



POLİTEKNİK DERGİSİ

JOURNAL of POLYTECHNIC

ISSN: 1302-0900 (PRINT), ISSN: 2147-9429 (ONLINE)

URL: <http://dergipark.org.tr/politeknik>



RLSE aktivasyon fonksiyonu tasarımının derin sinir ağlarının performansındaki etkisi

The impact of RLSE activation function design on the performance of deep neural networks

Yazar(lar) (Author(s)): İsmihan Gül ÖZELOĞLU¹, Edd AKMAN AYDIN², Necaattin BARIŞÇI³

ORCID¹: 0000-0001-9469-9550

ORCID²: 0000-0002-9887-3808

ORCID³: 0000-0002-8762-5091

To cite to this article: Özeloğlu İ. G., Akman Aydın E. and Barışçı N., “RLSE Aktivasyon Fonksiyonu Tasarımının Derin Sinir Ağlarının Performansındaki Etkisi”, *Journal of Polytechnic*, *(*) : *, (*).

Bu makaleye şu şekilde atıfta bulunabilirsiniz: Özeloğlu İ. G., Akman Aydın E. ve Barışçı N., “RLSE Aktivasyon Fonksiyonu Tasarımının Derin Sinir Ağlarının Performansındaki Etkisi”, *Politeknik Dergisi*, *(*) : *, (*).

Erişim linki (To link to this article): <http://dergipark.org.tr/politeknik/archive>

DOI: 10.2339/politeknik.1601441

RLSE Aktivasyon Fonksiyonu Tasarımının Derin Sinir Ağlarının Performansındaki Etkisi

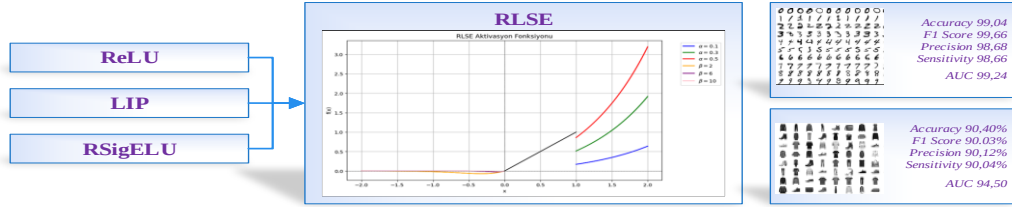
The Impact of RLSE Activation Function Design on the Performance of Deep Neural Networks

Önemli noktalar (Highlights)

- ❖ Yeni bir çift parametrelili RLSE aktivasyon fonksiyonu önerilmiştir. / A novel dual-parameter RLSE activation function is proposed.
- ❖ RLSE, kaybolan gradyan ve ölmekte olan ReLU problemlerini azaltmak üzere tasarlanmıştır. / RLSE is designed to mitigate the vanishing gradient and dying ReLU problems.
- ❖ RLSE, literatürdeki diğer aktivasyon fonksiyonlarından daha yüksek doğruluk oranı sağlamıştır. / RLSE achieved higher accuracy than other activation functions in the literature.

Grafik Özet (Graphical Abstract)

RLSE aktivasyon fonksiyonu, ReLU, LIP ve ELU bileşenlerini birleştirerek derin sinir ağlarında doğruluk ve öğrenme kararlılığını artırmaktadır. / The RLSE activation function combines ReLU, LIP, and ELU components to enhance accuracy and learning stability in deep neural networks.



Şekil X. RLSE aktivasyon fonksiyonu, ReLU, LIP ve ELU bileşenlerini birleştirmektedir./ **Figure.** The RLSE activation function combines ReLU, LIP, and ELU components.

Amaç (Aim)

Bu çalışmanın amacı, derin sinir ağlarında kaybolan gradyan ve negatif bölge problemlerini azaltmak için yeni bir RLSE aktivasyon fonksiyonu tasarlamaktır. / The aim of this study is to design a novel RLSE activation function to reduce vanishing gradient and negative region problems in deep neural networks.

Tasarım ve Yöntem (Design & Methodology)

Önerilen RLSE aktivasyon fonksiyonu, pozitif, doğrusal ve negatif bölgelerde farklı eğim parametreleri kullanılarak tasarlanmıştır. / The proposed RLSE activation function was designed using different slope parameters for positive, linear, and negative regions.

Özgünlük (Originality)

RLSE aktivasyon fonksiyonu, ReLU, LIP ve ELU fonksiyonlarının avantajlarını tek bir yapıda birleştirerek kaybolan gradyan ve ölmekte olan ReLU problemlerine çözüm getiren özgün bir yaklaşımdır. / The RLSE activation function offers an original approach that combines the advantages of ReLU, LIP, and ELU functions in a single structure, addressing vanishing gradient and dying ReLU problems.

Bulgular (Findings)

RLSE aktivasyon fonksiyonu, MNIST veri kümesinde %99,04 ve Fashion-MNIST veri kümesinde %90,40 doğruluk oranı elde ederek diğer aktivasyon fonksiyonlarından daha iyi performans göstermiştir. / The RLSE activation function achieved accuracy rates of 99.04% on the MNIST dataset and 90.40% on the Fashion-MNIST dataset, outperforming other activation functions.

Sonuç (Conclusion)

RLSE aktivasyon fonksiyonu, derin sinir ağlarının doğruluğunu ve kararlılığını artıran etkili bir yöntemdir. / The RLSE activation function is an effective method that enhances the accuracy and stability of deep neural networks.

Etik Standartların Beyanı (Declaration of Ethical Standards)

Bu makalenin yazarları çalışmalarında kullandıkları materyal ve yöntemlerin etik kurul izni ve/veya yasal-özel bir izin gerektirmediğini beyan ederler. / The authors of this article declare that the materials and methods used in this study do not require ethical committee permission and/or legal-special permission.

