

TÜRKİYE CUMHURİYETİ  
YILDIZ TEKNİK ÜNİVERSİTESİ  
FEN BİLİMLERİ ENSTİTÜSÜ

GENİŞ ÖLÇEKLİ VERİLER ÜZERİNDE SINIFLANDIRMA  
VE BÖLÜTLEME AMAÇLI EVRİŞİMSEL SİNİR AĞI VE  
İSTATİSTİKSEL MODELLERİN GELİŞTİRİLMESİ

Nurullah ÇALIK

DOKTORA TEZİ  
Elektronik ve Haberleşme Mühendisliği Anabilim Dalı  
Haberleşme Programı

Danışman  
Prof. Dr. Lutfiye DURAK ATA

Temmuz, 2019

**TÜRKİYE CUMHURİYETİ**  
**YILDIZ TEKNİK ÜNİVERSİTESİ**  
**FEN BİLİMLERİ ENSTİTÜSÜ**

**GENİŞ ÖLÇEKLİ VERİLER ÜZERİNDE SINIFLANDIRMA VE**  
**BÖLÜTLEME AMAÇLI EVRİŞİMSEL SİNİR AĞI VE İSTATİSTİKSEL**  
**MODELLERİN GELİŞTİRİLMESİ**

Nurullah ÇALIK tarafından hazırlanan tez çalışması 22.07.2019 tarihinde aşağıdaki jüri tarafından Yıldız Teknik Üniversitesi Fen Bilimleri Enstitüsü Elektronik ve Haberleşme Mühendisliği Anabilim Dalı Haberleşme Programı **DOKTORA TEZİ** olarak kabul edilmiştir.

Prof. Dr. Lütfiye DURAK ATA  
İstanbul Teknik Üniversitesi  
Danışman

**Jüri Üyeleri**

Prof. Dr. Lütfiye DURAK ATA, Danışman  
İstanbul Teknik Üniversitesi

Prof. Dr. M. S. Ufuk TÜRELİ, Üye  
Yıldız Teknik Üniversitesi

Dr. Öğr. Üyesi Yusuf YASLAN, Üye  
İstanbul Teknik Üniversitesi

Doç. Dr. Umut Engin AYTEN, Üye  
Yıldız Teknik Üniversitesi

Dr. Öğr. Üyesi Mehmet Akif YAZICI, Üye  
İstanbul Teknik Üniversitesi

---

---

---

---

---

Danışmanım Prof. Dr. Lutfiye DURAK ATA sorumluluğunda tarafımca hazırlanan Geniş Ölçekli Veriler Üzerinde Sınıflandırma ve Bölütleme Amaçlı Evrişimsel Sinir Ağı ve İstatistiksel Modellerin Geliştirilmesi başlıklı çalışmada veri toplama ve veri kullanımında gerekli yasal izinleri aldığımı, diğer kaynaklardan aldığım bilgileri ana metin ve referanslarda eksiksiz gösterdiğimi, araştırma verilerine ve sonuçlarına ilişkin çarpıtma ve/veya sahtecilik yapmadığımı, çalışmam süresince bilimsel araştırma ve etik ilkelerine uygun davrandığımı beyan ederim. Beyanımın aksinin ispatı halinde her türlü yasal sonucu kabul ederim.

Nurullah ÇALIK

İmza

*Kıymetli aileme ve  
değerli çalışma arkadaşlarıma ithafen.*



## TEŞEKKÜR

---

Doktora eğitimi, verilen vasıf itibarıyla, literatüre katkı sunulan algoritmanın, yöntemin veya modelin geliştirilmesinin yanı sıra ilgilenilen konu ile alakalı felsefeyi de kapsamaktadır. İçinde derin bir fikirsel sistematığı barındıran teknik konuların ele alınması esnasında bir yol göstericiye ihtiyaç duyulur. Sadece ortaya koyulan çalışmanın seyrinde kaybolmamak için değil, ayrıca bilimsel etiğin ve erdemin öğrenileceği bir danışman elzemdir. Bu noktada, bana akademinin ahlakını ve ahkâmını öğrettiği için değerli Hocam Prof. Dr. Lütfiye DURAK ATA ile tanışmayı büyük bir kıymet ve önemli bir kısmet olarak addederim. Ayrıca Tez İzleme Komitesinde bulunarak tezime sundukları değerli katkılarından ve yönlendirmelerinden dolayı Prof. Dr. M. S. Ufuk TÜRELİ Hocam ile Doç. Dr. Behçet Uğur TÖREYİN Hocam'a teşekkürlerimi sunarım. Yapılan çalışmalar esnasında, Medipol Hastanesi Patoloji Birimine ve Argenit Şirketinden Dr. Abdülkerim ÇAPAR'a sağladıkları veriler ve verdikleri teknik imkanlardan dolayı müteşekkirim.

Doktora sürecinde birçok inişler ve çıkışlar yaşanmaktadır. Olmayacağının düşünüldüğü anlarda sizleri destekleyen yol arkadaşlarının bulunması, bu çetin dönemde hedefe varmak için önemli bir etmendir. Durgunluğun hasıl olduğu böylesi anlarda beni motive eden ve bu çalışmanın gidişatında önemli katkıları bulunan Onur Can KURBAN, Ali Rıza YILMAZ'a; ayrıca doktoramın son dönemlerinde her ihtiyaç duyduğumda benden yardımlarını esirgemeyen Muhammet Ali KARABULUT, Hüseyin ÇUKUR ve Hüsamettin UYSAL'a teşekkürü bir borç bilirim.

Bir ailenin çocuklarına bırakabileceği en güzel miras, onların eğitimi ve öğretimidir. Bu değerler bireyin kişiliğini imâr eden en önemli unsurlar arasında yer alır. Diğer taraftan, insanın eğitim ve öğretim serüveni zorlu bir yolculuktur. Bu yolculuk esnasında beni, canlarından özge bir can bilip maddi ve manevi desteklerini esirgemeyen sevgili aileme minnettarım. Bu çalışmanın nihayete erişmesinde büyük bir hisse sahibi olan ve üstün bir fedakarlıkla benden desteğini hiç eksik etmeyen eşime şükranlarımı sunarım.

Nurullah ÇALIK

# İÇİNDEKİLER

---

|                                                                                                                      |             |
|----------------------------------------------------------------------------------------------------------------------|-------------|
| <b>SİMGE LİSTESİ</b>                                                                                                 | <b>vii</b>  |
| <b>KISALTMA LİSTESİ</b>                                                                                              | <b>viii</b> |
| <b>ŞEKİL LİSTESİ</b>                                                                                                 | <b>x</b>    |
| <b>TABLO LİSTESİ</b>                                                                                                 | <b>xiii</b> |
| <b>ÖZET</b>                                                                                                          | <b>xiv</b>  |
| <b>ABSTRACT</b>                                                                                                      | <b>xvi</b>  |
| <b>1 Giriş</b>                                                                                                       | <b>1</b>    |
| 1.1 Literatür Özeti . . . . .                                                                                        | 3           |
| 1.2 Tezin Amacı . . . . .                                                                                            | 5           |
| 1.3 Hipotez . . . . .                                                                                                | 6           |
| <b>2 Teorik Arka Plan</b>                                                                                            | <b>8</b>    |
| 2.1 Evrişimsel Sinir Ağları . . . . .                                                                                | 8           |
| 2.1.1 Evrişimsel Sinir Ağlarının Gelişimi . . . . .                                                                  | 9           |
| 2.1.2 Özgün Katmanlar . . . . .                                                                                      | 11          |
| 2.1.3 ESA için Örnek Ağ Yapıları . . . . .                                                                           | 16          |
| 2.1.4 Güncelleme Algoritmaları . . . . .                                                                             | 19          |
| 2.2 İstatistiksel Analiz Metotları . . . . .                                                                         | 21          |
| 2.2.1 Entropi . . . . .                                                                                              | 21          |
| 2.2.2 Kayıp Fonksiyonları ve Dağılım Metrikleri . . . . .                                                            | 22          |
| 2.3 Eğitici-siz Kümeleme Algoritmaları . . . . .                                                                     | 24          |
| 2.3.1 K-Ortalama Algoritması . . . . .                                                                               | 24          |
| 2.3.2 DBSCAN Algoritması . . . . .                                                                                   | 25          |
| 2.4 Sonuç . . . . .                                                                                                  | 27          |
| <b>3 Çok Sınıflı Geniş Ölçek Verilerin Az Sayıda Örneklem ile ESA Kullanılarak Modellenmesi ve Sınıflandırılması</b> | <b>28</b>   |
| 3.1 ESA Tabanlı Ağ Yapısı ve SMTS Algoritması . . . . .                                                              | 28          |

