



**Oftalmologik tasvirlarni tahlil qiluvchi veyvletlarga asoslangan neyron tarmoq
arxitekturasini tasvirlarni raqamli ishlash orqali takomillashtirish**

Raximov Baxtiyor Saidovich

Toshkent tibbiyot akademiyasi Urganch filiali

“Biofizika va axborot texnologiyalari” kafedrası mudiri

O‘razmatov Tohir Quronbayevich

Muhammad al-Xorazmiy nomidagi TATU Urganch

filiali “Axborot texnologiyalari” kafedrası katta o‘qituvchisi

tohir20314@gmail.com

Annotatsiya: Maqolada oftalmologik tasvirlarni tahlil qilishda raqamli ishlash texnologiyalari qo'llanilishi va EfficientNetB0 algoritmining afzalliklari tahlil qilinadi. EfficientNetB0 algoritmi chuqurlik, kenglik va rezolyutsiya kabi parametrlarni samarali balanslash orqali neyron tarmoqning ishlash unumdorligini oshiradi. Shuningdek, Squeeze-and-Excite (SE) bloki va Swish aktivatsiya funksiyasi kabi usullar yordamida tarmoq resurslarini tejash va tasvirlardan xususiyatlarni aniq ajratish imkoniyati oshiriladi.

Kalit so‘zlar: Oftalmologiya, raqamli tasvir ishlash, neyron tarmoq, EfficientNetB0, konvolyutsion neyron tarmoq, Swish aktivatsiya funksiyasi, Squeeze-and-Excite mexanizmi.



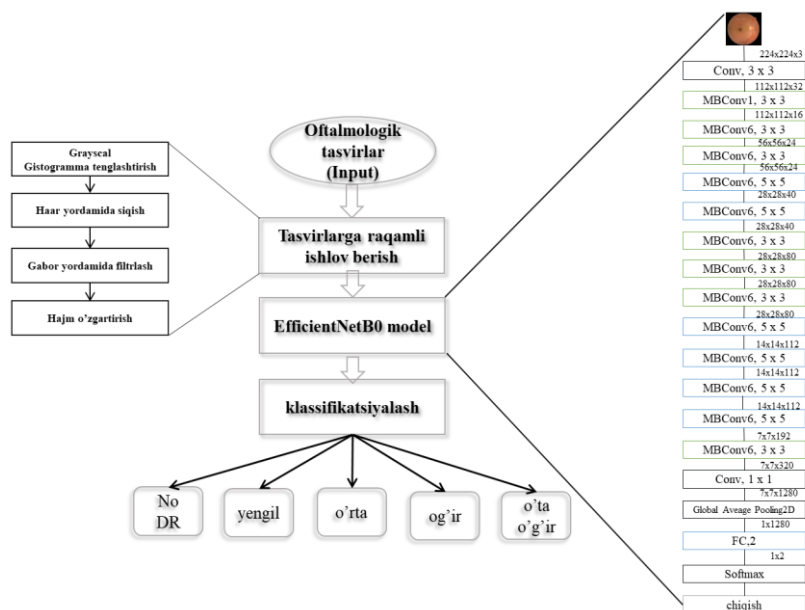
EfficientNetB0 algoritmi uchta asosiy parametрни miqyoslaydi: chuqurlik (depth), kenglik (width) va rezolyutsiya (resolution). Bu uch parametr o‘zaro balanslashgan holda, tarmoq samaradorligi oshiriladi:

- **Chuqurlik miqyosi (depth scaling):** Tarmoqning chuqurligini oshirish orqali modelning aniqligi ko‘payadi, ya’ni modelda konvolyutsion qatlamlar soni ortadi.
- **Kenglik miqyosi (width scaling):** Tarmoqning kengligini oshirish, ya’ni qatlamlardagi neyronlar sonini ko‘paytirish.
- **Rezolyutsiyani miqyoslash (resolution scaling):** Kiritilayotgan tasvirning rezolyutsiyasini oshirish orqali yuqori aniqlikda ishlashga imkon beradi.