RLSE Aktivasyon Fonksiyonu Tasarımının Derin Sinir Ağlarının Performansındaki Etkisi

Araştırma Makalesi / Research Article

İsmihan Gül ÖZELOĞLU^{1*}, Eda AKMAN AYDIN², Necaattin BARIŞCI³

¹Fen Bilimleri Enstitüsü, Elektrik-Elektronik Müh. Bölümü, Gazi Üniversitesi, Türkiye

²Teknoloji Fakültesi, Elektrik Elektronik Müh. Bölümü, Gazi Üniversitesi, Türkiye

³Teknoloji Fakültesi, Bilgisayar Müh. Bölümü, Gazi Üniversitesi, Türkiye

(Geliş/Received : 01.01.2025 ; Kabul/Accepted : 26.09.2025 ; Erken Görünüm/Early View : 23.11.2025)

ÖZ

Aktivasyon fonksiyonu derin sinir ağlarının performansı üzerinde kritik etkisi olan bir bileşendir. Bu çalışmada, derin sinir ağlarında, yüksek sınıflandırma doğruluğu ve düşük kayıp elde etmek için yeni bir aktivasyon fonksiyonu önerilmektedir. Önerilen RLSE (ReLU-LIP-Sigmoid-ELU kombinasyonu) aktivasyon fonksiyonu ile, kaybolan gradyan sorunu ve ölmekte olan ReLU probleminin üstesinden gelinmesi hedeflenmektedir. RLSE aktivasyon fonksiyonunun performansı MNIST ve Fashion-MNIST veri kümeleri üzerinde değerlendirilmiş, literatürde bulunan yeni geliştirilmiş aktivasyon fonksiyonlarıyla ve günümüzde en çok kullanılan aktivasyon fonksiyonlarıyla karşılaştırılmıştır. RLSE aktivasyon fonksiyonunun kullanılması ile, bu çalışmada tasarlanan Evrişimsel Sinir Ağı (ESA) mimarisinde MNIST veri kümesi için %99,04 ve Fashion MNIST veri kümesi için %90,40 doğruluk oranları elde edilmiştir. Sonuçlar, RLSE aktivasyon fonksiyonunun diğer aktivasyon fonksiyonlarından daha iyi performans gösterdiğini ortaya koymaktadır.

Anahtar Kelimeler: Aktivasyon fonksiyonu, evrişimsel sinir ağı, derin sinir ağları.

The Impact of RLSE Activation Function Design on the Performance of Deep Neural Networks

ABSTRACT

The activation function is a critical component that significantly impacts the performance of deep neural networks. In this study, a novel activation function, RLSE (a combination of ReLU-LIP-Sigmoid-ELU), is proposed to achieve high classification accuracy and low loss in deep neural networks. The RLSE activation function aims to address the vanishing gradient problem and the dying ReLU issue. The performance of the RLSE activation function was evaluated on the MNIST and Fashion-MNIST datasets and compared with newly developed activation functions found in the literature and the most commonly used activation functions today. Using the RLSE activation function, the Convolutional Neural Network (CNN) architecture designed in this study achieved accuracy rates of 99.04% for the MNIST dataset and 90.40% for the Fashion-MNIST dataset. The results demonstrate that the RLSE activation function outperforms other activation functions.

Keywords: Activation function, convolutional neural network, deep neural networks.

1. GİRİŞ (INTRODUCTION)

Derin öğrenme modelleri, görüntü işleme, doğal dil işleme, nesne tanıma, çeviri ve finansal tahmin gibi birçok farklı alanda yaygın bir şekilde kullanılmakta ve bu modellerin önemi gün geçtikçe artmaktadır [1]. Eğitim verilerini kullanarak model geliştiren derin öğrenme algoritmaları, girdilerden bilinmeyen sonuçları tahmin etmek için tasarlanmıştır [2]. Derin öğrenme mimarileri, evrişim, havuzlama, normalleştirme, tam bağlantılı ve etkinleştirme katmanları gibi çeşitli katmanlardan oluşur [3]. Aktivasyon fonksiyonları, dropout ve benzeri tekniklerle birlikte, bu katmanların organizasyonları, genellikle son katmanlarda bulunan tam bağlı katmanlarla birlikte eksiksiz bir ESA yapısını oluşturur [4].

Aktivasyon fonksiyonu, özellikle bir ağı çok katmanlı yapısının bulunduğu derin ESA' da önemli bir rol oynayan, genellikle doğrusal olmayan bir fonksiyondur

[5]. Bir aktivasyon fonksiyonu, derin sinir ağlarının temel bileşeni olarak kabul edilir, çünkü bu fonksiyon, ağı verilerdeki karmaşık kalıpları öğrenmesine olanak tanıyan çok katmanlı algılayıcıların ve ESA'nın yapı taşını oluşturur. Sinir ağlarının doğruluk, hesaplama karmaşıklığı ve kayıp gibi performans metrikleri, kullanılan aktivasyon fonksiyonunun özelliklerine doğrudan bağlıdır [6]. Yaygın olarak kullanılan doğrusal olmayan aktivasyon fonksiyonları arasında sigmoid, tanh, rectified linear unit (ReLU) [7], [8], leaky ReLU (LReLU) [9], exponential linear unit (ELU) [10], scaled-exponential linear unit (SELU) [11], Swish [12] ve GELU [13] yer almaktadır. Sigmoid ve tanh, düzgün bir şekilde değişen doğrusal olmayan aktivasyon fonksiyonları olup, evrişimli sinir ağlarında kaybolan gradyan problemi nedeniyle etkin bir şekilde kullanılamazlar. ReLU, kaybolan gradyan sorununu çözerken, süreksizlik nedeniyle sıfır noktasında türevlenemez ve bu da ağı yorumlanmasını güçleştirir.

*Sorumlu Yazar (Corresponding Author)

e-posta : ismihangrbrz@gmail.com

Ayrıca, düşük hesaplama karmaşıklığına rağmen, girdiden bağımsız olarak yalnızca sıfır çıktısı verir; bu durum, “ölmekte olan ReLU” problemi olarak bilinir [9], [14]. LReLU, parametrik ReLU (PReLU) ve rastgele ReLU (RReLU) gibi aktivasyon fonksiyonları, ölmekte olan ReLU sorununu çözmekle birlikte, deaktivasyon durumlarında gürültüye karşı dayanıksızdırlar [10]. Sigmoid lineer birim (SiLU), ölmekte olan ReLU problemini ortadan kaldırmak için sigmoid geçitleme tekniğini kullanmakta, ancak sigmoid fonksiyonunun güçlü işlem gereksinimleri zaman karmaşıklığını arttırmaktadır [15]. Swish, SiLU’nun bir modifikasyonu olup, daha iyi doğruluk elde etmek amacıyla ek parametre eğitimine ihtiyaç duymakta ve bunun sonucunda daha yüksek hesaplama karmaşıklığına yol açmaktadır [12]. GELU ise negatif değerleri yumuşak bir şekilde azaltarak ölmekte olan ReLU probleminin üstesinden gelir. GELU negatif girişleri tamamen sıfırlamadığı için verinin tüm dağılımını kullanmasına olanak tanır ve öğrenme sürecini daha istikrarlı hale getirir [13].