|          |                                                                                                      |           |
|----------|------------------------------------------------------------------------------------------------------|-----------|
| 3.1.1    | ESA Tabanlı Önerilen Model . . . . .                                                                 | 29        |
| 3.1.2    | SMTS Algoritması . . . . .                                                                           | 31        |
| 3.2      | Uygulama Alanı: Çevrim-dışı İmza Tanıma . . . . .                                                    | 32        |
| 3.2.1    | Literatür Özeti . . . . .                                                                            | 32        |
| 3.2.2    | Veri Seti ve Yöntemin Uyarlanması . . . . .                                                          | 35        |
| 3.2.3    | Başarım Değerlendirmesi . . . . .                                                                    | 37        |
| 3.3      | Sonuç . . . . .                                                                                      | 46        |
| <b>4</b> | <b>Az Sınıflı Yüksek Çözünürlüklü Görsel Verilerin İstatistiksel Modele Dayalı Sınıflandırılması</b> | <b>48</b> |
| 4.1      | Literatür Özeti . . . . .                                                                            | 48        |
| 4.2      | Yerel Histograma Dayalı Simetrik Kullback-Leibler Diverjansı . . . . .                               | 50        |
| 4.3      | Uygulama Alanı: Serviks Kanserinde Öncü Lezyonların Sınıflandırılması                                | 51        |
| 4.3.1    | Literatür Özeti . . . . .                                                                            | 53        |
| 4.3.2    | Katkılar . . . . .                                                                                   | 54        |
| 4.4      | Veri Seti ve Yöntemin Uyarlanması . . . . .                                                          | 55        |
| 4.4.1    | Veri Seti . . . . .                                                                                  | 55        |
| 4.4.2    | Ham Servikal KEP Görüntüsü için Ön İşleme Süreci . . . . .                                           | 56        |
| 4.4.3    | Dokuların İstatistiksel Öznitelikleri . . . . .                                                      | 58        |
| 4.4.4    | Kullback – Leibler Diverjansı Tabanlı Ağırlıklı $k$ -NN Algoritmasının Uyarlanması . . . . .         | 61        |
| 4.5      | Algoritmaların Analizi ve Başarım Karşılaştırmaları . . . . .                                        | 63        |
| 4.5.1    | Önerilen Algoritmaların Başarım Analizi . . . . .                                                    | 64        |
| 4.5.2    | İÖFY Algoritmasının Morfometrik Algoritmalar ile Kıyaslanması                                        | 66        |
| 4.6      | Sonuç . . . . .                                                                                      | 68        |
| <b>5</b> | <b>Yüksek Çözünürlüğe Sahip Görüntülerde Çakışık Nesnelerin Bölütlenmesi</b>                         | <b>69</b> |
| 5.1      | Literatür Özeti . . . . .                                                                            | 70        |
| 5.2      | Piksel Seyreltmeye Dayalı Yönlü $k$ -Ortalama Algoritması . . . . .                                  | 71        |
| 5.2.1    | Sentetik Problem Üzerinde YKA'nın Testi . . . . .                                                    | 77        |
| 5.3      | Uygulama Alanı: Çakışık Hücrelerin Bölütlenmesi . . . . .                                            | 79        |
| 5.3.1    | Literatür Özeti . . . . .                                                                            | 79        |
| 5.3.2    | YKA ile Problemin Çözümü . . . . .                                                                   | 81        |
| 5.4      | Sonuç . . . . .                                                                                      | 82        |
| <b>6</b> | <b>Sonuç ve Öneriler</b>                                                                             | <b>84</b> |
|          | <b>Kaynakça</b>                                                                                      | <b>86</b> |
|          | <b>Tezden Üretilmiş Yayınlar</b>                                                                     | <b>99</b> |

## SİMGE LİSTESİ

---

|                 |                             |
|-----------------|-----------------------------|
| $\mathcal{T}_i$ | i. sınıfa ait epitel kümesi |
| $\Re$           | Reel sayılar kümesi         |
| $\nabla$        | Gradyan                     |
| $\mathbb{S}_i$  | Epitel dokuya ait i. bölge  |

## KISALTMA LİSTESİ

---

|       |                                                     |
|-------|-----------------------------------------------------|
| BVA   | Büyük Veri Analitiği                                |
| BN    | Yığın Normalizasyonu (Batch Normalization)          |
| CAD   | Bilgisayar Destekli Tanı (Computed Aid Diagnosis)   |
| ÇEK   | Çapraz-Entropi Kaybı                                |
| ÇYH   | Çekirdek-Yoğunluk Haritası                          |
| ÇYÖ   | Çekirdek-Yoğunluk Özniteliği                        |
| DN    | Doku Normalizasyonu                                 |
| DÖ    | Derin Öğrenme                                       |
| EFV   | Gömülü Öznitelik Vektörü (Embedding Feature Vector) |
| ESA   | Evrşimsel Sinir Ağları                              |
| FC    | Tüm-bağlantı (Fully-Connected)                      |
| FCM   | Bulanık c-Ortalama (Fuzzy c-Means)                  |
| G-ESA | Geniş-ölçek Evrşimsel Sinir Ağı                     |
| GMM   | Gauss Karışım Model (Gaussian Mixture Model)        |
| HDFS  | Hadoop Distributed File System                      |
| H&E   | Hemotoksilen ve Eozin                               |
| HMM   | Saklı Markov Modelleri (Hidden Markov Models)       |
| HoG   | Histogramların Gradyanları (Histogram of Gradient)  |
| HPV   | İnsan Papilloma Virüsü (Human Papilloma Virus)      |
| KEP   | Küçük Epitel Parçası                                |
| KLD   | Kullback–Leibler Diverjansı                         |
| k-NN  | k-En Yakın Komşu (k-Nearest Neighbour)              |
| MLP   | Çok Katmanlı Perseptron (Multi-Layer Perceptron)    |

|      |                                                                       |
|------|-----------------------------------------------------------------------|
| IoT  | Nesnelerin İnterneti (Internet of Things)                             |
| OHE  | Tek Sınıf Tabanlı İkili kodlama (One-Hot Encoding)                    |
| OKF  | Olasılık Kütle Fonksiyonu                                             |
| PCA  | Temel Bileşen Analizi (Principal Component Analysis)                  |
| ReLU | Doğrultulmuş Doğrusal Birim (Rectified Linear Unit)                   |
| SAR  | Sentetik Açıklık Radarı                                               |
| SGD  | Stokastik Gradyan İniş (Stochastic Gradient Descent)                  |
| SIFT | Ölçek-bağımsız Öznitelik Dönüşümü (Scale-Invariant Feature Transform) |
| SIU  | Sinyal İşleme ve Uygulamaları Kurultayı                               |
| SMO  | Sıralı Küçültme En İyilemesi (Sequential Minimization Optimization)   |
| SMTS | Sınıf Merkezi Tabanlı Sınıflandırıcı                                  |
| SVM  | Destek Vektör Makinesi (Support Vector Machine)                       |
| UPS  | Uyarlamalı Piksel Seyreltme                                           |
| VGG  | Visual Geometry Group                                                 |
| YHÖ  | Yerel Histogram Özniteliği                                            |
| YHA  | Yerel Histogram Algoritması                                           |
| YKA  | Yönlü k-Ortalama Algoritması                                          |

## ŞEKİL LİSTESİ

|           |                                                                                                                                                                                                                      |    |
|-----------|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|----|
| Şekil 2.1 | YSA'dan ESA'ya geçiş: (a) YSA'ya ait tüm-bağlantı yapısı (b) Seyrek bağlantı tanımı (c) Ağırlık paylaşımı sonrası katmanlar arası bağlantı                                                                           | 9  |
| Şekil 2.2 | Evrişim operatörünün üç boyutlu veriler üzerindeki işlemi                                                                                                                                                            | 11 |
| Şekil 2.3 | Yığın normalizasyonunun işlem sırasının iki boyutlu gösterimi: (a) Aktivasyon öncesi yığın dağılımı (b) Yığın normalizasyonu (c) Yığın istatistiklerinin aktivasyona göre ayarlanması                                | 15 |
| Şekil 2.4 | Görüntü sınıflandırma problemine yönelik tasarlanan örnek bir ESA yapısı                                                                                                                                             | 15 |
| Şekil 2.5 | GoogleNet ile birlikte ileri sürülen <i>inception</i> modülüne dair bir örnek                                                                                                                                        | 17 |
| Şekil 2.6 | VGG-16 ağının blok şeması                                                                                                                                                                                            | 17 |
| Şekil 2.7 | ResNet-tipi ağlarda kullanılan atlamalı bağlantı mekanizmasının blok şeması. $F(\mathbf{X})$ farklı tasarımlarda olabilmektedir                                                                                      | 18 |
| Şekil 2.8 | Adam algoritmasının diğer güncelleme algoritmaları ile karşılaştırmaları: Adam'ın uyarlamalı moment özelliği daha etkili bir öğrenme sağlamaktadır [61]                                                              | 20 |
| Şekil 2.9 | DBSCAN algoritmasının çalışmasına $N_p = 4$ ve $R = 2$ parametreleri için bir örnek. A ve B noktaları R yarıçapı içinde $N_p$ komşuya sahip oldukları için küme noktası olurken C noktası gürültü olarak belirlenir. | 25 |
| Şekil 3.1 | Tablo 3.1 içindeki her <i>conv</i> bloğunun iç yapısı: <i>conv</i> bloğunun başarımı BN bloğunun olup olmamasına göre değişkenlik gösterdiği için sınırı kesikli çizgilerle çizilmiştir                              | 30 |
| Şekil 3.2 | Kullanılan veri setlerinden alınan bazı örnekleri: İlk satır GPDS-4k'ya, ikinci satır CEDAR'a ve üçüncü satır ise MCYT'ye karşılık gelir                                                                             | 35 |
| Şekil 3.3 | Ön-işleme adımı ve G-ESA'nın çevrim-dışı imza tanıma için uyarlanması ile elde edilen sınıflandırıcı modeli                                                                                                          | 36 |
| Şekil 3.4 | MCYT ve CEDAR veri setlerinin eğitim fazındaki öğrenme eğrileri                                                                                                                                                      | 40 |
| Şekil 3.5 | MCYT, CEDAR ve MCYT+CEDAR veri setleri üzerinden eğitim aşamasında kişi başına ilk $\theta \in \{2, 4, 6, 8, 10\}$ imzanın alınması ile eğitilen modellerin test başarımlarının kıyaslanması                         | 41 |

