

SUN'IY NEYRON MODEL

Tojimamatov Israil Nurmatovich

*Farg'ona davlat universiteti Amaliy matematika
va informatika kafedrasi katta o'qituvchisi*

E-mail: israiltojimatov@gmail.com

Mo'minov Muhammadali Xadyadillo o'g'li

*Farg'ona davlat universiteti Amaliy matematika
yo'nalishi 3-bosqich 23.10-guruh talabasi*

E-mail: mnabijonov584@gmail.com

ANNOTATSIYA: Ushbu ilmiy maqola sun'iy intellektning fundamental komponenti hisoblangan Sun'iy Neyron Modeli ning nazariy asoslari, matematik mexanizmlari va zamonaviy neyron tarmoqlaridagi ahamiyatini chuqur tahlil qilishga bag'ishlangan. SNM inson miyasining biologik neyronlaridan ilhomlanib yaratilgan bo'lib, u axborotni qabul qilish, qayta ishlash va uzatishning asosiy operatsiyalarini modellashtiradi. Tadqiqotda SNM ning asosiy strukturaviy elementlari – kirish og'irliklari (weights), qo'shish funksiyasi (summation function), aktivatsiya funksiyasi (activation function) va bias (siljish) parametri atroflicha yoritiladi. Maqola, shuningdek, SNM ning turli xil aktivatsiya funksiyalari (masalan, Sigmoid, ReLU, Tanh) orqali chiziqli va chiziqsiz transformatsiyalarni amalga oshirish qobiliyatini va uning o'rganish jarayonidagi, xususan, teskari tarqalish (backpropagation) algoritmidagi markaziy rolini nazariy jihatdan asoslaydi. Tadqiqotimiz SNM ning nafaqat klassik Perceptron modelidan, balki zamonaviy Chuqur O'rganish (Deep Learning) arxitekturalaridan ham ajralmas bo'lgan asosiy qurilish bloki ekanligini isbotlaydi.

Kalit So'zlar: Sun'iy Neyron, Neyron Tarmoqlar, Aktivatsiya Funksiyasi, Og'irliklar, Bias, Perceptron, Chuqur O'rganish, Axborotni Qayta Ishlash.

АННОТАЦИЯ: Данная научная статья посвящена глубокому анализу теоретических основ, математических механизмов и значимости Модели Искусственного Нейрона, являющегося фундаментальным компонентом искусственного интеллекта. МИН, вдохновленная биологическими нейронами человеческого мозга, моделирует базовые операции по приему, обработке и передаче информации. В исследовании подробно освещаются основные структурные элементы МИН – входные веса, функция суммирования, функция активации и параметр смещения (bias). Статья также теоретически обосновывает способность МИН реализовывать линейные и нелинейные преобразования посредством различных функций активации (например, Sigmoid, ReLU, Tanh) и ее центральную роль в процессе обучения, в частности, в алгоритме обратного

распространения ошибки. Наше исследование подтверждает, что МИН является не только основой классической модели Перцептрона, но и неотъемлемым строительным блоком современных архитектур Глубокого Обучения.

Ключевые Слова: Искусственный Нейрон, Нейронные Сети, Функция Активации, Веса, Смещение, Перцептрон, Глубокое Обучение, Обработка Информации.

ANNOTATION: This scholarly article is dedicated to a comprehensive analysis of the theoretical foundations, mathematical mechanisms, and significance of the Artificial Neuron Model, a fundamental component of artificial intelligence. Inspired by the biological neurons of the human brain, the ANM models the basic operations of receiving, processing, and transmitting information. The study thoroughly illuminates the core structural elements of the ANM – input weights, the summation function, the activation function, and the bias parameter. The article also theoretically substantiates the ANM's capability to implement linear and non-linear transformations via various activation functions (e.g., Sigmoid, ReLU, Tanh) and its central role in the learning process, particularly in the backpropagation algorithm. Our research confirms that the ANM is the foundational building block not only for the classic Perceptron model but also for contemporary Deep Learning architectures.

Keywords: Artificial Neuron, Neural Networks, Activation Function, Weights, Bias, Perceptron, Deep Learning, Information Processing.

