prime, excitatory_prime,
model=eqs_glut, on_pre=eqs_pre_glut, method='euler'
,name="ExcitatoryLoopbackV4" )
# excit noronlari --- (0--- inhib noronlari sinapslari
Sexc_inh_prime = Synapses(excitatory_prime, inhibitory_prime,
model=eqs_glut, on_pre=eqs_pre_glut, method='euler'
,name = "InhibitorExcitingV4")
# inhib noronlari --- (0--- excit noronlari sinapslari
Sinh_exc_prime = Synapses(inhibitory_prime, excitatory_prime,
model=eqs_gaba, on_pre=eqs_pre_gaba, method='euler'
,name = "ExcitatoryInhibitingV4")
# inhib noronlari --- (0--- excit noronlari sinapslari
Sinh_inh_prime = Synapses(inhibitory_prime, inhibitory_prime,
model=eqs_gaba, on_pre=eqs_pre_gaba, method='euler'
,name = "InhibitoryLoopbackV4")

# exc noronlari --- (0--- exc noronlari sinapslari
Sexc_exc_prime.connect('i != j', p=1)
# Sexc_exc.w = '0.1*rand()+2.1'
Sexc_exc_prime.w[:] = 1
# seed()

# all connections to selective
Sexc_exc_prime.w['j < 800 and 640 < j'] = w_n_prime

# selective connections to selective
Sexc_exc_prime.w['i < 800 and 640 < i and j < 800 and 640 < j'] = w_minus

# internal current subpopulation to current subpopulation
Sexc_exc_prime.w['i < 720 and 640 < i and j < 720 and 640 < j'] = w_plus
Sexc_exc_prime.w['i < 800 and 720 < i and j < 800 and 720 < j'] = w_plus

# exc noronlari --- (0--- inh noronlari sinapslari
Sexc_inh_prime.connect(p=1)
#Sexc_inh.w = '0.1*rand()+2.1'
Sexc_inh_prime.w[:] = 1
seed()

# inh noronlari --- (0--- exc noronlari sinapslari
Sinh_exc_prime.connect(p=1)
Sinh_exc_prime.w[:] = w_i_prime

# inh noronlari --- (0--- inh noronlari sinapslari
Sinh_inh_prime.connect('i != j', p=1)
Sinh_inh_prime.w[:] = 1

```

V2 ve V4 arasındaki bağlantılar

```

Sv2_v4 = Synapses(excitatory, excitatory_prime,
model=eqs_glut_v2, on_pre=eqs_pre_glut, method='euler', name="V2toV4" )

Sv2_v4.connect('i < 800 and 640 < i and j < 800 and 640 < j', p=1)

# feedforward connections within associated selective pools

```

```

Sv2_v4.w['i < 800 and 720 < i and j < 800 and 720 < j'] = Jf
Sv2_v4.w['i < 720 and 640 < i and j < 720 and 640 < j'] = Jf

# feedforward connections within non associated selective pools
Sv2_v4.w['i < 720 and 640 < i and j < 800 and 720 < j'] = Kf
Sv2_v4.w['i < 800 and 720 < i and j < 720 and 640 < j'] = Kf

Sv4_v2 = Synapses(excitatory_prime, excitatory,
model=eqs_glut_v2, on_pre=eqs_pre_glut, method='euler', name="V4toV2" )
Sv4_v2.connect('i < 800 and 640 < i and j < 800 and 640 < j', p=1)

# feedforward connections within associated selective pools
Sv4_v2.w['i < 800 and 720 < i and j < 800 and 720 < j'] = Jb
Sv4_v2.w['i < 720 and 640 < i and j < 720 and 640 < j'] = Jb

# feedforward connections within non associated selective pools
Sv4_v2.w['i < 720 and 640 < i and j < 800 and 720 < j'] = Kb
Sv4_v2.w['i < 800 and 720 < i and j < 720 and 640 < j'] = Kb

```

Genel tanımlamalar ve koşturma

```

# baslangic degerleri gerilimleri -65mV +- 2mV
excitatory.v = '(rand()*2-65)*mV'
inhibitory.v = '(rand()*2-65)*mV'

# baslangic degerleri gerilimleri -65mV +- 2mV
excitatory_prime.v = '(rand()*2-65)*mV'
inhibitory_prime.v = '(rand()*2-65)*mV'