|           |                                                                                                                                                                                                                                                                                                                                                                                                                                                            |    |
|-----------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|----|
| Şekil 3.6 | Taban, BN ekli ve SMTS uygulanan ağların başarımlarını karşılaştırmaları: (a) ve (d): sırasıyla 25-75 için Taban - BN ekli ağlar ile BN ekli - SMTS uygulanmış ağların sonuçları. (b) ve (e) ise, 50-50 için Taban - BN ekli ağlar ile BN ekli - SMTS uygulanmış ağların sonuçlarını ifade eder. Son olarak (c) ve (f): $\alpha$ değerlerine göre 25-75 ve 50-50 için G-ESA_v2 sonuçlarıdır . . . . .                                                      | 45 |
| Şekil 3.7 | Taban ve BN ekli ağların eğitim hızlarının karşılaştırılması: (a) ve (b): sırasıyla 25-75 ve 50-50 oranları için yakınsaklık hızı analizidir. BN eklenen ağların bir eğitim devrindeki zaman tüketimi BN eklenmemiş ağlara nazaran daha fazla zaman tüketimi olmasada (Tablo 3.5), tüm eğitim sürecindeki öğrenme hızları, eklenmemiş ağlardan çok daha yüksektir . . . . .                                                                                | 46 |
| Şekil 4.1 | Servikal lezyonların sınıflandırılması için önerilen istatistiksel özniteliklere dayalı algoritmanın blok diyagramı . . . . .                                                                                                                                                                                                                                                                                                                              | 53 |
| Şekil 4.2 | Veri seti içinde kullanılan kavramların açıklanması: (a) Yüksek çözünürlüklü H&E slayt görüntüsü, (b) Tanı için incelenecek olan epitel doku, (c) Küçük Epitel Parçaları (KEP) . . . . .                                                                                                                                                                                                                                                                   | 56 |
| Şekil 4.3 | Veri setindeki dokuların sınıflara göre dağılımı: (a) 471 Normal, (b) 240 CIN-1, (c) 107 CIN-2, and (d) 139 CIN-3 . . . . .                                                                                                                                                                                                                                                                                                                                | 57 |
| Şekil 4.4 | 4-derinlikli SegNet ağının eğitimi için üretilen veri setinden örnekler: 100 tane $512 \times 512$ büyüklüğündeki görüntüler ile eğitilen SegNet, sadece hücreleri bulmak için kullanılmaktadır . . . . .                                                                                                                                                                                                                                                  | 57 |
| Şekil 4.5 | Ön-işleme adımı sonrası çıktılar: Doku Normalizasyonu ve ortanca filtre kullanılarak RGB KEP görüntüsü (a), Gri-seviye görüntüye dönüştürülmüştür (b). SegNet kullanılarak Çekirdek Yoğunluk Haritası (c) elde edilmiştir. Mavi: bazal membran, Kırmızı: Epitel yüzey, and Sarı: Papilla sınırları . . . . .                                                                                                                                               | 58 |
| Şekil 4.6 | $\mathbf{P}_j \in \mathbb{R}^{N_w \times N_w}$ ızgara görüntüsü (magenta) dokuya ait katmanlar üzerinden elde edilir. Bazal sınıra olan uzaklık $d_{ds}$ , yeşil ok ile epitel yüzeye olan sınır $d_{up}$ ise turuncu ok ile işaretlenmiştir. $r$ değerinin aralığına göre ( $[0, 1/4)$ , $[1/4, 3/4]$ , $(3/4, 1]$ ) $\mathbf{P}_j$ , basal ( $\mathbb{S}_B$ ), orta ( $\mathbb{S}_C$ ) or üst ( $\mathbb{S}_U$ ) katman olarak sınıflandırılır . . . . . | 60 |
| Şekil 4.7 | CIN seviyelerine göre katmanlardan elde edilen histogram grafikleri                                                                                                                                                                                                                                                                                                                                                                                        | 60 |
| Şekil 4.8 | Parametre uzayında ızgara tabanlı en iyileme sonuçları: $k = \{1, 3, 5, 7, 9, 11, 13, 15\}$ ve $\gamma = \{0, 0, 02, \dots, 1, 50\}$ değerleri arasında en iyi doğruluk başarımlarını $75,55 \pm 4,02$ , norm: $\ell_1$ $\gamma$ -Parameter: 0,3 ve $k$ -NN: 3 konfigürasyonu ile elde edilmiştir . . . . .                                                                                                                                                | 65 |
| Şekil 4.9 | Veri setinin önerilen metrik uzayından t-SNE algoritması ile iki boyutlu düzleme haritalandırılarak görselleştirilmesi . . . . .                                                                                                                                                                                                                                                                                                                           | 65 |

|            |                                                                                                                                                                                                                                                                                                                                                                                              |    |
|------------|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|----|
| Şekil 4.10 | İÖFY algoritmasının farklı test oranları altındaki başarımların analizi . .                                                                                                                                                                                                                                                                                                                  | 66 |
| Şekil 5.1  | YKA'nın işleyiş sıralaması . . . . .                                                                                                                                                                                                                                                                                                                                                         | 72 |
| Şekil 5.2  | Sınır piksellerinin uyarlamalı olarak seyreltilmesi. $\theta_{\text{eşik}} = 10^\circ$ ve $d_{\text{eşik}} = 5$ olarak seçilmiştir . . . . .                                                                                                                                                                                                                                                 | 73 |
| Şekil 5.3  | $\theta_t$ eşik değerine göre YKA'nın uzaklık değerlendirmesi (Sarı: En yakın, Mavi: En uzak, Yön vektörü: $\mathbf{d} = [1, 0]^T$ ): (a) $k$ -Ortalama algoritması isotropik olarak mesafe değerlendirmesi yaparken YKA, (b) $\theta_t = 60^\circ$ eşik değerinde $120^\circ$ , (c) $\theta_t = 30^\circ$ 'ken $60^\circ$ 'lik bir kapsam açıklığına sahiptir . . . . .                     | 75 |
| Şekil 5.4  | Sentetik olarak oluşturulan bir nesne: Yapısında 5 tane çakışık nesnenin bulunduğu şekil (a) üzerinden ilk olarak sınır pikselleri çıkartılmıştır (b). Sonrasında UPS algoritması kullanılarak sınır pikselleri indirgenmiştir (c). Bu sayede algoritma 1690 tane sınır pikseli yerine 260 örnek üzerinden problemi çözmektedir . . . . .                                                    | 77 |
| Şekil 5.5  | Öbeğin sınır piksellerinin sınıflandırılması: YKA, girdi olarak seyreltilmiş sınır pikselleri ile öbek içinde kaç nesnenin bulunduğu bilgisin aldıktan sonra (a)'yı üretmektedir. (b) ise $k$ -Ortalama algoritmasının çıktısıdır . . . . .                                                                                                                                                  | 78 |
| Şekil 5.6  | Sınır piksellerinin sınıflandırılmasında aykırı sınıflandırılmış örnekler çıkmaktadır. Bu örnekler elips uydurma algoritmasına bozucu etkileri olur (a). Bunun giderilmesi için DBSCAN ile aykırı değerler elenmiştir. Bu aşamadan sonra YKA'nın nihayi çıktısı (b) elde edilir. (c) ise aynı problem için su kuyusu algoritmasının çıktısıdır 78                                            |    |
| Şekil 5.7  | YKA'nın [149] ve su-kuyusu algoritmaları ile karşılaştırılması: İlk sırada florasen hücrelerden elde edilmiş ikili görüntüler bulunmaktadır. İkinci ve üçüncü sıralarda ise [149] ve su-kuyusu algoritmalarının çıktıları ve dördüncü satırda da YKA'nın sonuçları gösterilmiştir. (a-c) olan hücreler yüksek çözünürlüğe sahipken (d) ve (e) düşük çözünürlük ile elde edilmiştir . . . . . | 82 |

## TABLO LİSTESİ

|           |                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                             |    |
|-----------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|----|
| Tablo 3.1 | Ağların parametre bilgileri   @: Kullanılan filtre sayısı $p, s$ : Sıfır ekleme ve filtre kaydırma miktarları . . . . .                                                                                                                                                                                                                                                                                                                                                                                                                                                                                     | 29 |
| Tablo 3.2 | Ağların eğitimi için kullanılan üst-parametre listesi . . . . .                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                             | 39 |
| Tablo 3.3 | MCYT ve CEDAR veri setleri üzerinden rasgele oluşturulan beş farklı alt kümenin kullanılması ile elde edilen sonuçların ortalamaları . .                                                                                                                                                                                                                                                                                                                                                                                                                                                                    | 39 |
| Tablo 3.4 | MCYT, CEDAR ve MCYT+CEDAR veri setlerinde bulunan kişilerden ilk $\theta \in \{2, 4, 6, 8, 10\}$ imzanın eğitim için alınarak yapılan testlerin sonuçları . . . . .                                                                                                                                                                                                                                                                                                                                                                                                                                         | 41 |
| Tablo 3.5 | GPDS-4k veri seti için ağların detaylı karşılaştırılması. Başarım metriklerindeki $\% \Delta 0,1$ 'lik değişim, 25-75 oranlarında 75 imza ve 50-50 oranları için 48 imzaya karşılık gelir. <i>Eğitim</i> sütununda [sn] cinsinden bulunan değerler, ağların iki farklı oran için bir eğitim devrinde tükettikleri zaman miktarlarıdır. <i>Test</i> sütunu ise ağların 5k örnek ile test edilmesi için gereken süreyi ifade eder. <i>Ağ [MB]</i> parametresi, ağın toplam parametre büyüklüğünü verirken, <i>Veri [MB]</i> ise eğitim fazında ağın ihtiyaç duyduğu veri işleme alanı ile ilgilidir . . . . . | 43 |
| Tablo 4.1 | Patologların CIN ve SIL tabanlı derecelendirme sistemlerine göre karışıklık matrisi: Gözlemciler arası değişkenlik kesin tanıya göre yüzde olarak gösterilmiştir . . . . .                                                                                                                                                                                                                                                                                                                                                                                                                                  | 64 |
| Tablo 4.2 | Önerilen algoritmaların $\gamma = 0,5$ $k = 3$ parametreleri için başarım analizi . . . . .                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                 | 64 |
| Tablo 4.3 | İÖFY'nin farklı sınıflandırıcılar ile sınıflandırılan morfometrik öznitelikler ile karşılaştırılması . . . . .                                                                                                                                                                                                                                                                                                                                                                                                                                                                                              | 67 |
| Tablo 4.4 | İÖFY ve morfometrik öznitelikler arasında en iyi sonucu veren MÖV-SVM <sub>poly</sub> algoritmalarının karışım matrisleri . . . . .                                                                                                                                                                                                                                                                                                                                                                                                                                                                         | 67 |

# Geniş Ölçekli Veriler Üzerinde Sınıflandırma ve Bölütleme Amaçlı Evrimsel Sinir Ağı ve İstatistiksel Modellerin Geliştirilmesi