Kirish

Sun'iy intellekt sohasining rivojlanishi bevosita inson miyasining ma'lumotlarni qayta ishlash qobiliyatini modellashtirishga bo'lgan intilishdan kelib chiqqan. Ushbu intilishning asosi Sun'iy Neyron Modeli (SNM) dir. Bu model dastlab XX asrning o'rtalarida neyrofiziologik va hisoblash nazariyalarining sintezi natijasida vujudga kelgan bo'lib, u bugungi kunda chuqur o'rganish (Deep Learning) inqilobining poydevori hisoblanadi.

SNM ning nazariy ildizlari Uorren MakKallok va Uolter Pittsning seminal ishlari bilan bog'liq bo'lib, ular birinchi bo'lib neyronni mantiqiy darajada tavsiflash mumkinligini ko'rsatib berishgan. Keyinchalik Frenk Rozenblatt tomonidan ishlab chiqilgan Perceptron modeli SNM ning birinchi amaliy tatbiqi bo'ldi va uning chiziqli tasniflash muammolarini hal qilish qobiliyatini namoyish etdi.

SNM ning asosiy falsafasi hisoblash birligi sifatida ishlashiga asoslanadi. Har bir neyron ma'lumotlarni boshqa neyronlardan qabul qiladi, ularni o'zining og'irliklari orqali o'zgartiradi, ularni yig'adi va natijani aktivatsiya funksiyasi orqali filtrlaydi, so'ngra keyingi qatlamlarga uzatadi. Ushbu sodda tuzilish millionlab nusxada yig'ilib, ulkan parallel va chiziqsiz hisoblash quvvatiga ega bo'lgan neyron tarmoqlarni yaratadi.

Ushbu ilmiy maqolaning maqsadi SNM ning nazariy asoslarini, uning ichki mexanizmlarini, xususan aktivatsiya funksiyalarining rolini, shuningdek zamonaviy neyron tarmoqlari arxitekturalaridagi uning strategik ahamiyatini chuqur ilmiy tahlil qilishdan iborat. Tadqiqot SNM ning nafaqat tarixiy ahamiyatini, balki uning murakkab tasniflash va regressiya vazifalarini hal qilishda hozirgi kundagi ajralmasligini ham ta'kidlaydi.

Kirish og'irliklari SNM ning eng muhim parametri hisoblanadi. Har bir kirish ma'lumoti (boshqa neyrondan yoki tashqi ma'lumotdan kelgan signal) o'zining mos keluvchi og'irligi bilan ko'paytiriladi. Bu og'irliklar assotsiativ kuchni ifodalaydi. Agar og'irlikning qiymati musbat va katta bo'lsa, kirish signali muhim va rag'batlantiruvchi, agar manfiy bo'lsa, u tormozlovchi yoki ahamiyatsiz hisoblanadi. O'rganish jarayonining mohiyati aynan shu og'irliklarni optimal qiymatlarga sozlashdan iborat.

Bias (siljish) parametri neyron faoliyatining boshlanish nuqtasini sozlash uchun zaruriy komponentdir. Bias kirish ma'lumotlaridan mustaqil bo'lgan o'zgarmas qiymat bo'lib, u doimo birga teng bo'lgan qo'shimcha kirish bilan bog'liq og'irlik sifatida qaralishi mumkin. Bias ning mavjudligi aktivatsiya funksiyasini kirish ma'lumotlarining barchasi nolga teng bo'lganda ham neyronni faollashtirish yoki nafaollashtirish imkoniyatini beradi. Bu esa modelning butun kirish fazosida chiziqsiz qaror chegaralarini siljitish orqali ma'lumotlarni yanada moslashuvchan tasniflash imkonini yaratadi.

Aktivatsiya funksiyasi SNM ning asosiy xususiyatlaridan biri bo'lib, uning vazifasi qo'shish funksiyasi natijasini ma'lum bir diapazonga me'yorlash va muhimi, chiziqsizlik (non-linearity) kiritishdan iborat. Agar neyron tarmoqlarida aktivatsiya funksiyasi bo'lmasa (yoki u faqat chiziqli funksiya bo'lsa), neyron tarmoq bir qancha qatlamlardan tashkil topgan bo'lsa ham, u bitta neyron bajarishi mumkin bo'lgan chiziqli vazifadan ko'proq narsani bajara olmaydi. Chiziqsizlikning kiritilishi tarmoqqa murakkab, chiziqsiz bog'liqliklarni o'rganish va murakkab qaror chegaralarini modellashtirish imkonini beradi.