Keyingi bosqichda MBConv Layer (Mobile Inverted Bottleneck Conv Layer) amalga oshiriladi. Bu qatlam modelni yengil va samarali qilish uchun mobil inverli bottle-neck strukturasidan foydalanadi[137; 778-b., 142; 10431-b.]. Bu yerda:

- **Spress (Squeeze) va Excite mexanizmi:** Kanalning hajmini qisqartirish va qayta kengaytirish orqali parametrlarni tejamkorlik bilan ishlatadi.
- **Depthwise Separable Convolution:** Standart konvolyutsion qatlamlardan farqli o‘laroq, har bir kanal alohida ishlov beriladi, bu esa resurslarni tejaydi.

Quyida EfficientNetB0 CNN ning ishlash sxemasi keltirilgan. Bu yerda har bir qatlamning tasvir va filtr o‘lchamlari bilan birga yozilgan.



1 rasm. EfficientNetB0 arxitekturası

EfficientNetB0 da **Swish** aktivatsiya funksiyasidan foydalaniladi. Bu funktsiya odatda ReLU dan yaxshiroq ishlaydi va uning ifodasi quyidagicha:

$$\text{Swish}(x) = x * \frac{1}{1 + e^{-x}} \quad (1)$$

Bu funksiyaning afzalligi shundaki, gradientlar kamroq kesiladi va model optimizatsiya jarayonida yaxshi o'rgatiladi.

Yuqorida ko'rsatilgandek kiruvchi tasvirlar modelga sozlangandan keyin model arxitekturasida bo'yicha qatlamlardagi jarayonlar sodir bo'ladi. Birinchi Conv qatlam- bu konvolyutsion neyron tarmoqlarda (CNN) tasvirlar va boshqa ma'lumotlar bilan ishlash berishda eng muhim qism hisoblanadi. Conv qatlam tasvirdan xususiyatlarni ajratib olish uchun ishlatiladi va bu qatlam tasvirda aniqlik va semantik jihatdan muhim bo'lgan xususiyatlarni tanlab oladi. Conv qatlam tasvirga **filtrlar** yoki **kernel** (yadro) yordamida ishlash beradi. Har bir filtr kichik matritsadan iborat bo'lib, tasvirning kichik qismi bilan ko'paytirilib, xususiyat xaritalari (feature map) hosil qiladi. Buning natijasida tasvirning turli tomonlari (chegaralar, burchaklar, naqshlar) ajratib olinadi.



Conv qatlamning asosiy komponenti **kernel** bo'lib, u kichik matritsa shaklida bo'ladi (odatda 3x3, 5x5 yoki 7x7). Har bir filtr tasvirning kichik qismi ustida ishlov beradi. Kernel har bir tasvirning bir qismini qamrab oladi va matematik operatsiyalarni bajaradi, bu yerda asosan **skalyar ko'paytma** va **summatsiya** bajariladi.

Kernel tasvir bo'ylab harakatlanadi va bu harakat **stride** deb ataladi. Stride bu kernel qancha bosqichda tasvir bo'yicha qadam tashlashini bildiradi. Misol uchun, agar stride = 1 bo'lsa, kernel har bir pozitsiyada 1 ta pikselni tashlab o'tadi, agar stride = 2 bo'lsa, 2 pikseldan keyin harakatlanadi. Stride katta bo'lsa, chiqish o'lchami kichikroq bo'ladi.

Agar tasvirning chegaralarida ma'lumot yo'q bo'lsa, ba'zi hollarda **padding** ishlatiladi. Padding bu tasvir atrofida 0 qiymatli elementlarni qo'shib tasvirning chegaralarini to'ldirishga imkon beradi. Bu yerda asosiy maqsad tasvirning o'lchamlarini saqlab qolishdir.