Venkatappareddy ve arkadaşlarının gerçekleştirdiği çalışmada, ortogonal Legendre polinomları tabanlı lineer parametre (LIP) modeli kullanarak yeni bir aktivasyon fonksiyonu önerilmiştir [16], [17]. Bu model, ortalama havuzlamaya kıyasla gürültülü verilerle maksimum havuzlamanın daha iyi performans sergilemesinin arkasındaki mekanizmaları anlamaya yardımcı olmuştur. Jahan ve arkadaşlarının araştırmasında ise, ReLU ve swish aktivasyon fonksiyonlarının sınırlamalarını aşmak amacıyla, kaybolan gradyan, nöron ölümü ve çıkış ofseti gibi yaygın sorunları çözeabilen, kendinden kapalı bir ReLU (SGReLU) fonksiyonu önerilmiştir [18]. Kılıçarslan ve arkadaşlarının çalışmasında ise, ReLU, sigmoid ve ELU aktivasyon fonksiyonlarının bir birleşimi olan yeni RSigELU aktivasyon fonksiyonları, RSigELUS ve RSigELUD olarak tanımlanmıştır [19]. Bu yeni aktivasyon fonksiyonları, pozitif, negatif ve lineer aktivasyon bölgelerinde etkili bir şekilde çalışabilmekte ve kaybolan gradyan problemi ile negatif bölge sorununu aşabilmektedir.

Bu çalışmada, ReLU, LIP [16] ve RSigELU [19] aktivasyon fonksiyonlarının kombinasyonu olan RLSE isimli yeni bir aktivasyon fonksiyonu oluşturulmuştur. Önerilen RLSE aktivasyon fonksiyonu, pozitif aktivasyon bölgesinde ReLU ve ELU fonksiyonlarının bir kombinasyonunun davranışını sergilerken, negatif aktivasyon bölgesinde LIP fonksiyonunun davranışını sergilemektedir. RLSE aktivasyon fonksiyonunun performansı MNIST, Fashion MNIST veri setleri kullanılarak özel tasarlanmış ESA mimarisi üzerinde değerlendirilmiştir.

Bu çalışmanın ikinci bölümünde RLSE aktivasyon fonksiyonu ve tasarlanan ESA mimarisi sunulmaktadır. Aktivasyon fonksiyonlarını ve veri setlerini değerlendirmek için kullanılan yaklaşımlar üçüncü bölümde tartışılmış, RLSE aktivasyon fonksiyonunun performansı diğer aktivasyon fonksiyonlarının

performansı ile karşılaştırılmıştır. Elde edilen sonuçlar sonuç bölümünde tartışılmıştır.

2. METOD (METHOD)

2.1. RLSE Aktivasyon Fonksiyonu (RLSE Activation Function)

Bu çalışma, kaybolan gradyan ve negatif bölge problemlerinin üstesinden gelmek için üç aktif bölgedeki (pozitif, doğrusal ve negatif) verilere göre çalışabilen yeni bir aktivasyon fonksiyonu önermektedir. Önerilen RLSE aktivasyon fonksiyonu, çift parametre gerektirir. Bu parametrelerden α pozitif bölgenin eğim katsayısını ifade ederken, β negatif bölgenin eğim katsayısını ifade etmektedir. Pozitif bölgede ELU, doğrusal bölgede ReLU ve negatif bölgede tanh fonksiyonları kullanılarak tasarlanmış RLSE aktivasyon fonksiyonunun matematiksel gösterimi Eşitlik 1’deki gibidir.

$$f(x) = \begin{cases} \alpha(e^x - 1) & \text{if } 1 < x < \infty \\ x & \text{if } 0 < x < 1 \\ \frac{x(1+\tanh(\beta x))}{2} & \text{if } -\infty < x < 0 \end{cases} \quad (1)$$

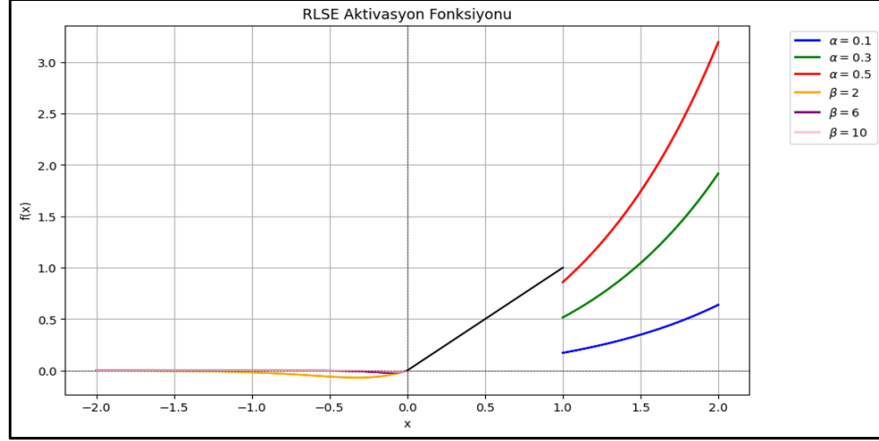
RLSE aktivasyon fonksiyonunun türevi Eşitlik 2’de görülmektedir. RLSE aktivasyon fonksiyonu her yerde türevlenebilir olduğundan öğrenme süreci sürekli. Böylelikle bir çok aktivasyon fonksiyonunun karşılaştığı kaybolan gradyan probleminin üstesinden gelebilir.

$$f(x) = \begin{cases} \alpha e^x & \text{if } 1 < x < \infty \\ 1 & \text{if } 0 < x < 1 \\ \tanh\left(\frac{\beta x}{2}\right) + \frac{\beta x}{2} \operatorname{sech}^2\left(\frac{\beta x}{2}\right) + 1 & \text{if } -\infty < x < 0 \end{cases} \quad (2)$$

Şekil 1, RLSE aktivasyon fonksiyonunun farklı α ve β parametreleri için davranışını göstermektedir. Grafikte, $1 < x < \infty$ bölgesinde fonksiyonun üstel bileşeni $\alpha(e^x - 1)$ kullanılarak pozitif girişlerde doğrusal olmayan bir büyüme sağlandığı görülmektedir. Bu bölgedeki farklı α değerleri fonksiyonun eğimini belirlemekte ve büyük α değerlerinin daha hızlı artışa yol açtığı gözlemlenmektedir.

$0 < x < 1$ aralığında fonksiyon, $f(x) = x$ olacak şekilde doğrusal bir geçiş sağlamaktadır. Bu yapı, küçük pozitif girişlerin doğrudan iletilmesini mümkün kılarak, türevlenebilirliği ve eğitim sürecindeki stabiliteyi artırmaktadır.

$-\infty < x < 0$ bölgesinde fonksiyon, $\frac{x(1+\tanh(\beta x))}{2}$ ifadesi ile belirlenmiş olup, burada β parametresi negatif girişlerin doygunluk seviyesini kontrol etmektedir. Düşük β değerleri daha kademeli bir geçiş sağlarken, yüksek β değerleri negatif girişlerde daha hızlı bir doyuma ulaşılmasına neden olmaktadır. Bu özellik, RLSE’ nin negatif girişlerde aşırı baskılanmayı önleyerek öğrenme sürecini iyileştirmesine olanak tanımaktadır.