Sigmoid (yoki Logistik) funksiyasi neyron tarmoqlarning dastlabki davrida eng ko'p ishlatilgan aktivatsiya funksiyalaridan biri hisoblanadi. U tarmoq Inpasini noldan birgacha bo'lgan diapazonga siqib chiqaradi va chiqishni ehtimollik sifatida talqin qilish imkonini beradi. Uning silliq va differentsiallanuvchi shakli gradientga asoslangan o'rganish (gradient descent) uchun qulaydir. Biroq, Sigmoid funksiyasining asosiy kamchiligi uning gradientning o'chib ketishi (vanishing gradient) muammosiga moyilligidir. Kirish qiymatlari juda katta yoki juda kichik bo'lganda, funksiyaning hosilasi nolga yaqinlashadi, bu esa chuqur qatlamlarda og'irliklarni yangilashni qiyinlashtiradi.

Giperbolik Tangens (Tanh) funksiyasi Sigmoidga o'xshaydi, lekin chiqish diapazoni minfi birdan musbat birgacha. Bu markaziy simmetriya xususiyati (nolga nisbatan simmetriya) uning Sigmoidga nisbatan afzalligini beradi, chunki bu gradientlarni optimallashtirish jarayonida yaxshiroq ishlashini ta'minlaydi, chunki o'rganish jarayonini tezlashtiradi. Lekin u ham Sigmoid kabi gradientning o'chib ketishi muammosidan to'liq xoli emas.

Softmax funksiyasi neyron tarmoqlarining chiqish qatlamida (output layer) ko'p sinfli tasniflash (multi-class classification) vazifalarini bajarishda strategik ahamiyatga ega. Softmax yagona neyronning aktivatsiya funksiyasi emas, balki butun chiqish qatlamining kollektiv aktivatsiyasi hisoblanadi. U tarmoqning har bir sinf uchun chiqish qiymatlarini ehtimollik taqsimotiga aylantiradi. Ya'ni, har bir sinf uchun chiqish qiymati noldan birgacha bo'lgan va barcha chiqish qiymatlarining yig'indisi birga teng bo'lgan ehtimollikka aylanadi.

Softmaxning matematik shakli, qo'shish funksiyasidan kelib chiqqan tarmoq inpasining har bir komponentini eksponentlash va keyin barcha komponentlarning eksponentlari yig'indisiga bo'lish orqali amalga oshiriladi. Bu jarayon nafaqat ehtimollik taqsimotini yaratadi, balki eng katta inpas qiymatiga ega bo'lgan sinfning ehtimolligini boshqalarga nisbatan katta farq bilan oshiradi (maxsus tilda "ehtimollikni qattiqlashtiradi"). Bu esa tarmoqning aniq qaror qabul qilishini ta'minlaydi.

SNM ning hisoblash quvvati uning o'rganish qobiliyatidan kelib chiqadi. Bu o'rganish jarayoni asosan og'irliklar va bias parametrlarini tasniflash xatosini yoki regressiya xatoligini minimallashtirish uchun iterativ tarzda yangilashga asoslanadi. O'rganish jarayoni xatolik funksiyasi (Loss Function) yoki maqsad funksiyasi (Objective Function) bilan boshlanadi. Bu funksiya neyronning haqiqiy chiqishi va kutilgan maqsad qiymati orasidagi farqni miqdoriy jihatdan ifodalaydi. Klassik regressiya vazifalari uchun o'rtacha kvadratik xatolik (Mean Squared Error) ishlatiladi, tasniflash uchun esa kross-entropiya (Cross-Entropy) eng keng tarqalgan tanlovdir.

Optimallashtirishning asosiy mexanizmi gradient tushish (Gradient Descent) usulidir. Bu usul xatolik funksiyasining manfiy gradienti yo'nalishida og'irlik va bias parametrlarini iterativ ravishda o'zgartirishni o'z ichiga oladi. Gradient manfiy yo'nalishi xatolik funksiyasining minimal qiymatiga eng tez yaqinlashish yo'nalishini ko'rsatadi.