Conv qatlamda har bir filtr o'zining ma'lum bir **kernelini** tasvir bo'yicha "sirg'alib" harakatlantiradi va tasvirning ustiga qo'yilgan qismi bilan **dot product (nuqta ko'paytma)** amalga oshiradi. Bu operatsiya quyidagicha:

$$Z(i, j) = (W \cdot X + b) \quad (2)$$

Bu yerda:

W-filtr(kernel) vaznlar to'plami;

X-kiruvchi tasvirning qismi;

b-bias parametric;

Z(i, j)- hosil bo'lgan xususiyat xaritasi.

Misol uchun: 5x5 o'lchamdagi tasvirga 3x3 kernel qo'llash jarayoni keltirilgan:

Kiritish tasviri:



$$\begin{bmatrix} 1 & 1 & 1 & 0 & 0 \\ 0 & 1 & 1 & 1 & 0 \\ 0 & 0 & 1 & 1 & 1 \\ 0 & 0 & 1 & 1 & 0 \\ 1 & 1 & 0 & 0 & 0 \end{bmatrix}$$

Kernel (filtr):

$$\begin{bmatrix} 1 & 0 & 1 \\ 0 & 1 & 0 \\ 1 & 0 & 1 \end{bmatrix}$$

Tasvirning birinchi 3x3 qismi bilan kernelning **dot product** hisoblanadi va summasi olinadi:

$$\begin{bmatrix} 1 & 1 & 1 \\ 0 & 1 & 1 \\ 0 & 0 & 1 \end{bmatrix}$$

Hisoblangan qiymatlar: $(1 \cdot 1) + (1 \cdot 0) + (1 \cdot 1) + (0 \cdot 0) + (1 \cdot 1) + (1 \cdot 0) + (0 \cdot 1) + (0 \cdot 0) + (1 \cdot 1) = 5$

Tasvirning butun qismi bo'yicha kernel harakatlanadi va natijada xarakteristikalar xaritasi hosil bo'ladi.

Keyingi qatlamda ya'ni MbVonv qatlamlarini umumiy ishlash prinsipi quyidagicha:

Expand bosqichi (Kengaytirish) Swish bilan. Bu bosqichda kiruvchi tensorning hajmi kengaytiriladi va Swish aktivatsiya funksiyasi qo'llaniladi.

$$X_{\text{exp}} = \text{Swish}(X_{\text{in}} \cdot W_{\text{exp}}) \quad (3)$$

Bu yerda:

X_{in} -kiruvchi tensor;

W_{exp} -kengaytiruvchi konvolyutsiya tortishishlari;



Swish-aktivatsiya funksiyasi.

Depthwise konvolyutsiya har bir kanal uchun alohida qo'llaniladi va Swish bilan aktivatsiyalanadi.

$$X_{\text{depth}} = \text{Swish}(\text{DepthConv}(X_{\text{exp}})) \quad (4)$$

DepthConv- depthwise konvolyutsiya amali;

Swish-aktivatsiya funksiyasi.

Squeeze bosqichi (Siqish) Swish bilan. Depthwise konvolyutsiyadan keyin Swish funksiyasi bilan siqish bosqichi amalga oshiriladi.

$$X_{\text{squeeze}} = \text{Swish}(X_{\text{depth}} \cdot W_{\text{squeeze}}) \quad (5)$$

Bu yerda:

W_{squeeze} -siqish uchun konvolyutsiya tortishishlari;

Swish-aktivatsiya funksiyasi.

Squeeze-and-Excitation (SE) bloki Swish bilan. Agar EfficientNet'da mavjud bo'lsa, SE blokida aktivatsiya funksiyasi ham Swishga o'zgartiriladi. SE blokining matematik ifodasi quyidagicha bo'ladi:

$$s = \sigma(W_2 \cdot \text{Swish}(W_1 \cdot \text{GlobalAvgPool}(X_{\text{depth}}))) \quad (6)$$

GlobalAvgPool- global o'rtacha pooling;

W_1 va W_2 -SE blokidagi tortishishlar;

s- kanallar bo'yicha tarqatilgan vazn koeffitsienti.