Şekil 1. RLSE aktivasyon fonksiyonu grafiği (RLSE activation function graph)

RLSE aktivasyon fonksiyonu, α ve β parametreleri aracılığıyla pozitif ve negatif girişler için esnek bir dönüşüm sağlayarak, derin sinir ağlarında farklı öğrenme dinamiklerine uyum sağlayabilecek bir yapı sunmaktadır.

2.2. Veri Kümesi (Dataset)

RLSE aktivasyon fonksiyonunun, incelenen 3 makaledeki aktivasyon fonksiyonlarıyla ve günümüzde en çok kullanılan ReLU, Swish ve GELU aktivasyon fonksiyonlarıyla performans açısından karşılaştırmayı yapmak için MNIST ve Fashion-MNIST veri setleri kullanılmıştır.

Şekil 2.a' da örneği görülen MNIST veri kümesi toplam 70000 el yazısı rakam ve 10 sınıftan oluşmaktadır [20]. Bu çalışmada, veri kümesinde yer alan 60000 adet görsel eğitim için, 10000 adet görsel ise test için kullanılmıştır. Veri kümesindeki her görüntü gri tonlamalı ve 28×28 piksel boyutundadır. Veri kümesindeki görüntüler csv formatında kullanılmıştır.

Zalando araştırma ekibi tarafından geliştirilen Şekil 2.b' deki Fashion-MNIST veri kümesi, MNIST veri kümesine alternatif olarak 70000 giyim görüntüsünden oluşmaktadır [21]. Bu çalışmada, veri kümesinde yer alan 60000 adet görsel eğitim için, 10000 adet görsel ise

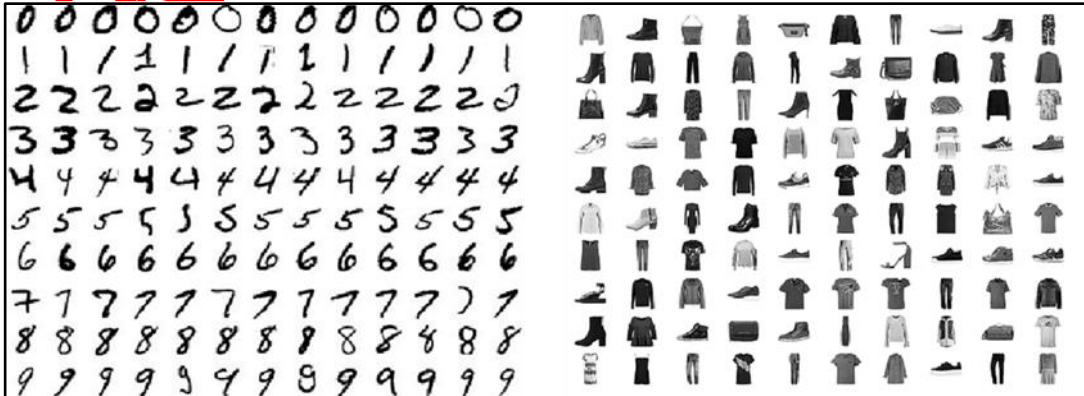
test için kullanılmıştır. Veri kümesi 10 sınıfa sahiptir. Her bir görüntü gri tonlamalı ve 28×28 piksel boyutundadır. Veri kümesindeki görüntüler csv formatında kullanılmıştır.

2.3. Evrişimsel Sinir Ağı Modeli (Convolutional Neural Network Model)

Evrişimsel sinir ağları, özellikle görüntü işleme ve bilgisayarla görme alanlarında yaygın olarak kullanılan derin öğrenme modelleridir [22]. CNN'ler, katmanlı yapıları sayesinde girdilerden anlamlı özellikler çıkararak sınıflandırma ve tanıma görevlerinde yüksek performans sergiler [23]. CNN modelleri genellikle evrişim katmanı, aktivasyon fonksiyonu, havuzlama katmanı, düzleştirme katmanı ve tam bağlantılı katman temel bileşenlerinden oluşur [24].

Evrişim katmanında, giriş verisinden özellikler çıkarmak için filtreler kullanılır [25]. Bu filtreler, giriş verisi üzerinde kaydırılarak belirli paternleri tespit eder. Her bir filtrenin çıktısı, özellik haritası olarak adlandırılır. Bu süreç, görüntüdeki kenarlar, dokular ve diğer temel özelliklerin yakalanmasını sağlar [26].

Evrişim işlemi sonrasında, doğrusal olmayanlık eklemek ve modeli daha karmaşık ilişkileri öğrenebilir hale getirmek için aktivasyon fonksiyonları kullanılır [27].



Şekil 2. (a) MNIST veri kümesi örneği (MNIST dataset example)

(b) Fashion MNIST veri kümesi örneği (Example of Fashion MNIST dataset)

Katman	Çıkış Boyutu	Parametre
Evrişim Katmanı	(None, 28, 28, 16)	416
Evrişim Katmanı	(None, 28, 28, 16)	6416
Yığın Normalleştirme	(None, 28, 28, 16)	64
Maksimum Havuzlama	(None, 14, 14, 16)	0
Evrişim Katmanı	(None, 14, 14, 32)	8224
Yığın Normalleştirme	(None, 14, 14, 32)	128
Maksimum Havuzlama	(None, 7, 7, 32)	0
Evrişim Katmanı	(None, 7, 7, 64)	18496
Yığın Normalleştirme	(None, 7, 7, 64)	256
Maksimum Havuzlama	(None, 3, 3, 64)	0
Evrişim Katmanı	(None, 3, 3, 128)	73856
Yığın Normalleştirme	(None, 3, 3, 128)	512
Maksimum Havuzlama	(None, 2, 2, 128)	0
Düzleştirme Katmanı	(None, 512)	0
Tam Bağlantılı Katman	(None, 1024)	525312
Sönümleme	(None, 1024)	0
Tam Bağlantılı Katman	(None, 10)	10250

Şekil 6. RLSE aktivasyon fonksiyonunu Fashion-MNIST veri kümesi üzerinde değerlendirmek için tasarlanan ESA mimarisinin çıkış boyutları ve parametre sayıları. (Output dimensions and parameter numbers of the ESA architecture designed to evaluate the RLSE activation function on the Fashion-MNIST dataset.)

ESA mimarisi, Şekil 3' teki ESA mimarisine göre daha fazla katmandan oluşmaktadır. Fashion-MNIST veri kümesinin eğitiminde kullanılan ESA mimarisi Şekil 5' te görülmektedir.