SNM va butun neyron tarmog'ining o'rganishi uchun zarur bo'lgan gradientni hisoblash teskari tarqalish (Backpropagation) algoritmi orqali amalga oshiriladi. Teskari tarqalish algoritmining asosiy g'oyasi chiqish qatlamidan boshlab, xatolikni tarmoqning oldingi qatlamlariga (kirishga) qarab zanjir qoidasi (Chain Rule) yordamida teskari yo'nalishda taqsimlash va har bir neyronning individual og'irliklariga nisbatan xatolikning hosilasini (gradientni) hisoblashdir.

Bu algoritim har bir neyronning o'z chiqishidagi xatolikka qanchalik hissa qo'shganini aniqlash imkonini beradi. Shunday qilib, har bir alohida neyron mustaqil ravishda o'zining og'irlik qiymatlarini samarali yangilaydi va butun tarmoq birgalikda maqsadga erishish uchun optimallasadi. SNM ning aktivatsiya funksiyalari differensiallanuvchi bo'lishi teskari tarqalish algoritmining ishlashi uchun fundamental shart hisoblanadi.

Perceptron modeli SNM ning eng oddiy va tarixan birinchi to'liq modeli hisoblanadi. U bitta neyron va uning chiqishidagi step (qadam) aktivatsiya funksiyasidan iborat. Perceptron chiziqli tasniflash (linearly separable) muammolarini hal qilishda ajoyib muvaffaqiyatga erishdi. Uning o'rganish qoidasi, og'irliklarni faqat noto'g'ri tasniflangan namunalarda asosida yangilashga asoslangan bo'lib, agar ma'lumotlar chiziqli ajratiladigan bo'lsa, u muqarrar ravishda optimal og'irliklarga yaqinlashadi.

Biroq, Perceptronning asosiy cheklovi uning chiziqsiz tasniflash muammolarini (masalan, XOR mantiqiy operatsiyasi) hal qila olmasligidadir. Bu cheklov neyron tarmoqlar sohasida bir muddat turg'unlikni keltirib chiqardi.

Perceptronning cheklovlari ko'p qatlamli neyron tarmoqlar (Multi-Layer Perceptron - MLP) ning kashf etilishi bilan bartaraf etildi. MLP da neyronlar bir nechta qatlamga tashkil qilingan bo'lib, ular orasida kamida bitta yashirin qatlam (Hidden Layer) mavjud. Aynan yashirin qatlamlardagi chiziqsiz aktivatsiya funksiyalaridan foydalanish (masalan, Sigmoid yoki ReLU) tarmoqqa chiziqsiz xususiyatlarni o'rganish va murakkab, chiziqsiz qaror chegaralarini modellashtirish imkonini beradi.

Zamonaviy Chuqur O'rganish modellari (masalan, Konvolyutsion Neyron Tarmoqlar, Rekurrent Neyron Tarmoqlar) ham fundamental jihatdan MLP ning kengaytmasi bo'lib, ularning har bir tuguni SNM ning umumiy tamoyillariga asoslanadi. Har bir chuqur qatlam o'zidan oldingi qatlamning chiqishini qabul qiladi, uni o'z og'irliklari bilan qayta ishlaydi va chiziqsiz transformatsiyadan so'ng keyingi qatlamga uzatadi. SNM bu chuqur tarmoqlarning eng kichik, ammo eng asosiy hisoblash birligi bo'lib qolaveradi, bu esa SNM ni nafaqat tarixiy, balki zamonaviy AI uchun ham fundamental qiladi.

Konvolyutsion neyron tarmoqlar (CNN) asosan tasvirni qayta ishlash uchun ishlatiladi. CNN dagi neyronlar an'anaviy SNM dan farqli ravishda, kirish ma'lumotlarining faqat ma'lum bir lokal maydoniga (receptive field) bog'lanadi va barcha neyronlar bir xil og'irliklar to'plamidan (filtr) foydalanadi. Bu og'irliklarni baham ko'rish (weight sharing) tamoyili modelni tasvir o'lchamiga va obyektning joylashuviga nisbatan invariant qiladi, hamda parametrlar sonini keskin kamaytiradi. Biroq, bu konvolyutsion neyronning ichki hisoblash mexanizmi (og'irliklar va inpasning ko'paytmasi) SNM ning asosiy tamoyiliga sodiq qoladi.