Chiqish esa quyidagicha modifikatsiyalanadi:

$$X_{\text{se}} = s \cdot X_{\text{depth}} \quad (7)$$



Keyingi bosqich Shortcut (Skip Connection) Swish bilan. Shortcut ham Swish aktivatsiya funksiyasi bilan birga ishlatiladi. Agar kiruvchi va chiqish o'lchamlari bir xil bo'lsa, shortcut bilan ma'lumotlar to'g'ridan-to'g'ri qo'shiladi[156; 966-b.].

$$X_{out} = X_{in} + \text{Swish}(X_{squeeze})\text{-agar o'lchamlar mos kelsa,}$$

aks holda:

$$X_{out} = \text{Swish}(X_{squeeze}) \quad (8)$$

Bu blokda Swish funksiyasi an'anaviy ReLU yoki ReLU6 ga qaraganda silliqroq chiqish beradi va gradientlar yo'qolishiga qarshi yaxshi himoya qiladi.

Xulosa

Oftalmologik tasvirlarni tahlil qilish uchun taklif etilgan EfficientNetB0 algoritimga asoslangan neyron tarmoq tasvirlarni samarali va yuqori aniqlik bilan qayta ishlash imkonini beradi. Algoritmning chuqurlik, kenglik va rezolyutsiyani miqyoslash imkoniyatlari, shuningdek, Squeeze-and-Excite mexanizmi va Swish aktivatsiya funksiyasi yordamida oftalmologik tasvirlarni tahlil qilish jarayonlari optimallashtiriladi. Bu usul resurslarni tejab, ma'lumotlarni aniq ajratishga xizmat qiladi va bu sohada yangi imkoniyatlar yaratadi.

Foydalanilgan adabiyotlar ro'yxati:

1. Tan, M., & Le, Q. (2019). *EfficientNet: Rethinking Model Scaling for Convolutional Neural Networks*. In Proceedings of the 36th International Conference on Machine Learning (pp. 6105-6114). ICML.
2. Hu, J., Shen, L., & Sun, G. (2018). *Squeeze-and-Excitation Networks*. IEEE Conference on Computer Vision and Pattern Recognition (CVPR), 7132-7141.



3. Krizhevsky, A., Sutskever, I., & Hinton, G. E. (2012). *ImageNet Classification with Deep Convolutional Neural Networks*. Advances in Neural Information Processing Systems, 25, 1097-1105.
4. LeCun, Y, Bengio, Y, & Hinton, G. (2015). *Deep Learning*. Nature, 521(7553), 436-444.
5. He K, Zhang X, Ren S, & Sun, J.(2016). *Deep Residual Learning for Image Recognition*. IEEE Conference on Computer Vision and Pattern Recognition (CVPR), 770-778
6. Howard, A. G., Zhu, M., Chen, B., Kalenichenko, D., Wang, W., Weyand, T., ... & Adam, H. (2017). *MobileNets: Efficient Convolutional Neural Networks for Mobile Vision Applications*. arXiv preprint arXiv:1704.04861.
7. Ioffe, S., & Szegedy, C. (2015). *Batch Normalization: Accelerating Deep Network Training by Reducing Internal Covariate Shift*. In Proceedings of the 32nd International Conference on Machine Learning (pp. 448-456). ICML.
8. Russakovsky, O., Deng, J., Su, H., Krause, J., Satheesh, S., Ma, S., ... & Fei-Fei, L. (2015). *ImageNet Large Scale Visual Recognition Challenge*. International Journal of Computer Vision, 115(3), 211-252.
9. Glorot, X., & Bengio, Y. (2010). *Understanding the Difficulty of Training Deep Feedforward Neural Networks*. In Proceedings of the Thirteenth International Conference on Artificial Intelligence and Statistics (AISTATS), 249-256.
10. Mishkin, D., & Matas, J. (2015). *All You Need is a Good Init*. arXiv preprint arXiv:1511.06422.