Nurullah ÇALIK

Elektronik ve Haberleşme Mühendisliği Anabilim Dalı

Doktora Tezi

Danışman: Prof. Dr. Lütfiye DURAK ATA

Bilginin yapı taşı olan veri, üretim kaynaklarının hacmi, çeşitliliği ve hızı nedeniyle devasa bir büyüklüğe ulaşmıştır. Meydana gelen bu nicelik içinden nitelikli faydanın elde edilmesi ise ayrı bir disiplin olarak ele alınması gereken bir yapıya dönüşmüştür. Bu sebeple yapay zeka, iletişim ağları, güvenlik, depolama ve gizlilik gibi birçok alt disiplini bünyesinde bütünleyen *Büyük Veri* olgusu doğmuştur. Veri biliminin ihtiyaç duyduğu tek çatı altında birleşme ile, bilgi birkez daha çeşitli sahalarda önemini perçinlemiştir. Büyük verilerin işlenmesi ile biyomedikal alanda doktorlar daha nesnel kararlar verebilirken, hastalara ait veriler daha sistematik bir şekilde planlanabilmektedir. Finans sektöründe yatırımlar sezgisel adımlar yerine, piyasaların duyarlılığının kesitirilmesi ile elde edilen veriler doğrultusunda yapılmaktadır. Endüstriyel uygulamalarda organizasyon planlamaları yine iş zekası çerçevesinde geliştirilmektedir. Bunlarla birlikte nesnelerin interneti ve bulut bilişim gibi teknolojiler, geniş-ölçek seviyesindeki veri akışını ve yönetimini mümkün kılarak büyük veri olgusunu geniş bir uygulama yelpazesine yaymaktadır.

Büyük veri kavramı, bütünleşik veya dağıtık verinin saklanması ve organize edilmesinin yanı sıra analitik olarak işlenmesini de içermektedir. Veri madenciliği alanında geliştirilen geleneksel algoritmalar ve donanımlar, işlenmesi gereken örnek miktarının üssel olarak artması nedeniyle hesaplama yükü ve depolama ihtiyaçları bakımından yetersiz kalmaktadır. Yeni nesil geliştirilen yöntemlerin en temel özelliklerinden birinin ölçeklenebilir olması kaçınılmaz olmuştur. Bu bağlamda,

öne sürülen algoritmaların paralel çalışma teknikleri veya grafik işlemcilerinin güçlü mimarileri üzerinden geliştirilmesiyle büyük veri analitiği etkin araçlara kavuşmuştur. Özellikle derin öğrenme temelinde üretilen modellerin çok fazla veri ile eğitilmesi ile sergiledikleri başarımlar, makine öğrenmesi ve büyük veri kombinasyonunun çığır açıcı etkilerini gözler önüne sermiştir.

Eğiticili sınıflandırma problemlerinde sınıf başına düşen eğitim örneği ne kadar fazla ise modellerin veriyi betimleme kabiliyetleri o derece iyi olmaktadır. Bu çerçevede, geniş-ölçek veriler büyük bir fırsat sağlamaktadır. Diğer taraftan çok sınıflı ( $\geq 1000$ ) ve sınıf başına az örnek ( $\leq 10$ ) kullanılarak modellerin eğitilmesi, üzerinde çalışılması gereken önemli bir konudur. Bu tanım, gerçek hayatta biyometrik verilerin sınıflandırılması problemi ile yakından ilgilidir. Dolayısıyla tez kapsamında, eğiticili öğrenme alanına yönelik olarak evrimsel sinir ağı temelinde özgün bir sınıflandırma modeli önerilmiştir. Az sınıflı veri setleri, içerdikleri verinin çok fazla ve yoğun olması nedeniyle istatistiksel yöntemlerin analiz menziline girebilmektedirler. Özellikle, veri setindeki sınıflar arasında ayırt edici öznelilikler istatistiksel olarak çıkartılabilmesi eğitim sürecini hızlandırmaktadır. Bu doğrultuda, eğiticili öğrenme içinde belirlenen bu alana yönelik yerel histogram tabanlı bir yöntem ortaya koyulmuştur. Yerel histogramlar öznelilik olarak ele alınmış ve örnekler arası benzerlik metriği için simetrik Kullback-Leibler Diverjansı kullanılmıştır. Sınıflandırmaya yönelik ağırlıklı K-En Yakın Komşu algoritması tercih edilmiştir. Bu adımlar bütününde oluşturulan algoritma, serviks dokularının sınıflandırılması probleminde test edilmiştir. Ayrıca, eğiticişiz öğrenme dahilinde var olan bölütleme problemi kapsamında, yüksek çözünürlüğe sahip görüntülerde çakışık nesnelerin ayrıştırılması konusu ele alınmış, uyarlamalı olarak veri azaltma tekniği temelinde etkili bir algoritma geliştirilmiştir. Bu algoritma, çakışık nesneleri sınır pikselleri üzerinden çözebilecek şekilde k-Ortalama algoritmasının temelini oluşturan kayıp fonksiyonunun yeniden tanımlanması ile ortaya konmuştur. Geliştirilen yöntem histopatoloji alanında önemli bir problem olan çakışık hücrelerin bölütlenmesine uyarlanarak literatüre katkılar sunulmuştur.

**Anahtar Kelimeler:** Derin öğrenme, evrimsel sinir ağları, geniş ölçekli veri, istatistiksel metotlar, çakışık nesnelerin bölütlenmesi.

# **Development of Convolutional Neural Network and Statistical Models for Classification and Segmentation on Large-Scale Data**

Nurullah ÇALIK

Department of Electronics and Communications Engineering

Doctor of Philosophy Thesis

Advisor: Prof. Dr. Lütfiye DURAK ATA

The data that constitutes the building block of information has reached a huge size due to the volume, variety and speed of production resources. Obtaining qualified benefit from this quantity turned into a structure that should be considered as a separate discipline. For this reason, the concept of *Big Data*, which integrates many sub-disciplines such as artificial intelligence, communication networks, security, storage and privacy, has emerged. With the unification under the single roof that data science needs, information has once again reinforced its importance in various fields. With the processing of large data, doctors can make more objective decisions in the biomedical field, while patient data can be planned in a more systematic way. Investments in the financial sector are made in line with the data obtained by crossing the sensitivity of the markets instead of intuitive steps. Organizational planning in industrial applications is also developed within the framework of business intelligence. In addition, technologies such as the Internet of Things and Cloud Computing make it possible for large-scale data flow and management to spread large data phenomena across a wide range of applications.

The concept of big data includes the storage and organization of integrated or distributed data as well as analytical processing. Traditional algorithms and equipment developed in the field of data mining are insufficient in terms of calculation load and storage requirements due to the exponential increase in the amount of sample to be processed. It is inevitable that one of the most basic features of the new generation developed methods would be scalable. In this context, large-scale

data analytics have gained effective tools through the development of proposed algorithms over parallel operating techniques or powerful architectures of graphic processors. In particular, the performances of models produced on the basis of deep learning has demonstrated the groundbreaking effects of machine learning and big data combination.

The more training examples per class in supervised classification problems, the better the ability of the models to describe the data. In this context, large-scale data provides a great opportunity. On the other hand, training of models using multi-class ( $\geq 1000$ ) and less sample per class ( $\leq 10$ ) is an important issue that needs to be studied. This definition is closely related to the problem of classification of biometric data in real life. Therefore, within the scope of the thesis, an original classification model based on convolutional neural network is proposed for the area of the supervised learning. Less-class data sets can enter the analysis range of statistical methods because the data they contain is too much and dense. Specifically, the fact that discriminative features can be extracted statistically between the classes in the data set accelerates the training process. In this respect, a local histogram-based method has been proposed for this area identified in supervised learning. Local histograms are considered as features and symmetrical Kullback-Leibler Divergence is used for the similarity metrics between samples. Weighted k-Nearest Neighbor algorithm is preferred for classification. The algorithm, which is created in the whole of these steps, is tested on the problem of classification of cervical tissues. In addition, within the scope of segmentation problem in unsupervised learning, the separation of overlapping objects in high resolution images has been dealt with and an adaptive algorithm based on data reduction technique has been developed. This algorithm is introduced by redefining the loss function which forms the basis of the k-mean algorithm to solve overlapping objects over boundary pixels. The developed method is adapted to segmentation of overlapped cells which is an important problem in histopathology field and contributions to the literature are presented.

**Keywords:** Deep learning, convolutional neural network, large-scale data, statistical methods, overlapped object segmentation.

# 1

## Giriş

---

İnsanlık tarihi boyunca bilgi, dünyayı ve toplumları şekillendiren önemli bir kaynak olmuştur. Teknolojik gelişmeler, asırlar boyu üretilen ve nesiller boyunca aktarılan muazzam bilgi mirasının üzerine bina edilmişlerdir. Son yüzyıl boyunca ortaya çıkan teknolojik ilerlemeler, sahip oldukları gücün perde arkasında olan kaynağını gün yüzüne çıkartarak bilgi çağının başlamasına sebep olmuşlardır. Bu çağda bilginin serbest dolaşım hızına en büyük katkıları sağlayan internet ve hücrel haberleşme sistemleri, farklı fikirlerin farklı zihinler ile buluşmasının önünü açmıştır. Bu sayede zincirleme veri reaksiyonu tetiklenmiştir.