# at 1s, select population 1
C_selection = 1
rate_selection = rate_input
stimuli1 = TimedArray(np.r_[np.ones(10)*rate_attention/Hz,
np.ones(10)*(rate_selection/Hz+rate_attention/Hz),
np.ones(100)*rate_attention/Hz], dt=25 * ms)
stimuli2 = TimedArray(np.r_[np.zeros(10), np.ones(10),
np.zeros(100)], dt=25 * ms)

input1 = PoissonInput(excitatory[640:800], 'S_AMPAext',
C_selection, rate_selection, 'stimuli2(t)')
input2 = PoissonInput(excitatory[720:800], 'S_AMPAext',
C_selection, rate_attention, '1')

# at 4s, reset selection
#rate_selection = 25 * Hz
#stimuli_reset = TimedArray(np.r_[np.zeros(60), np.ones(10),
np.zeros(100)], dt=25 * ms)
#input_reset_I = PoissonInput(excitatory, 'S_AMPAext', 800,
rate_selection, 'stimuli_reset(t)')
#input_reset_E = PoissonInput(inhibitory, 'S_AMPAext', 800,
rate_selection, 'stimuli_reset(t)')

```

```

net = Network([excitatory, inhibitory, excitatory_prime, inhibitory_prime,
Sinh_exc, Sexc_inh, Sexc_exc, Sinput_exc, Sinput_inh, Sinh_inh,
Sinh_exc_prime, Sexc_inh_prime, Sexc_exc_prime,
Sinput_exc_prime, Sinput_inh_prime, Sinh_inh_prime,
Sv2_v4, Sv4_v2,
input1, input2,
#input_reset_I, input_reset_E
])

# Durum ve Spike monitorleri
Mspike_exc = SpikeMonitor(excitatory[0:800])
Mspike_inh = SpikeMonitor(inhibitory[0:200])
# Durum ve Spike monitorleri
Mspike_exc_prime = SpikeMonitor(excitatory_prime[0:800])
Mspike_inh_prime = SpikeMonitor(inhibitory_prime[0:200])
net.add([Mspike_exc, Mspike_inh, Mspike_exc_prime, Mspike_inh_prime])
Pop_inh = PopulationRateMonitor(inhibitory)
Pop_exc = PopulationRateMonitor(excitatory)
Pop_inh_prime = PopulationRateMonitor(inhibitory_prime)
Pop_exc_prime = PopulationRateMonitor(excitatory_prime)
Pop_v4_s1 = PopulationRateMonitor(excitatory_prime[720:800])
Pop_v4_s2 = PopulationRateMonitor(excitatory_prime[640:720])
net.add([Pop_exc, Pop_inh, Pop_inh_prime, Pop_exc_prime, Pop_v4_s1, Pop_v4_s2])

net.run(2000*ms, report='stdout')

print ("V2 Excitatory Spikes:" + str(Mspike_exc.num_spikes))
print ("V2 Inhibitory Spikes:" + str(Mspike_inh.num_spikes))

print ("V4 Excitatory Spikes:" + str(Mspike_exc_prime.num_spikes))
print ("V4 Inhibitory Spikes:" + str(Mspike_inh_prime.num_spikes))

```

Çizdirme

```

%matplotlib inline

import brian2tools as b2t
import numpy as np
from matplotlib import *

rcParams['figure.figsize'] = (18,18)

fig, ax1 = plt.subplots()
b2t.plot_raster(Mspike_exc.i+201, Mspike_exc.t,
time_unit=second, marker=',', color='k')
b2t.plot_raster(Mspike_inh.i+1, Mspike_inh.t,
time_unit=second, marker=',', color='#444444')
plt.yticks([0,200,1000])
plt.grid(True, 'major', 'y', ls='--', lw=.5, c='k', alpha=.3)
plt.title(u"V2 Modulunde Uyarilan bir populasyonun

```

```

    tarama oruntusu ve ateslenme sikliklari")
plt.xlabel(u"Zaman (saniye)")
plt.ylabel(u"Noronlar")
ylim((0,1000))
plt.plot()

ax2 = ax1.twinx()
ax2.plot(Pop_exc.t,Pop_exc.smooth_rate(width=10*ms)/ Hz
        ,label=u"Uyarici" )
ax2.plot(Pop_exc.t,Pop_inh.smooth_rate(width=10*ms)/ Hz
        ,label=u"Baskilayici" )
ylim((0,30))
xlim((.150,.750))
#b2t.plot_rate(Pop_exc.t,Pop_exc.smooth_rate(window='flat'
        , width=25*ms),label="Excitatory")
#ax = b2t.plot_rate(Pop_inh.t,Pop_inh.smooth_rate(window='flat'
        , width=25*ms),label="Inhibitory")
ax2.tick_params(axis='y',which='minor')
plt.xlabel(u"Zaman (saniye)")
plt.ylabel(u"Ateslenme sikligi (spiking rate) [Hz]")
plt.plot()
plt.legend()
plt.grid(b=True, which='major', ls='--', lw=0.5, c='r', alpha=.2)
plt.grid(b=True, which='minor', ls='--', lw=0.5, c='r', alpha=.2)
show()