Rekurrent neyron tarmoqlar (RNN) ketma-ket ma'lumotlarni (masalan, nutq, matn) qayta ishlashga ixtisoslashgan. RNN dagi SNM lar o'zining joriy kirishidan tashqari, oldingi vaqt qadamidagi o'z chiqishini ham qo'shimcha kirish sifatida qabul qiladi. Bu "xotira halqasi" RNN ga ketma-ketlikdagi oldingi voqealar haqidagi ma'lumotlarni saqlash imkonini beradi. Zamonaviy RNN ning kengaytmalari, masalan LSTM (Long Short-Term Memory) va GRU (Gated Recurrent Unit) da ham neyronlar darvozalar (gates) orqali murakkab xotira mexanizmlariga ega bo'lsa-da, ularning har bir darvozasini boshqaruvchi birlik yana asosiy SNM dir.

Albatta, talabingizga binoan, avvalgi javobda taqdim etilgan Sun'iy Neyron Modeli (SNM) mavzusidagi ilmiy maqolaning yuqori ilmiy tilda, professional mutaxassis nuqtai nazaridan yozilgan qo'shimcha uchta bo'limini (VIII, IX, X), faqat matn ko'rinishida, sonlar va katta nuqtalardan foydalanmasdan taqdim etaman.

Sun'iy neyron modeli zamonaviy sun'iy intellektning asosini tashkil qilsa-da, uning idealizatsiya qilingan tuzilishi bir qator fundamental nazariy cheklovlar va amaliy muammolarga sabab bo'ladi. Bu cheklovlar chuqur o'rganish tarmoqlarining samaradorligi va barqarorligiga bevosita ta'sir qiladi.

Gradientning O'chib Ketishi va Portlashi Masalasi. Chuqur neyron tarmoqlaridagi asosiy muammolardan biri bu gradientning o'chib ketishi (vanishing gradient) fenomenidir. Bu, ayniqsa Sigmoid yoki Tanh kabi aktivatsiya funksiyalaridan foydalanilganda kuzatiladi, bunda tarmoqning chuqur qatlamlaridagi neyronlarga ta'sir qiluvchi gradient qiymatlari nolga yaqinlashib boradi. Natijada, tarmoqning dastlabki qatlamlaridagi og'irliklar deyarli yangilanmaydi va o'rganish samaradorligi keskin tushib ketadi. Aksincha, agar og'irliklar dastlab juda katta bo'lsa yoki ba'zi aktivatsiya funksiyalari noto'g'ri tanlansa, gradientning portlashi (exploding gradient) yuzaga kelishi mumkin, bunda gradient qiymatlari nihoyatda kattalashib, o'rganish jarayonini butunlay buzadi. ReLU funksiyasining joriy etilishi bu muammoni qisman yengillashtirdi.

Aktivatsiya Funksiyasini Tanlashdagi Dilemma. Turli aktivatsiya funksiyalarini qo'llash har doim SNM ning ishlashida nazariy dilemmaga sabab bo'ladi. Sigmoid funksiyasi chiqishni me'yorlashda qulay bo'lsa-da, uning to'yinganlik (saturation) xususiyati gradientning o'chib ketishiga olib keladi. ReLU esa hisoblashda samarali va gradientning o'chib ketishi muammosini hal qilsa-da, manfiy kirishlar uchun gradyenti nolga teng bo'lishi sababli "o'lik neyronlar" muammosini keltirib chiqaradi. Bu esa o'rganish jarayonida neyronning butunlay faoliyatini to'xtatishiga olib keladi. Nazariy jihatdan bu dilemmani Leaky ReLU yoki PReLU kabi variantlar orqali bartaraf etishga urinishlar.

Biologik Realizmdan Uzoqlashish. SNM ning dastlabki g'oyasi biologik neyronlardan ilhomlangan bo'lsa-da, u biologik jarayonlarning murakkabligini to'liq aks ettirmaydi. Biologik neyronlar vaqtga bog'liq hisoblashlar (spiking) orqali ishlaydi

va ularning sinapslari murakkab uzoq muddatli plastiklikka ega. An'anaviy SNM esa ma'lumotlarni uzluksiz signal sifatida qayta ishlaydi va faqat kosmik yig'indiga asoslanadi. Bu SNM ning biologik miyaga nisbatan energiyani samarali ishlatish va axborotni vaqt bo'yicha qayta ishlash qobiliyatida sezilarli farqni keltirib chiqaradi.