Bilgi çağının yapı taşı olan veri, günümüzde tarih boyunca hiç olmadığı kadar öneme sahip olmuştur. Özellikle dijital çağın başlaması ile veri üretimi üssel olarak artmaktadır. Şuanki veri miktarı, son beş yıl içinde kopyalanan veya üretilen veri miktarı yaklaşık dokuz katına çıkmış ve bu rakam en az iki yılda bir ikiye katlanarak devam edeceği öngörülmektedir [1]. Veri miktarlarının devasa boyutlara ulaşması ile birlikte *Büyük Veri* kavramı ortaya çıkmıştır. Veri analitiği konularının ele alındığı büyük veri zemini, bilgi işleme alanlarında öncü bir terim haline gelse de standart bir tanıma halen bulunmamaktadır [2, 3]. Büyük veri ifadesiyle, sadece veri miktarının büyüklüğü değil, aynı zamanda bu alana yönelik geliştirilen çözümler de kastedilmektedir. Geleneksel veri analitiği teknikleri ve istatistiksel yöntemler, işlem hızları ve bellek gereksinimleri gibi nedenlerden dolayı yetersiz kalmaktadır. Bu sebeple, büyük veri kapsamında yenilikçi algoritmaların üretilmesine ihtiyaç duyulmuştur. Diğer taraftan mevcut veri organizasyonu şemaları yerine dağıtık sistemler geliştirilmiş, donanımsal olarak güçlü tepki hızına sahip işlemciler ile birlikte grafik işlemcilerinin paralel veri işleme yetileri önemli bir rol kazanmıştır. Ayrıca ızgara tabanlı hesaplama (grid-computing) teknikleri, bulut tabanlı bilgi depolama sistemleri ile veri analitiği için yeni bir çağ başlamıştır [4].

Tüm bu kavramların tek çatıda toplandığı büyük veri ekosistemi için bazı temel özellikler tanımlanmıştır. Üretilen veriler ve geliştirilen algoritmalar, *yüksek hacim* (volume), *çeşitlilik* (variety) ve *hız* (velocity) (3V) temelinde ele alınmaktadır

[5]. Günümüzde sosyal medya, kredi kartları, akıllı arabalar veya cep telefonları ile yoğun bir bilgi üretimi söz konusudur. Bu büyüklük *yüksek hacim* özelliği altında ele alınmakta ve veri yönetiminin sağlanması için dağıtık mimarili sistemler geliştirilmektedir. Bilgi miktarı ile birlikte *çeşitlilik* özelliği de üretilen veriyi tanımlamak için gereken kavramlardan birisidir. Gelişen teknoloji, sanal gerçeklik, hologram veya LIDAR gibi alışık olunmayan veri tiplerini sunmaktadır. Yeni üretilen veri formatları çeşitlilik altında incelenmektedir. Veriyi betimleyen diğer bir temel özellik ise onun üretim hızıdır. Çok çeşitli alanlardan sürekli bir bilgi akışı olmakta ve bu konu büyük verinin *hız* özelliği başlığında ele alınmaktadır. Buna örnek olarak nesnelerin interneti (internet of things (IoT)) verilebilir [6]. Veri miktarının artması ile bu üç temel özelliğin yanında *doğruluk* veya *güvenilirlik* (veracity) özelliği de araştırma alanları arasına girmiştir. Üretilen bilginin ne derece güvenilir olduğunun irdelendiği bu alan çerçevesinde büyük veri özellikleri 4V olarak tanımlanmıştır [7, 8]. Veri miktarının çok olması yeni fırsatların doğmasına ve etkili algoritmaların gelişimine imkan sağlamıştır. Bu noktada ise verinin *değeri* (value) önemli bir kriter olarak karşımıza çıkar. Üretilen bilginin birçoğu gereksiz ve kullanıma elverişsizdir. Belli bir probleme yönelik geliştirilen çözümlerin temellendirildiği veri, ilgili problemi genelleyebileme değerinde olmalıdır. Bu bağlamda büyük veri analitiğinin (BVA) temeli 5V olarak ortaya koyulmuştur [9]. Öne sürülen bu tanımlamalar çerçevesinde BVA, birçok farklı alanda kendine yer bulmuştur.

Günümüzde IoT ile birlikte BVA, endüstriyel alanlardan finans sektörüne, ticari faaliyetlerden askeri uygulamalara kadar birçok alana nüfuz etmiştir. Akıllı şehir konsepti çerçevesinde şehir planlamaları yeniden organize edilerek taşımacılık, ulaşım ve yaşam alanları gibi konularda yenilikçi çözümler sunulmuştur [10, 11]. Ayrıca BVA, kentsel alt yapı işlerinde [12], sosyal ağ uygulamalarında [13] ve acil durum iletişim ağlarında kullanılmaktadır [14]. Özellikle akıllı ormancılık alanında doğal hayatın izlenmesi ve bitki örtüsünün değerlendirilmesi gibi konularda etkin olarak tercih edilmektedir [4]. BVA'nın bir diğer önemli kullanım alanı ise finans sektörüdür. Borsa üzerinden üretilen veriler üzerinden hisse senedi hareketleri kestirilmeye çalışılmakta [15] veya piyasaların duyarlılık analizi (sentiment analysis) yapılabilmektedir [16]. Bunlarla birlikte sağlık alanında giyilebilir aygıtlar sayesinde toplanan verilerden hasta takip ve teşhis uygulamalarında çığır açıcı etkiler oluşturmıştır [17]. Uygulama alanındaki bu çeşitlilik karşısında BVA, sorun çözme yeteneğini ve esneklik kabiliyetini arka planında yatan güçlü algoritmalara ve geliştirme platformlarına borçludur.

Devasa boyutlarda ve çok farklı kanallardan gelen verinin saklanması ve işlenmesi mevcut veri tabanı sistemlerini aşan bir durumdur. Veriyi tek disk içinde sakalamak yerine birden fazla sunucuda barındırıp, sunucular arasındaki işleyişin planlanması için bir organizasyon çerçevesi gerekmektedir. Bu amaç doğrultusunda geliştirilen

Hadoop platformu verileri bloklar olarak *Hadoop Distributed File System* (HDFS) formatında farklı diskler içinde saklamaktadır [17]. Hadoop, tanımlanan işi *bölümleme* ve *bütünleme* (map-reduce) mantığında işlemektedir [18]. Bölümleme adımı işi, alt parçalara ayrılarak farklı birimlerde işlenir ve ardından bütünleme adımı birim çıktıları birleştirilerek sonuç oluşturulur. Birbirinden bağımsız birimlerde veri, yığınlar (batch) olarak işlendiği için bu strateji paralel işlemeye olanak sağlar. Analiz platformlarından sonra BVA için diğer bir kritik konu ise var olan veri madenciliği tekniklerinin hızlandırılmasıdır. Literatürde eğiticili ve eğitici-siz öğrenme konularında üstün başarımlar sağlamış olan algoritmalar BVA için ölçeklendirilmişlerdir. Ayrıca yapay sinir ağlarının (YSA) yeni nesil versiyonu olan evrişimsel sinir ağları (ESA) ve derin öğrenme (DÖ) teknikleri gelişen GPU teknolojisi sayesinde büyük verilerin işlenmesinde yükselen bir eğilim olmuşlardır.

Tüm bu katkılar ışığında BVA, farklı disiplinlerin eş zamanlı gelişimi ve katkısı ile ortaya çıkmıştır. Bu yeni disiplin, veri organizasyonu için ortaya atılan çerçevelerin yanı sıra çözülmesi istenen probleme yönelik geliştirilen algoritmalar ve bu algoritmaların çalışacağı donanımlara kadar birçok sac ayağı üzerine kurulmuştur. Geliştirilen bu algoritmalar içinde eğiticili ve eğitici-siz öğrenme alanında olanlar yapay zeka teknolojisinin lokomotifidir. Bundan dolayı tezin yol haritası bu alanlar doğrultusunda kurgulanmış ve literatür taraması da yine bu bağlamda yapılmıştır.

## 1.1 Literatür Özeti

Eğiticili öğrenme alanında en yaygın kullanılan algoritmalar örnek olarak Destek Vektör Makinesi (Support Vector Machine (SVM)) verilebilir. SVM sınıflar arasında en yüksek marjı sağlayacak üst-düzlemi (hyper-plane) iç-bükey (convex) en iyileme (optimization) teknikleri ile bulmaktadır. Doğrusal olarak ayrıştırılamayan veriler çekirdek teknikleri sayesinde üst-uzaya haritalandırılarak sınıflandırılabilir. SVM algoritmasının dayandığı teorik temeller, algoritmanın genelleme yetisinin gücünü oluşturmaktadır. Diğer taraftan SVM  $O(n^3)$  işlem karmaşıklığı nedeniyle büyük veriler için pratik bir yöntem değildir. Buna yönelik, SVM'in en iyilemesi için bölümleme ve bütünleme stratejisinde farklı uygulamalarda kullanılan sıralı küçültme en iyilemesi (sequential minimization optimization (SMO)) geliştirilmiştir [19–21]. Yeni önerilen SMO, veri seti parçaları üzerinde üst-düzlem hesabı yaptıktan sonra elde edilen alt-modellerden üst modeli oluşturmaktadır. SVM tabanlı geliştirilen diğer bir yöntem ise [22]'de ileri sürülen PEGASOS algoritmasıdır. SVM'in birincil (primal) formundan hareketle geliştirilen algoritma, kayıp fonksiyonunu stokastik gradyan azalımı (Stochastic Gradient Descent (SGD)) ile minimize etmektedir. İlk olarak veri

yığınlara bölünür ve sonrasında her yığın için model parametreleri hesaplanmaktadır. Ardışıl hesaplanan parametre değerleri güncellenerek arasında güncelleme yapılarak genel model kestirilmektedir [23]. çalışmasında Rebentrost ve ark. kuantum mekaniği tabanlı bir SVM algoritması öne sürmüşlerdir. Yapılan bu çalışma ile SVM'in hesaplama karmaşıklığı  $O(\log(DN))$  seviyelerine getirilmiştir. Burada  $D$  öznitelik ve ötürün boyutu,  $N$  ise örnek sayısıdır. [24] çalışmasında tek sınıflı SVM sınıflandırıcıları için çok büyük veri setlerine yönelik bir yöntem geliştirilmiştir. Çekirdek uzayında hedef kümesinin etrafına en küçük yarı çapa sahip üst-küre (hypersphere) çevrelemesi üzerine çalışan algoritma, eğitim verisi olarak tüm seti kullanmak yerine öz-küme (core-set) örnekleri ile problemi hızlıca çözmektedir.