Fashion-MNIST veri kümesinin eğitiminde kullanılan ESA mimarisinin her katmanının çıkış boyutu ve parametre sayısı Şekil 6' da görülmektedir. Fashion-MNIST veri kümesinin eğitiminde toplamda 1930832 parametre kullanılarak %90,40 doğruluk oranına ulaşılmıştır.

3. DENEYSEL SONUÇLAR (EXPERIMENTAL RESULTS)

Çalışmada β eğim katsayısı $0 < \beta < 10$ aralığında değerler alabilir ve α eğim katsayısı $0 < \alpha < 1$ aralığında değerler almaktadır. Bu değerler pozitif ve negatif bölgeleri kontrol etmek için kullanılmaktadır. MNIST veri kümesinde yapılan deneylerde geçerleme kayıp ve doğruluk oranları için en iyi ölçek değerleri Tablo 1' de sunulmuştur. Deneysel çalışmalarda elde edilen sonuçlara göre en iyi α değeri 0,3 , en iyi β değeri 6 olarak bulunmuş ve tüm deneylerde kullanılmıştır.

Çizelge 1'e göre, en iyi başarı oranı RLSE aktivasyon fonksiyonunda $\alpha = 0,3$ ve $\beta = 6$ parametre değerleri ile elde edilmiş olup, geçerleme doğruluk ve geçerleme kayıp değerleri sırasıyla 0,9904 ve 0,0426 olarak bulunmuştur. Çizelge 2' de MNIST veri kümesi için diğer aktivasyon fonksiyonları ve RLSE aktivasyon fonksiyonun kullanılması durumunda eğitim ve geçerleme kayıp ve doğruluk oranları görülmektedir. Çizelge 2' de de görüldüğü gibi, RLSE aktivasyon fonksiyonu MNIST veri kümesi için en iyi başarı oranına sahiptir ve geçerleme doğruluğu açısından diğer aktivasyon fonksiyonlarından daha iyi performans göstermektedir.

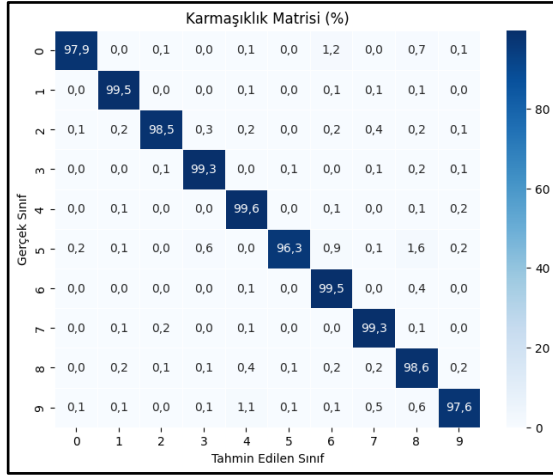
Şekil 7' de MNIST veri kümesi görüntülerinin RLSE aktivasyon fonksiyonu kullanılarak oluşturulan ESA mimarisinde gerçekleştirilen onlu sınıflandırmasının karmaşıklık matrisi bulunmaktadır. Şekil 7' deki karmaşıklık matrisinde de görüldüğü gibi MNIST veri kümesi üzerinde gerçekleştirilen onlu sınıflandırmada 1918 örnekten 1911' ini doğru tahmin ederek (%99,6) en çok ayırt edilen sınıf 4. sınıf olmuştur. 5. sınıf ise en düşük doğrulukla (%96,3) belirlenmiştir.

Çizelge 1. MNIST veri kümesi için RLSE aktivasyon fonksiyonunun deneysel sonuçları.

$\beta \backslash \alpha$	2	2	6	6	10	10
	Geçerleme Kaybı	Geçerleme Doğruluğu	Geçerleme Kaybı	Geçerleme Doğruluğu	Geçerleme Kaybı	Geçerleme Doğruluğu
0,1	0,0710	0,9850	0,0733	0,9842	0,0571	0,9855
0,2	0,0608	0,9868	0,0475	0,9887	0,0457	0,9890
0,3	0,0464	0,9894	0,0426	0,9904	0,0712	0,9841
0,4	0,0436	0,9895	0,0596	0,9872	0,0661	0,9832
0,5	0,0533	0,9866	0,0508	0,9875	0,0656	0,9839

Çizelge 2. MNIST veri kümesi için RLSE aktivasyon fonksiyonunun en iyi deneysel sonuçları

Aktivasyon Fonksiyonu	Epoch	Eğitim Kaybı	Geçerleme Kaybı	Eğitim Doğruluğu	Geçerleme Doğruluğu
ReLU	10	0,0087	0,0529	0,9970	0,9875
Swish		0,0076	0,0571	0,9974	0,9863
Gelu		0,0074	0,0507	0,9975	0,9879
RSigELU [19]		0,0064	0,0641	0,9981	0,9835
SGReLU [18]		0,0078	0,0671	0,9971	0,9856
LIP [16]		0,0062	0,0548	0,9978	0,9880
RLSE		0,0063	0,0426	0,9979	0,9904

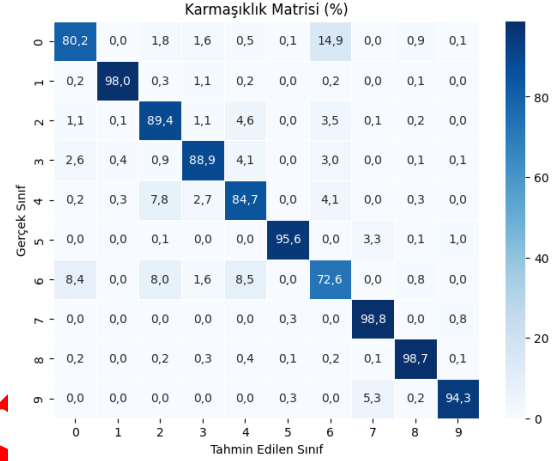


Şekil 7. RLSE aktivasyon fonksiyonunu MNIST veri kümesi üzerinde değerlendirmek için tasarlanan ESA mimarisinin onlu sınıflandırma karmaşıklık matrisi (Decimal classification complexity matrix of the ESA architecture designed to evaluate the RLSE activation function on the MNIST dataset)

Fashion-MNIST veri kümesinde yapılan deneylerde geçerleme kayıp ve doğruluk oranları için en iyi ölçek değerleri Çizelge 3' te sunulmuştur. Deneysel çalışmalarda elde edilen sonuçlara göre en iyi α değeri 0,2 , en iyi β değeri 6 olarak bulunmuş ve tüm deneylerde kullanılmıştır.