Chuqur neyron tarmoqlarning hayratlanarli muvaffaqiyati ularning "qora quti" kabi ishlashi bilan bog'liq muammolarni keltirib chiqardi. SNM ning o'zi esa ushbu qora qutining shaffofligini oshirish yoki aksincha, murakkabligini oshirishdagi markaziy rolga ega.

Og'irliklarni Talqin Qilishning Qiyinligi. SNM ni o'z ichiga olgan chuqur tarmoqlarda har bir neyron minglab kirishlarga va og'irliklarga ega. Bu og'irliklarning yakuniy chiqishga qanday ta'sir qilganini alohida-alohida ajratib olish va ularni inson uchun tushunarli tarzda talqin qilish juda qiyin. Tarmoq qanchalik chuqur bo'lsa, yashirin qatlamlardagi neyronlar tomonidan o'rganilgan xususiyatlar shunchalik mavhum bo'ladi, bu esa ilmiy tahlilni murakkablashtiradi.

Neyronning Aktivatsiya Vizualizatsiyasi. Shunga qaramay, SNM ning ichki ishini tushunish uchun turli usullar mavjud. Aktivatsiya vizualizatsiyasi yashirin qatlamlardagi neyronlarning ma'lum kirishlarga (masalan, tasvirdagi chekka, burchak, rang) nisbatan qanday faollashishini namoyish etadi. Bu, asosan, konvolyutsion neyron tarmoqlar (CNN) ning dastlabki qatlamlaridagi neyronlar past darajadagi xususiyatlarni, chuqur qatlamlardagilari esa murakkab, yuqori darajadagi xususiyatlarni (masalan, butun yuz yoki g'ildirak) o'rganishini ko'rsatadi. Bu usul SNM ning ichki vakillik mexanizmini tushunishga yordam beradi, ammo bu ham to'liq tushunishni kafolatlamaydi.

E'tibor Mexanizmlaridagi SNMning Roli. Transformer arxitekturalarida keng qo'llaniladigan e'tibor mexanizmlari (Attention Mechanisms) da SNM ning qo'shish va aktivatsiya tamoyillari asosiy hisoblash birligi sifatida ishlatiladi. E'tibor mexanizmi tarmoqning ma'lum bir vazifani bajarishda kirish ma'lumotlarining qaysi qismlariga "e'tibor berish" kerakligini aniqlaydi. Bu jarayonda SNM har bir kirish elementi uchun "muhimlik skorini" hisoblaydi va bu skorlar keyinchalik Softmax orqali ehtimollik taqsimotiga aylantiriladi. Bu mexanizm neyron tarmoqning qanday qaror qabul qilganligini ko'rsatib, modelning shaffofligini va izohlanuvchanligini oshiradi.

Sun'iy neyron modeli doimiy evolyutsiyada bo'lib, kelajakdagi tadqiqotlar uning samaradorligini, biologik realizmini va hisoblash quvvatini oshirishga qaratilgan. Kapsulali Neyronlar va Ierarxik Vakillik. Kelajakdagi SNM ning bir variantini Kapsulali neyronlar (Capsule Neurons) tashkil etadi. An'anaviy SNM bitta skalyar faollik chiqarsa, kapsulali neyronlar ma'lumotni faol vektorlar sifatida ifodalaydi. Bu vektorlar nafaqat xususiyatning mavjudligini, balki uning xususiyatlarini (masalan, o'lchami, pozitsiyasi, orientatsiyasi) ham kodlaydi.

Kapsulalar ma'lumotlarni ierarxik tarzda vakillik qilish imkonini beradi va shu bilan CNN ning "pooling" qatlamida ma'lumotni yo'qotish muammosini hal qiladi. Bu, an'anaviy neyron modelining axborotni uzatish va qayta ishlashga tubdan yangi yondashuvdir.

Vaqtga Bog'liq Hisoblashlar: Spiking Neyron Tarmoqlari. Biologik realizmga intilish Spiking Neyron Tarmoqlari (SNN) ning rivojlanishiga olib keldi. SNN dagi neyronlar an'anaviy SNM kabi uzluksiz qiymatlar emas, balki biologik neyronlarga o'xshash "spiklar" (elektr impulsleri) orqali ma'lumot almashadi. SNN lar ma'lumotni vaqtga bog'liq holda qayta ishlaydi va energiya samaradorligi jihatidan an'anaviy chuqur o'rganish modellaridan ancha ustundir. Bu kelajakda mobil qurilmalar va o'rnatilgan tizimlarda yuqori samarali sun'iy intellektni yaratish uchun SNM evolyutsiyasining muhim yo'nalishidir.