fig.savefig("ExcitedRaster.eps", format="eps")

fig, ax1 = plt.subplots()
b2t.plot_raster(Mspike_exc_prime.i+201, Mspike_exc_prime.t,
        time_unit=second, marker=',', color='k')
b2t.plot_raster(Mspike_inh_prime.i+1, Mspike_inh_prime.t,
        time_unit=second, marker=',', color='#444444')
plt.yticks([0,200,1000])
plt.grid(True, 'major', 'y', ls='--', lw=.5, c='k', alpha=.3)
plt.title(u"V4 Modulunde Uyarilan bir populasyonun tarama
        oruntusu ve ateslenme sikliklari")
plt.xlabel(u"Zaman (saniye)")
plt.ylabel(u"Noronlar")
ylim((0,1000))
plt.plot()

ax2 = ax1.twinx()
ax2.plot(Pop_exc_prime.t
        ,Pop_exc_prime.smooth_rate(width=25*ms)/ Hz ,label=u"Uyarici" )
ax2.plot(Pop_v4_s1.t,
        Pop_v4_s1.smooth_rate(width=5*ms)/ Hz ,label=u"V4 S1" )
ax2.plot(Pop_v4_s2.t,
        Pop_v4_s2.smooth_rate(width=5*ms)/ Hz ,label=u"V4 S2" )
ax2.plot(Pop_inh_prime.t,
        Pop_inh_prime.smooth_rate(width=25*ms)/ Hz ,label=u"Baskilayici" )
ylim((0,100))
#b2t.plot_rate(Pop_exc.t,
        Pop_exc.smooth_rate(window='flat', width=25*ms),label="Excitatory")
#ax = b2t.plot_rate(Pop_inh.t,
        Pop_inh.smooth_rate(window='flat', width=25*ms),label="Inhibitory")
ax2.tick_params(axis='y',which='minor')

```

```
plt.xlabel(u"Zaman (saniye)")
plt.ylabel(u"Ateslenme sikligi (spiking rate) [Hz]")
plt.plot()
plt.legend()
plt.grid(b=True, which='major', ls='-', lw=0.5, c='r', alpha=.2)
plt.grid(b=True, which='minor', ls='--', lw=0.5, c='r', alpha=.2)
xlim((.150,.750))
show()

fig.savefig("ExcitedRaster2.eps", format="eps")
```



ÖZGEÇMİŞ

Ad Soyad: Mehmet Ali Anıl

Doğum Tarihi ve Yeri: 5 Haziran 1988 Viyana/Avusturya

E-Posta:m.a.anil@ieee.org

ÖĞRENİM DURUMU:

- **Lisans:** 2011, İstanbul Teknik Üniversitesi, Elektrik ve Elektronik Fakültesi, Elektronik Bölümü
- **Lisans:** 2012, İstanbul Teknik Üniversitesi, Fen Edebiyat Fakültesi, Fizik Mühendisliği Bölümü

MESLEKİ DENEYİMLER VE ÖDÜLLER:

- Grus BV de yöneticilik
- Pavo Tasarım da donanım mühendisliği
- Joint Institute of Laboratory Aeurnautics (JILA)'da Laboratuar asistanı
- University of Colorado Boulder'da Doktora Araştırmacısı

YAYINLAR, SUNUMLAR VE PATENTLER:

- Anıl, M. A. (2013, April). A Josephson Parametric Amplifier (JPA) for the ADMX-HF Experiment. In APS Meeting Abstracts.
- Dağ, C. B., Anıl, M. A., & Serpengüzel, A. (2015). Meandering waveguide distributed feedback lightwave circuits. Journal of Lightwave Technology, 33(9), 1691-1702.
- Danacı, B., Anıl, M. A., & Erzan, A. (2014). Motif statistics of artificially evolved and biological networks. Physical Review E, 89(6), 062719.
- Shokair, T. M., Root, J., Van Bibber, K. A., Brubaker, B., Gurevich, Y. V., Cahn, S. B., ... & Reed, A. (2014). Future directions in the microwave cavity search for dark matter axions. International Journal of Modern Physics A, 29(19), 1443004.