SVM haricinde yaygın kullanılan bir diğer yöntem ise  $k$ -En Yakın Komşu ( $k$ -Nearest Neighbour ( $k$ -NN)) algoritmasıdır.  $k$ -NN *tembel* öğrenici (lazy learner) diye tanımlanan veriye bir model oturtmak yerine eğitim için veri setinin tamamını kullanmaktadır. Test örneğinin sınıflandırılması, eğitim kümesindeki örnekler ile arasındaki mesafeye bağlıdır. En yakın  $k$  örneğin sınıfına bakıldıktan sonra baskın olan sınıf doğrultusunda test örneği için karar verilir.  $k$ -NN algoritması, eğitim örnekleri üzerinden model uydurma ile ilgilenmediği için eğitim fazında her hangi bir işlem yüküne ihtiyaç yoktur. Diğer taraftan test fazında test örneği eğitim kümesindeki tüm örnekler ile karşılaştırılması büyük veri setleri için problem oluşturur. Buna yönelik Deng ve ark. [25] çalışmasında  $k$ -Ortalama algoritması kullanarak veri setini her sınıf için belli öbeklere toplamayı önermişlerdir. Her öbek için hesaplanan merkezler üzerinden  $k$ -NN test işlemini sürdürmektedir. Bu sayede karşılaştırma sayısı büyük ölçüde azaltılmıştır. [26] çalışmasında ise  $k$ -NN algoritması, Spark platformu kullanılarak bölümle ve bütünle stratejisi temelinde yeniden tanımlanmıştır. Paralelize edilen algoritma veri yığınları üzerinden çalışmaktadır. İteratif olarak sonuca ulaşan algoritma 11 milyonluk bir veri seti ile test edilmiştir. Ayrıca makalede, algoritmanın Hadoop versiyonuna nazaran 10 kat daha hızlı çalıştığı rapor edilmiştir.

Eğitici-siz öğrenme alanında ise  $k$ -Ortalama algoritması [27] kendisinden sonra birçok algoritmaya ilham olmuş köklü bir algoritmadır [28–31].  $k$ -Ortalama algoritması önceden tanımlanan sınıf sayısı kadar veriyi porsiyonlara böler. Bu bölme işlemini örneklerin sınıf merkezlerine olan uzaklıklarının toplamı ile ifade edilen bir kayıp fonksiyonunu minimize ederek yapmaktadır. İteratif olarak çalışan algoritma her iterasyonda tüm veri setini taramaktadır. Bu sebeple geniş ölçekli veriler söz konusu olduğunda yüksek hesaplama maliyetine ihtiyaç duyulur. Cai ve ark. [32] çalışmasında çok-görüşlü  $k$ -Ortalama algoritmasını önermişlerdir. Klasik algoritmayı graf matrisleri ve farklı uzaklık fonksiyonları üzerinden tanımladıktan sonra  $\ell_{2,1}$ -norm üzerinden probleme çözüm sunmuşlardır. Bu sayede aykırı değerlere (outliers) karşı seyrek çözüm (sparse solution) üzerinden daha kararlı bir yapı ortaya konmuştur.

Algoritma ayrıca çok çekirdekli işlemciler için paralel işlemeye elverişli hale getirilerek işlem yükü en aza indirilmiştir. Cui ve ark. ise bölümle ve bütünle temelinde [33] çalışmasını önermişlerdir. Veri seti öbeklere ayrıştırılarak alt-işçilerde (sub-workers)  $k$ -Ortalama kayıp fonksiyonuna göre işlendikten sonra sonuçlar genel çıktı için toparlanmıştır.

Geleneksel algoritmalarda ortaya çıkan bu gelişmelere paralel olarak YSA çevresinde de veri işleme gidişatını değiştirecek önemli iyileşmeler yaşanmıştır. DÖ olarak ortaya atılan yeni nesil YSA mimarileri özellikle görüntü işleme, sınıflandırma ve bölütleme alanlarında çığır açıcı başarılarla imza atmışlardır. 90'lı yılların sonlarında öne sürülen çalışmalar ile temelleri atılan ESA'lar [34], AlexNet'in [35] IMAGENET [36] yarışmasında sergilediği üstün başarımla araştırmacıların dikkatlerini üstüne çekmiştir. Ardından geliştirilen GoogleNet [37], VGG-16 [38] ve ResNet [39] gibi ağ önerileri ile ESA, makine görmesi alanında güçlü bir yer edinmiştir. Özellikle semantik bölütleme gibi zor bir alanda en üst başarımlar ESA'lar tarafından kazanılmıştır [40]. Özellikle, gelişen grafik kart teknolojisi ile birlikte ESA kabiliyetleri artırılmış ve dudak okuma [41], görselin içeriğini metne dönüştürme [42] ve üç boyutlu yüz görsellerinin modellenmesi [43] gibi alanlarda kayda değer sonuçlar yakalamışlardır. Özellikle BVA çerçevesinde ele alınan geniş ölçekli veriler ESA'ların başarımlarını iyileştirmelerine önemli katkılar sunmuşlardır. Diğer taraftan, sınıf başına az veri miktarının bulunduğu durumlarda eğitim için yeni tekniklerin geliştirilmesi gerekmektedir. Problem karmaşıklığı sınıf sayısının fazla olması ile artan veri setleri için, sınıf başına düşen örnek miktarının azlığı ESA'lar için aşılması gereken bir konudur. Ayrıca ESA'lar her ne kadar bir çok alanda en üst başarımları yakalasalar da gerek duydukları kaynaklar itibariyle spesifik problemler için analitik çözümlerin bulunmasını gerekli kılmaktadırlar. İstatistiksel metotlar ile modellenen problemler için yenilikçi sınıflandırma algoritmalarının bulunması bu anlamda önemlidir. Bu bağlamda tezin ana çerçevesi, özgün ESA metotlarının geliştirilmesi ve geniş ölçekli verilerin istatistiksel modellemesi olarak belirlenmiştir.

Tespit edilen problemler çerçevesinde geliştirilen model ve yöntemlere yönelik literatür özetleri ilgili bölümler içinde ayrıca verilmiştir. Yapılan bu düzenleme ile konuyla alakalı çalışmalar bir bütün olarak ele alınmış ve sunulan katkının literatürdeki yeri ve niteliği hakkında daha anlaşılır bir planlama sergilenmiştir.

## 1.2 Tezin Amacı

Veri büyüklüğü ile birlikte sınıf sayısında oluşan artışlar, problemlerin hesaplama karmaşıklığı ile birlikte algoritmaların modelleyebilme sınırlarını da zorlamaktadır.

Literatürde öne sürülen teknikler her ne kadar geniş ölçekli verilerin üstesinden gelseler de sınıf sayılarının görece az ( $\leq 1k$ ) olması önemli bir unsurdur. Diğer taraftan, gerçek hayatta sınıf başına az bir örnek ile çok fazla sınıfın modellenmesini gerektiren uygulamalar mevcuttur. Bu alanların başında biyometri gelmektedir. Kişilerin imzası, parmak izi veya yüzünün tanınmasına yönelik geliştirilen sistemlerde kişi başına onlarca eğitim örneği almak pratik değildir. Ayrıca böyle sistemlerde çok fazla ( $\geq 1k$ ) kişi üzerinden sınıflandırma yapılması gerekebilir. Bu çerçevede, tez kapsamında öncelikle sınıf başına az örnekle çok sınıflı problemlerin çözümüne yönelik bir modelin geliştirilmesi amaçlanmaktadır.

Geniş ölçekli veriler, bilgi yoğunluğu nedeniyle istatistiksel olarak modellemeye oldukça elverişli bir zemin sağlarlar. İstatistiksel öznitelikler klasik öznitelik çıkartma yöntemlerinin aksine hızlı olarak elde edilebilirler. Az miktardaki veriler için betimleyici olmayan bu öznitelikler, veri miktarının artması ile ayırt edici olmaktadır. Sınıf sayısının çok yüksek olması, istatistiksel modellerin birbirleri arasındaki girişimi arttıracak için az sayıda sınıfa sahip bir veri seti ile çalışılması gerekmektedir. Bu amaca yönelik, az sınıflı ve sınıf-içi yüksek veri miktarına sahip veri setleri için istatistiksel öznitelikler tabanında bir sınıflandırma sistematığı önerilmesi planlanmıştır.

Büyük veriye eklenen değer kavramı oldukça önemlidir. İşlenmesi gereken veri içinde hangi örneklerin değerli olduğunun tespiti ve buna bağlı olarak gereksiz örneklerin elenmesi, algoritmaları hızlandıran etmenler arasında yer alır. Bu olgu, bölütleme problemleri dahilinde ele alınan çakışık nesnelerin ayrıştırılmasında özellikle kendini göstermektedir. Bu doğrultuda, gereksiz örneklerin uyarlamalı seyreltilmesi ile çakışık nesnelerin hızlı ve etkin olarak çözülmesine yönelik bir algoritma sunulması hedeflenmiştir.

### 1.3 Hipotez

Çok sınıflı ve sınıf başına az örnek alınarak yapılacak olan sınıflandırmalar için ESA tabanlı bir çözümleme öne sürülmüştür. Geliştirilen bu model, ham verileri doğrudan işlemek yerin, verileri ağ parametrelerine uygun hale getirecek bir ön-işleme adımından geçirmektedir. Bu sadeye yapılan bir standartlaştırma, ilgili probleme etkin bir çözüm olmuştur. Ayrıca sınıflandırma için ESA modelinin çıktısı yerine gömülü öznitelik vektörleri üzerinden 1-NN tabanlı bir sınıflandırma şeması sunulmuştur. Geliştirilen yöntem  $4k$  sınıflı ve eğitim için sınıf başına 10 taneden az örnek alınarak çevrim-içi (offline) imza tanıma uygulamasında kullanılmıştır.

Az sınıflı ve sınıf başına yüksek miktarda örnek içeren veri setleri için, istatistiksel

öznitelik olarak genel histogram bilgisi ayırt edici bir yeterliliğe sahip değildir. Fakat, yerel histogram kullanarak örüntü-bağımsız rasgele dağılıma sahip veri öbekleri modellenenebilir. Ayrıca, bu öznitelikler elde edildikten sonra Kullback-Leibler diverjansı gibi metrikler üzerinden benzerlik ölçütleri hesaplandıktan sonra sınıflandırma yapılabilir. Öne sürülen istatistiksel metot, rahim ağzı kanseri olan serviks dokularının derecelendirilmesi için uyarlanmıştır. Bu sayede, doku yayılımından bağımsız hızlı ve etkin bir yöntem ile literatüre katkı sunulmuştur.