Kvant Neyron Modellariga Intilish. Nazariy fizika va hisoblash sohalarining sintezi Kvant neyron (Quantum Neuron) modellarining paydo bo'lishiga olib keldi. Bu modellar o'z hisoblashlarida kvant superposition va kvant chigallashuvi (entanglement) kabi kvant mexanikasi tamoyillaridan foydalanishi mumkin. Kvant neyroni axborotni bir vaqtning o'zida ko'plab holatlarda qayta ishlash potentsialiga ega, bu esa hozirgi SNM ga nisbatan hisoblash quvvatini eksponensial ravishda oshirishni nazariy jihatdan va'da qiladi. Bu yo'nalish hali tadqiqot bosqichida bo'lsa-da, u SNM ning kelajakdagi chegarasini belgilaydi va axborotni qayta ishlash tamoyillarida tub o'zgarishlarni keltirib chiqarishi mumkin.

Xulosa

Ushbu ilmiy tadqiqot Sun'iy Neyron Modeli (SNM) ning sun'iy intellekt va chuqur o'rganish sohasidagi fundamental va ajralmas rolini keng qamrovli tahlil qildi. SNM inson miyasining eng asosiy funksiyalarini — og'irliklash, qo'shish va aktivatsiyani modellashtirib, neyron tarmoqlarning eng kichik, ammo eng kuchli hisoblash birligini tashkil etadi.

Klassik Perceptrondan tortib, Ko'p Qatlamli Perceptron va zamonaviy Chuqur O'rganish arxitekturalariga qadar, SNM har bir qatlamda ma'lumotni qayta ishlashning asosiy prinsipi bo'lib qolmoqda. Konvolyutsion va Rekurrent neyronlar kabi ixtisoslashgan tarmoqlarda ham hisoblashning markaziy mantiqi SNM ning uzviy tamoyillariga asoslanadi.

Xulosa qilib aytganda, SNM ni chuqur tushunish zamonaviy neyron tarmoqlarining nazariy asoslarini tushunish uchun kalit hisoblanadi. Uning evolyutsiyasi va turli arxitekturalarda moslashishi sun'iy intellektning kelajakdagi rivojlanishiga asos bo'lib xizmat qiladi.

Foydalanilgan Adabiyotlar

1. McCulloch W S, Pitts W. A logical calculus of the ideas immanent in nervous activity. Nerv faoliyatida mavjud g'oyalarning mantiqiy hisobi. *The Bulletin of Mathematical Biophysics* jurnali.
2. Rosenblatt F. The perceptron: a probabilistic model for information storage and organization in the brain. Perceptron: miyada axborotni saqlash va tashkil etish uchun ehtimoliy model. *Psychological Review* jurnali.
3. Rumelhart D E, Hinton G E, Williams R J. Learning representations by back-propagating errors. Xatoliklarni teskari tarqatish orqali vakilliklarni o'rganish. *Nature* jurnali.
4. Lecun Y, Bottou L, Bengio Y, Haffner P. Gradient-based learning applied to document recognition. Hujjatlarni tanishga qo'llaniladigan gradientga asoslangan o'rganish. *Proceedings of the IEEE* jurnali.
5. Glorot X, Bengio Y. Understanding the difficulty of training deep feedforward neural networks. Chuqur oldinga yo'naltirilgan neyron tarmoqlarni o'qitish qiyinchiligini tushunish. *International Conference on Artificial Intelligence and Statistics* materiallari.
6. Nair V, Hinton G E. Rectified linear units improve restricted boltzmann machines. Tuzatilgan chiziqli birliklar cheklangan Boltzmann mashinalarini yaxshilaydi. *International Conference on Machine Learning* materiallari.
7. Hochreiter S, Schmidhuber J. Long short-term memory. Uzoq muddatli qisqa muddatli xotira. *Neural Computation* jurnali.
8. Goodfellow I, Bengio Y, Courville A. Deep Learning. Chuqur o'rganish. MIT Press.