Çizelge 4' te görüldüğü gibi, RLSE aktivasyon fonksiyonu Fashion-MNIST veri kümesi için de en iyi başarı oranına sahiptir. En iyi başarı oranı RLSE aktivasyon fonksiyonunda $\alpha = 0,2$ ve $\beta = 6$ parametre değerleri ile elde edilmiş olup, geçerleme doğruluk ve geçerleme kayıp değerleri sırasıyla 0,9040 ve 0,5653 olarak bulunmuştur. RLSE aktivasyon fonksiyonu,

geçerleme doğruluğu açısından diğer aktivasyon fonksiyonlarından daha iyi performans göstermektedir. Şekil 8' de Fashion-MNIST veri kümesi görüntülerinin RLSE aktivasyon fonksiyonu kullanılarak oluşturulan ESA mimarisinde gerçekleştirilen onlu sınıflandırmasının karmaşıklık matrisi bulunmaktadır. Şekil 8' deki karmaşıklık matrisinde de görüldüğü gibi Fashion-MNIST veri kümesi üzerinde gerçekleştirilen onlu sınıflandırmada 2014 örnekten 1989' unu doğru tahmin ederek en çok ayırt edilen sınıf 7. sınıf (%98,8) olmuştur. 6. sınıf ise en düşük doğrulukla (%72,6) belirlenmiştir.



Şekil 8. RLSE aktivasyon fonksiyonunu Fashion-MNIST veri kümesi üzerinde değerlendirmek için tasarlanan ESA mimarisinin onlu sınıflandırma karmaşıklık matrisi (Decimal classification complexity matrix of the ESA architecture designed to evaluate the RLSE activation function on the Fashion-MNIST dataset)

Çizelge 3. Fashion-MNIST veri kümesi için RLSE aktivasyon fonksiyonunun deneysel sonuçları. (Experimental results of RLSE activation function for Fashion-MNIST dataset.)

β \ α	2	2	6	6	10	10
	Geçerleme Kaybı	Geçerleme Doğruluğu	Geçerleme Kaybı	Geçerleme Doğruluğu	Geçerleme Kaybı	Geçerleme Doğruluğu
0,1	0,4373	0,9027	0,5268	0,8954	0,4506	0,9035
0,2	0,4101	0,9015	0,5653	0,9040	0,5298	0,9029
0,3	0,5024	0,9032	0,4301	0,9032	0,4490	0,8942
0,4	0,4881	0,8972	0,4694	0,9038	0,4369	0,8988
0,5	0,4261	0,8971	0,4443	0,8959	0,4837	0,8989

Çizelge 4. Fashion- MNIST veri kümesi için RLSE aktivasyon fonksiyonunun en iyi deneysel sonuçları. (Fashion- Best experimental results of RLSE activation function for MNIST dataset.)

Aktivasyon Fonksiyonu	Epoch	Eğitim Kaybı	Geçerleme Kaybı	Eğitim Doğruluğu	Geçerleme Doğruluğu
ReLU	20	0,0530	0,5692	0,9809	0,8828
Swish		0,0464	0,4917	0,9820	0,8972
Gelu		0,0532	0,5370	0,9802	0,8902
RSigELU [19]		0,0471	0,5242	0,9816	0,8976
SGReLU [18]		0,0578	0,5253	0,9786	0,8942
LIP [16]		0,0579	0,4795	0,9775	0,8981
RLSE		0,0461	0,5653	0,9828	0,9040

Çizelge 5' te görülen sonuçlar, RLSE aktivasyon fonksiyonunu MNIST ve Fashion- MNIST veri kümeleri üzerinde değerlendirmek için tasarlanan ESA mimarisinin sınıflandırma problemi üzerindeki istatistiksel performans metriklerini sunmaktadır.

Çizelge 5. RLSE aktivasyon fonksiyonu ile MNIST ve Fashion-MNIST veri kümelerinin sınıflandırma performansını değerlendirme metrikleri (Metrics for evaluating the classification performance of MNIST and Fashion-MNIST datasets with RLSE activation function)

Veri Seti	Kesinlik (%)	F1 Skoru (%)	Duyarlılık (%)	ROC-AUC Skoru (%)
MNIST	98,68	98,66	98,66	99,24
Fashion MNIST	90,12	90,03	90,04	94,50

Bu metrikler, modellerin sınıflandırma görevlerindeki doğruluğunu ve güvenilirliğini değerlendirmek için kullanılır.

Kesinlik, modelin doğru pozitif tahminlerinin tüm pozitif tahminlere oranını ölçer. Yüksek bir kesinlik, modelin yanlış pozitif sayısını azaltmada iyi olduğunu göstermektedir. Kesinlik, Eşitlik 3' te görüldüğü gibi hesaplanmaktadır.

$$Kesinlik = \frac{Doğru Pozitif (DP)}{Doğru Pozitif (DP) + Yanlış Pozitif (YP)} \quad (3)$$

Çizelge 5' te, RLSE aktivasyon fonksiyonu ile gerçekleştirilen onlu MNIST sınıflandırmasında %98,68 kesinlik değeri elde edilmiştir, bu da modelin yanlış pozitif tahminlerden kaçınmada oldukça etkili olduğunu göstermektedir.

Duyarlılık, modelin doğru pozitif tahminlerinin tüm gerçek pozitiflere oranını ifade etmekte ve modelin yanlış negatifleri en aza indirme yeteneğini ölçmektedir. Duyarlılık, Eşitlik 4' te görüldüğü gibi hesaplanmaktadır.

$$Duyarlılık = \frac{Doğru Pozitif (DP)}{Doğru Pozitif (DP) + Yanlış Negatif (YN)} \quad (4)$$

RLSE aktivasyon fonksiyonu ile gerçekleştirilen onlu MNIST sınıflandırmasında %98,66 duyarlılık oranı ile yüksek performans sergilenmiştir, bu da modelin gerçek pozitifleri yakalama konusunda başarılı olduğunu göstermektedir.

F1 skoru, bir modelin doğruluğunu ve hassasiyetini dengeleyen bir ölçüttür. Yüksek F1 skoru, modelin hem doğru pozitifleri yakalamada hem de yanlış pozitiflerden kaçınmada iyi performans gösterdiğini ifade etmektedir. F1 skoru, Eşitlik 5' te görüldüğü gibi hesaplanmaktadır.

$$F1 Skoru = 2 \times \frac{Kesinlik \times Duyarlılık}{Kesinlik + Duyarlılık} \quad (5)$$

Çizelge 5' teki verilere göre, RLSE aktivasyon fonksiyonu kullanılarak MNIST veri kümesinin F1 skoru %98,66 olarak belirlenirken Fashion MNIST veri kümesinin F1 skoru %90,03 olarak belirlenmiştir.