Sınıflandırma problemlerinin yanı sıra tez kapsamında yüksek çözünürlüğe sahip görüntüler üzerinde çakışık nesnelerin bölütlenmesi problemi ele alınmıştır. Çakışık nesneler, sınır pikselleri kullanılarak ayrılabilirler. Yüksek çözünürlük altında sınır pikselleri çok fazla değersiz nokta içerir. Bu durum, iteratif olarak çalışan algoritmalar için işlem yüküne neden olur. Tanımlanan bu probleme yönelik uyarlamalı piksel seyreltme ile birlikte modifiye edilmiş yeni bir  $k$ -Ortalama algoritması geliştirilmiştir. Ortaya koyulan algoritma, histopatolojik görüntülerde çakışık hücre problemi üzerinden test edilmiştir.

## 2 Teorik Arka Plan

---

Geniş-ölçek verilerin işleminde gerekli olan matematiksel arkaplan zengin bir içeriğe sahiptir. Problemlerin temellendirildiği çerçevenin (framework) YSA, istatistiksel modellemeye yönelik veya makine öğrenmesi tabanlı olmasına bağlı olarak çok çeşitli matematiksel araçlara gerek duyulmaktadır.

Tezin konuları içerisinde temel alınan algoritmalar ve matematiksel araçlara ait detaylı açıklamalar bu bölümde ele alınmıştır. Konular içindeki bütünlüğün bozulmaması için ana bölümlere ait temel kavramlar, bu başlık altında üç kısımda işlenmiştir. Her kısım ana bölümlerin sıralamasına denk gelecek şekilde gerekli olan teorik arka planı içermektedir. Bu sayede, tez kapsamında yapılan çalışmaların anlatıldığı bölümlerde detaylara girilmeden sadece geliştirilen algoritmalara dair bilgiler verilmiştir.

Bu bölüm, üç alt bölümden oluşmaktadır. *Evrişimsel Sinir Ağları* başlığında ESA'ların gelişimi ve önerilen yenilikçi katmanlar ve mimariler ele alınmıştır. *İstatistiksel Analiz Metotları*'da ise tezin içinde kullanılan istatistiksel araçların temelleri irdelenmiştir. *Eğitici-siz Kümeleme Algoritmaları* kısmında yine tez kapsamında kullanılan *k-Ortalama* ve *Yoğunluğa Dayalı Konumsal Uygulamaların Gürültü ile Birlikte Kümelenebilirliği* (Density-Based Spatial Clustering of Applications with Noise (DBSCAN)) algoritmalarına yer verilmiştir.

### 2.1 Evrişimsel Sinir Ağları

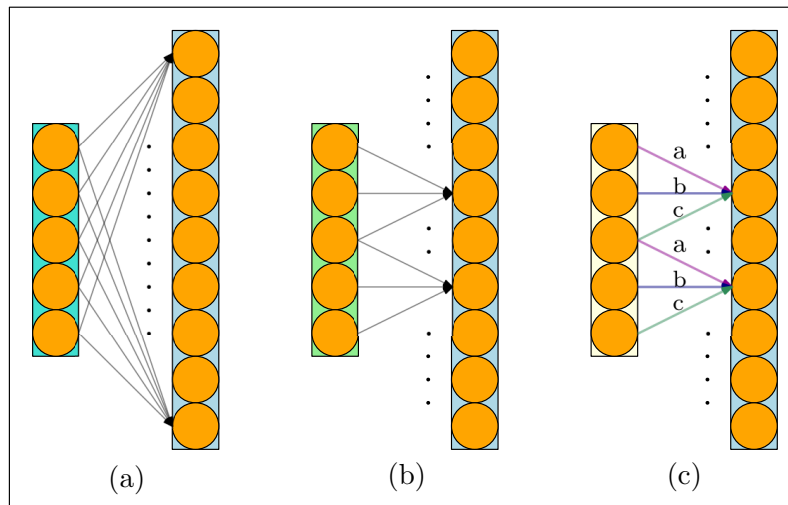
Veri işleme alanında dönüm noktası olan YSA fikri, ilk nöron yapısı olan perseptrondan bu yana birçok farklı varyasyona dönüşerek gelişimini sürdürmüştür. İlk olarak sınıflandırma problemi özelinde ortaya çıkan bu yöntem sonraları regresyon ve karmaşık fonksiyonların modellemesine kadar birçok alanda kullanılmıştır. Yapısal olarak tek katmanlı nöronlardan çok katmanlı mimarilere geçilmiş, ardından radyal temelli fonksiyon, Hopfield gibi ağlar türetilmiştir. Son on yılda ESA mimarilerinin geliştirilmesi, grafik işlemcilerindeki performans artışı ve büyük veri olgusu ile YSA, *Derin Öğrenme* gibi bir paradigma kazanmıştır.

### 2.1.1 Evrimsel Sinir Ağlarının Gelişimi

ESA, YSA'nın lokal bilgi işlemeye dayalı olarak geliştirilmiş farklı bir türüdür. ESA'nın biyolojik olarak esin kaynağı olan Hubel ve Wiesel'in çalışmasında [44], kedilerin görsel korteksinin belirli alt bölgeleri işleyen karmaşık hücre yığınlarından oluştuğu ortaya koyulmuştur. Bu bölgeler, diğer hücrelerden bağımsız olarak algıladıkları alana dair bilginin işlenmesinden sorumludurlar. Alt bölgeleri oluşturan hücreler, kendilerine girdi olarak gelen görüntüleri yerel filtreler ile analiz ederek konumsal ilintiler (korelasyon) üzerinden öznitelik çıkartırlar. Bu fikir doğrultusunda, YSA'nın bağlantı yapısı ve tanımları değiştirilerek ESA mimarisi öne sürülmüştür. Bu bağlamda (i) *seyrek ağırlıklandırma* ve (ii) *ağırlık paylaşımı*, ESA'lar için temel fikri oluşturmaktadır.

(i) **Seyrek Ağırlıklandırma:** YSA'da iki katman arasındaki matematiksel bağıntı  $\mathbf{o}^{l+1} = f(\mathbf{W} \times \mathbf{o}^l + \mathbf{b})$  olarak ifade edilir.  $\mathbf{W}, \mathbf{b}, f(\cdot)$  sırasıyla ağırlık matrisi, bias vektörü ve aktivasyon fonksiyonudur. Şekil 2.1(a)'da görüleceği üzere bu bağıntı ile iki katman arasındaki tüm nöronlar birbirlerine ağırlık çarpanları ile bağlıdır. Bu tip bağlantılara tüm-bağlı (fully-connected) denilmektedir. Yerel bilgi işlemeye dayalı geliştirilen fikir ile  $\mathbf{o}^{l+1}$ 'de bulunan bir nöronun kendisine yakın (algı alanı içinde)  $N_k$  tane *komşu* nöron üzerinden bağlantı kurması tanımlanmıştır (Şekil 2.1(b)'de  $N_k = 3$ ). Bu sayede  $\mathbf{o}^{l+1}$  katmanındaki her bir nöron  $\mathbf{o}^l$  katmanında bulunan  $N_k$  komşu nöron için özelleşmiştir. Bu bağlantı yapısının başka bir avantajı ise işlem yükünü azaltmasıdır.

(ii) **Ağırlık Paylaşımı:** ESA'nın ortaya çıkmasını sağlayan diğer bir nokta ise seyrek ağırlıklandırma ile belirlenen komşu nöronlara ait ağırlıkların  $\mathbf{o}^{l+1}$  katmanındaki



**Şekil 2.1** YSA'dan ESA'ya geçiş: (a) YSA'ya ait tüm-bağlantı yapısı (b) Seyrek bağlantı tanımı (c) Ağırlık paylaşımı sonrası katmanlar arası bağlantı

tüm nöronlar için aynı olmasıdır. Şekil 2.1(c)'de aynı renk ve harfe ait ağırlıklar paylaşımlıdır. Paylaşımlı ağırlıklar sayesinde  $\mathbf{o}^l$  ve  $\mathbf{o}^{l+1}$  arasındaki bağıntı sadece  $N_k$  tane nöron üzerinden tanımlanır. Geri yayılım esnasında gradyan akışı, paylaşımlı ağırlığın bağlantı sağladığı nöronlara gelen geri yayılım değerlerinin toplamı kadar olmaktadır. Bu adımlar sonrasına iki katman arasındaki bağıntı *evrişimsel* olarak tanımlanabilmektedir. (2.1)'de seyrek bağlantı ve ağırlık paylaşımı sonrasında oluşan ağırlık matrisi  $\mathbf{W}$ 'nin yapısı verilmiştir.  $\mathbf{o}^l$  katmanındaki nöronların  $\mathbf{W}$  ile çarpılması, aynı katmanın  $\mathbf{w} = [a, b, c]^T$  ile evrişimine ( $\otimes$ ) denk düşmektedir.

$$\mathbf{W} = \begin{bmatrix} a & b & c & 0 & 0 & 0 & \dots \\ 0 & a & b & c & 0 & 0 & \dots \\ 0 & 0 & a & b & c & 0 & \dots \\ & & \vdots & & \ddots & & \end{bmatrix} \quad \mathbf{W} \times \mathbf{o}^l \triangleq \mathbf{w} \otimes \mathbf{o}^l \quad \mathbf{w} = [a, b, c]^T \quad (2.1)$$

Bir boyut üzerinden verilen bu örnek iki boyutlu uygulamalar içinde gayet uygundur. İki boyutlu nöron dağılımlarının bağıntısı,