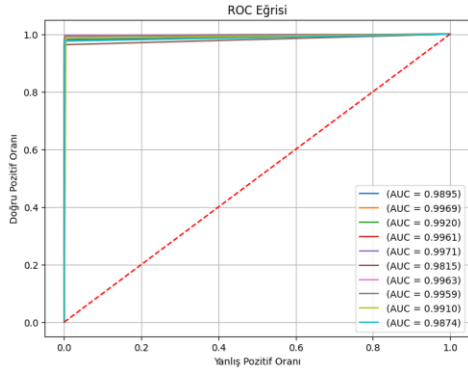
ROC-AUC skoru, bir modelin sınıflandırma yeteneğinin genel bir ölçüsüdür ve modelin pozitif ve negatif sınıfları ne kadar iyi ayırt edebildiğini göstermektedir. ROC bir olasılık eğrisidir ve altında kalan alan olan AUC ayrılabilirliğin derecesini veya ölçüsünü temsil eder. ROC eğrisinde x ekseninde YPO (Yanlış Pozitif Oran), y ekseninde ise DPO (Doğru Pozitif Oranı) bulunmaktadır. Eğrinin altında kalan arttıkça sınıflar arasında ayırt etme performansı artmaktadır. DPO, duyarlılık değeridir. YPO ise hatalı tahmin yapma oranıdır. YPO Eşitlik 6' da görüldüğü gibi hesaplanmaktadır.

$$YPO = \frac{Yanlış Pozitif (YP)}{Yanlış Pozitif (YP) + Doğru Negatif (DN)} \quad (6)$$

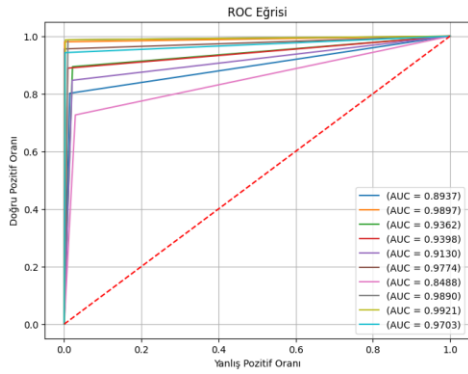
Yüksek bir ROC-AUC skoru, modelin tüm sınıflandırma eşiklerinde genel olarak iyi performans gösterdiğini ifade etmektedir. RLSE aktivasyon fonksiyonu ile gerçekleştirilen onlu MNIST sınıflandırmasında %99,24 ROC-AUC skoru elde edilmiştir, bu da modelin pozitif ve negatif sınıfları ayırt etmede iyi olduğunu gösterir. Şekil 9' da MNIST ve Fashion MNIST veri kümeleri üzerinde değerlendirmek için tasarlanan ESA mimarisinin onlu sınıflandırmadaki ROC eğrileri görülmektedir.

4. SONUÇ VE TARTIŞMA (CONCLUSION and DISCUSSION)

Aktivasyon katmanları, ağa gelen girdileri işler ve aktivasyon fonksiyonlarını kullanarak bu girdilere karşılık gelen çıktıyı üretir. Derin öğrenme mimarileri genellikle karmaşık problemlerle ilgilendiğinden, doğrusal aktivasyon fonksiyonları bu tür mimariler için uygun değildir. Bunun nedeni, aktivasyon fonksiyonlarının, sinir ağlarının değişkenler arasındaki karmaşık ve sürekli ilişkiyi öğrenmesi ve tahmin etmesi için temel oluşturmalarıdır. Bu nedenle derin öğrenme modellerinde doğrusal olmayan aktivasyon fonksiyonları ön plana çıkmaktadır. Bununla birlikte, yaygın olarak bilinen sigmoid ve tanjant aktivasyon fonksiyonlarında kaybolan gradyan sorunu vardır. Bu sorunun üstesinden gelmek için ReLU aktivasyon fonksiyonu önerilmiştir. Ancak, ReLU etkinleştirme işlevi, negatif değerler için sıfır verdiğinden, negatif bölge için doğru bir değer döndüremez. ReLU'daki negatif değer probleminin üstesinden gelmek için alternatif aktivasyon fonksiyonları önerilmiştir. Bu nedenle, çalışmamızda hem kaybolan gradyan hem de negatif değer problemlerinin üstesinden gelmek için yeni çift parametrelili RLSE aktivasyon fonksiyonu önerilmiştir. Önerilen RLSE aktivasyon fonksiyonu, eğitim sürecinde başarı oranını artırma özelliğiyle ön plana çıkmaktadır. RLSE aktivasyon fonksiyonu ile MNIST ve Fashion MNIST veri kümeleri üzerinde yapılan deneyler



(a)



(b)

Şekil 9. RLSE aktivasyon fonksiyonunu (a) MNIST ve (b) Fashion-MNIST veri kümeleri üzerinde değerlendirmek için tasarlanan ESA mimarisinin onlu sınıflandırmalardaki ROC Eğrisi grafikleri (ROC Curve plots of the ESA architecture designed to evaluate the RLSE activation function on (a) MNIST and (b) Fashion-MNIST datasets in decimal classifications.)

literatüre göre daha iyi sonuçlar elde edilmesini sağlamıştır. Gelecekte, bu fonksiyonun daha derin ve geniş ölçekli ağlarda test edilmesi, farklı optimizasyon algoritmalarıyla entegrasyonu ve transfer öğrenme senaryolarında performansının değerlendirilmesi önem arz etmektedir. Ayrıca, RLSE'nin farklı veri türleri üzerinde, özellikle doğal dil işleme ve tıbbi görüntüleme gibi alanlarda nasıl bir etki yarattığının araştırılması, fonksiyonun genel uygulanabilirliği açısından faydalı olacaktır. Bunun yanı sıra, RLSE'nin parametrik yapısının adaptif hale getirilerek öğrenilebilir parametreler aracılığıyla model tarafından optimize edilmesi, aktivasyon fonksiyonlarının dinamik olarak ayarlanabilir hale gelmesini sağlayabilir. Bu tür geliştirmeler, RLSE'nin daha geniş kullanım alanlarına yayılmasını ve derin öğrenme modellerinde daha esnek ve verimli bir aktivasyon fonksiyonu olarak konumlandırılmasını mümkün kılacaktır.

ETİK STANDARTLARIN BEYANI (DECLARATION OF ETHICAL STANDARDS)

Bu makalenin yazar(lar)ı çalışmalarında kullandıkları materyal ve yöntemlerin etik kurul izni ve/veya yasal-özel bir izin gerektirmediğini beyan ederler.

YAZARLARIN KATKILARI (AUTHORS' CONTRIBUTIONS)

İsmihan Gül ÖZELOĞLU: Kavramsallaştırma, araştırma, metodoloji geliştirme, yazılım ve deneylerin gerçekleştirilmesi, deney sonuçlarının analizi ve makalenin yazımını yürütmüştür.

Eda AKMAN AYDIN: Kavramsallaştırma, deney sonuçlarının analizi ve yorumlanması, makalenin gözden geçirilmesi ve düzenlenmesini gerçekleştirmiştir.

Necaattin BARIŞCI: Kavramsallaştırma, metodoloji geliştirme ve deney sonuçlarının yorumlanması çalışmalarını yürütmüştür

ÇIKAR ÇATIŞMASI (CONFLICT OF INTEREST)

Bu çalışmada herhangi bir çıkar çatışması yoktur.

KAYNAKLAR (REFERENCES)

- [1] Sengupta, S., Basak, S., Salkia, P., Paul, S., Tsalavoutis, V., Atiah, E., & Peters, A., "A review of deep learning with special emphasis on architectures, applications and recent trends", *Knowledge-Based Systems*, 194: 105596, (2020).
- [2] Sarker, I. H., "Machine learning: Algorithms, real-world applications and research directions", *SN computer science*, 2(3): 160, (2021).
- [3] Cong, S., & Zhou, Y., "A review of convolutional neural network architectures and their optimizations", *Artificial Intelligence Review*, 56(3): 1905-1969, (2023).
- [4] Zheng, Y., Gao, Z., Wang, Y., & Fu, Q., "MOOC dropout prediction using FWTS-CNN model based on fused feature weighting and time series", *IEEE Access*, 8: 225324-225335, (2020).
- [5] Dubey, S. R., Singh, S. K., & Chaudhuri, B. B., "Activation functions in deep learning: A comprehensive survey and benchmark", *Neurocomputing*, 503: 92-108, (2022).
- [6] Krichen, M., "Convolutional neural networks: A survey", *Computers*, 12(8): 151, (2023).
- [7] V. Nair, G.E. Hinton, Rectified linear units improve restricted boltzmann machines, in: Proc. International Conference on Machine Learning, Haifa, Israel, 807-814, (2010).
- [8] X. Glorot, A. Bordes, Y. Bengio, Deep sparse rectifier neural networks, in: Proc. International Conference on Artificial Intelligence and Statistics Conference, Ft. Lauderdale, FL, USA, (2011).
- [9] Apicella A., Donnarumma F., Isgrò F., Prevete R., "A survey on modern trainable activation functions", *Neural Netw.*, 138: 14-32, (2021).
- [10] Clevert D.-A., Unterthiner T., Hochreiter S., "Fast and accurate deep network learning by exponential linear units (ELUs)", *arXiv [cs.LG]*, (2015).
- [11] Klambauer G., Unterthiner T., Mayr A., Hochreiter S., "Self-normalizing neural networks", *arXiv [cs.LG]*, (2017).
- [12] Ramachandran, P., Zoph, B., & Le, Q. V., "Swish: a self-gated activation function", *arXiv preprint arXiv:1710.05941*, 7(1): 5, (2017).

- [13] Hendrycks, D., & Gimpel, K., "Gaussian error linear units (gelus)", *arXiv preprint arXiv:1606.08415*, (2016).
- [14] Lu, Lu, et al. "Dying relu and initialization: Theory and numerical examples." *arXiv preprint arXiv:1903.06733*, (2019).
- [15] Elfving, Stefan, Eiji Uchibe, and Kenji Doya., "Sigmoid-weighted linear units for neural network function approximation in reinforcement learning", *Neural networks*, 107: 3-11, (2018).
- [16] Venkatappareddy, P., Culli, J., Srivastava, S., & Lall, B., "A Legendre polynomial based activation function: An aid for modeling of max pooling", *Digital Signal Processing*, 115: 103093, (2021).
- [17] Carini, Alberto, et al., "Legendre nonlinear filters", *Signal Processing*, 109: 84-94, (2015).
- [18] Jahan, Israt, et al., "Self-gated rectified linear unit for performance improvement of deep neural networks", *ICT Express*, 9(3): 320-325, (2023).
- [19] Kiliçarslan, S., & Celik, M., "RSigELU: A nonlinear activation function for deep neural networks", *Expert Systems with Applications*, 174: 114805, (2021).
- [20] LeCun, Y., Bottou, L., Bengio, Y., & Haffner, P., "Gradient-based learning applied to document recognition", *Proceedings of the IEEE*, 86(11): 2278-2324, (1998).
- [21] Xiao, H., Rasul, K., & Vollgraf, R., "Fashion-mnist: a novel image dataset for benchmarking machine learning algorithms", *arXiv preprint arXiv:1708.07747*, (2017).
- [22] Zhao, X., Wang, L., Zhang, Y., Han, X., Deveci, M., & Parmar, M., "A review of convolutional neural networks in computer vision", *Artificial Intelligence Review*, 57(4): 99, (2024).
- [23] Fan, C. L., "Multiscale Feature Extraction by Using Convolutional Neural Network: Extraction of Objects from Multiresolution Images of Urban Areas", *ISPRS International Journal of Geo-Information*, 13(1): 5, (2023).
- [24] Taye, M. M., "Theoretical understanding of convolutional neural network: Concepts, architectures, applications, future directions", *Computation*, 11(3): 52, (2023).
- [25] Zhao, L., & Zhang, Z., "A improved pooling method for convolutional neural networks", *Scientific Reports*, 14(1): 1589, (2024).
- [26] Elizar, E., Zulkifley, M. A., Muharar, R., Zaman, M. H. M., & Mustaza, S. M., "A review on multiscale-deep-learning applications", *Sensors*, 22(19): 7384, (2022).
- [27] Dubey, S. R., Singh, S. K., & Chaudhuri, B. B., "Activation functions in deep learning: A comprehensive survey and benchmark", *Neurocomputing*, 503: 92-108, (2022).
- [28] Özdemir, C., "Avg-topk: A new pooling method for convolutional neural networks", *Expert Systems with Applications*, 223: 119892, (2023).
- [29] Zafar, A., Aamir, M., Mohd Nawi, N., Arshad, A., Riaz, S., Alruban, A., ... & Almotairi, S., "A comparison of pooling methods for convolutional neural networks", *Applied Sciences*, 12(17): 8643, (2022).
- [30] Jeczminek, E., & Kowalski, P. A., "Flattening layer pruning in convolutional neural networks", *Symmetry*, 13(7): 1147, (2021).
- [31] Ullah, U., Jurado, A. G. O., Gonzalez, I. D., & Garcia-Zapirain, B., "A fully connected quantum convolutional neural network for classifying ischemic cardiopathy", *IEEE Access*, 10: 134592-134605, (2022).