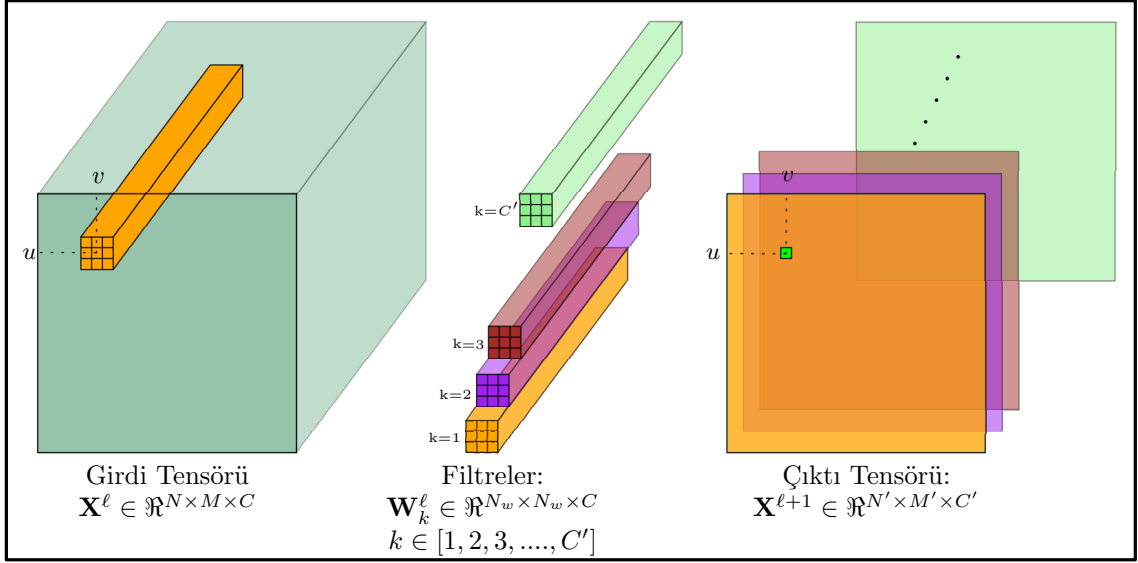
$$\mathbf{O}^{l+1} = f(\mathbf{O}^l \otimes \mathbf{W}^l + \mathbf{B}^l) \quad (2.2)$$

(2.2) ile ifade edilmektedir. Sırasıyla  $\mathbf{O}^l, \mathbf{O}^{l+1}, \mathbf{W}^l, \mathbf{B}^l$ , giriş ve çıkış nöronları, ağırlık çekirdek matrisi ve bias matrisini temsil eder. Evrişim işlemi ile görüntülerde yerel örüntülerin işlenmesine olanak sağlanmıştır. YSA uygulamalarında ham görüntünün vektör haline getirilerek ağırlık girişine uygulanması sık kullanılan bir yöntemdir. Piksellerin komşuluk bilgisinin kaybolduğu bu uygulama yerine bölgesel bilginin farklı filtreler ile işlediği ESA, ham görüntülerin işlenmesi için daha etkin bir araçtır.

Evrişim operatörü sadece iki boyutlu görüntüler için değil aynı zamanda renkli görüntülerin ele alındığı üç boyutlu tensörler için de tanımlanmaktadır.

$$\mathbf{X}^{l+1}(u, v, k) = \sum_{c=1}^C \sum_{i=-N_w/2}^{\lfloor N_w/2 \rfloor} \sum_{j=-N_w/2}^{\lfloor N_w/2 \rfloor} \mathbf{X}^l(u-i, v-j, k) \times \mathbf{W}_k^l(i, j, c) \quad (2.3)$$

(2.3)'te verilen eşitlikte  $\mathbf{X}^l \in \Re^{N \times M \times C}$  ve  $\mathbf{X}^{l+1} \in \Re^{N' \times M' \times C'}$  girdi ve çıktı tensörleri olsun.  $\mathbf{W}_k^l \in \Re^{N_w \times N_w \times C}$  boyutlarındaki  $k \in \{1, 2, 3, \dots, C'\}$  tane filtre kullanılarak üç boyutlu veriler evrişimsel olarak işlenebilmektedir. Burada dikkat edilmesi gereken nokta ise kayma işleminin uzamsal boyutlarda olurken kanal boyutunda tüm-bağlı bir yapı mevcuttur. Şekil 2.2'de girdi sensörünün herhangi bir filtre ile evrişiminden iki



**Şekil 2.2** Evrişim operatörünün üç boyutlu veriler üzerindeki işlemi

boyutlu bir *öznitelik haritası* (feature map) çıktığı görülmektedir. Bu nedenle, filtre sayısı çıktı tensörünün kanal sayısını belirler. Evrişim işlemi evrişim katmanında yer almaktadır. Bu katmanın çıktısı, YSA'da olduğu gibi bir aktivasyon fonksiyonundan geçer. Ardından yine bir evrişim katmanı ile bu işlem devam eder. Ardışık evrişim katmanlarında filtre sayısının artarak ilerlediği ESA tasarımlarında, çıktı tensörünün kanal sayısı da artmaktadır. Uzamsal boyutların korunarak devam etmesi durumunda çok büyük boyutlarda bir tensör ile karşılaşılır. Bundan dolayı her katman sonrasında hesaplama yükü artarak devam eder. Ayrıca aktivasyon fonksiyonu YSA'da yaygın olarak kullanılan *sigmoid* veya *tanh* ise aktivasyon fonksiyonlarının içerdiği üssel ifadelerden dolayı hesaplama maliyeti daha da büyür. Sürekli artan kanal ve fazla sayıda nöronun tetiklediği bir diğer problem ise *doymadır* (saturation). Bu tarz problemlerin çözümü için evrişim katmanının yanı sıra ESA'lar için özgün katmanlar geliştirilmiştir.

### 2.1.2 Özgün Katmanlar

Evrişim katmanı ile birlikte, hem YSA'ların üstesinden gelmeye çalıştıkları problemlerin hem de ESA'lar ile birlikte gelen sorunların çözümü için yeni katmanların türetilmesine ihtiyaç duyulmuştur. Bu çerçevede, yenilikçi aktivasyon fonksiyonları tanımlanmış, ayrıca havuzlama ve yığın normalizasyonu gibi sinir ağlarının eğitim kalitesini arttıracak yeni katmanlar öne sürülmüştür.

**(i) Aktivasyon Fonksiyonları:** YSA'ların birçok problemin üstesinden gelebilmesini sağlayan en önemli unsurlardan birisi kullanılan aktivasyon fonksiyonlarının *doğrusal olmayan* bir yapıda olmasıdır [45]. Bu sayede YSA'lar, birçok modeli genelleştirebilme

kabiliyetine sahiptirler. Ayrıca aktivasyon fonksiyonları ağ yapılarına, verinin daha etkin ve gürbüz (robust) işlenmesi için uyarlamalı bir mekanizma sağlar. YSA tasarımı bu fonksiyonların en çok tercih edilenleri arasında *sigmoid* ve *tanh* gelmektedir.

$$f_{sig}(x) = \frac{1}{1 + e^{-x}} , \quad f_{tanh}(x) = \frac{e^x - e^{-x}}{e^x + e^{-x}} \quad (2.4)$$

Gradyanların sönümlenmesi (vanishing gradients) problemi ağ tasarımlarında karşılaşılan ve başarımı olumsuz yönde etkileyen bir durumdur [46, 47]. Ağırlıkların güncellenmesi sırasında aktivasyon çıktılarının türevin sıfır olduğu bölgelerde birikmesi olarak ortaya çıkar. (2.4)'teki fonksiyonlar  $|x| \rightarrow \infty$  durumunda fonksiyonların türevleri ortadan kalkar. Bundan dolayı öğrenme işlemi devam edemez. Probleme çözüm olarak ESA mimarilerinde (2.5)'te ifade edilen *Düzeltilmiş Doğrusal Birim* (Rectified Linear Unit (ReLU)) ile *Sızdıran ReLU* kullanılması önerilmiştir [48].

$$ReLU(x) = \begin{cases} x & x \geq 0 \\ 0 & \text{diğer} \end{cases} , \quad \text{Sızdıran ReLU}(x) = \begin{cases} x & x \geq 0 \\ \alpha x & \text{diğer} \end{cases} \quad (2.5)$$

(2.4)'te tanımlanan aktivasyonların  $x \rightarrow \infty$  türev değerinin küçülmesinin aksine ReLU fonksiyonunun pozitif bölgesinde türev her zaman sabittir. Bu sebeple gradyanların sönümlenmesi problemi oluşmaz. Ayrıca türev ifadesi küçülmediğinden dolayı (2.4)'teki fonksiyonlara nazaran daha hızlı öğrenme imkanı sunar. ReLU hesaplama açısından da birçok fayda sağlamaktadır. Ağın derinliği arttıkça nöron sayıları da artmakta ve buna bağlı olarak işlem yükü çoğalmaktadır. Ağın ileri ve geri yayılımlarında ReLU, algoritmaların içinde basit bir *eğer* ifadesi ile hesaplanabilmesi nedeniyle işlem karmaşıklığını azaltmaktadır. Ayrıca ağ büyüklüğüne bağlı olarak (2.4)'te belirtilen fonksiyonlar hızlı bir doyuma giderken ReLU daha karardır [48]. ReLU'nun bir başka katkısı ise katmanlar arasında veriyi öğrenilebilen *seyrek* (sparse) tanım alanlarına (domain) haritalandırarak modellemesidir. Diğer taraftan ReLU'nun negatif bölgesi gradyan sönümlenmesi problemine açıktır. Bunu iyileştirmek için *Sızdıran ReLU* önerilmiştir [49].

ESA içinde ReLU kullanımı, derin ağların öğrenme yetisini arttırırken, gradyanların sönümlenmesi ve doyum gibi problemlere karşı daha kararlı ağ tasarımlarını mümkün kılmaktadır. Her ne kadar işlem yükünü azaltsa da sabit uzamsal boyutlar nedeni ile ağ içinde akan bilgi tensörünün kanal sayısındaki artış, hesaplama maliyetini hâlâ makul

seviyelerin dışında tutmaktadır. Bu probleme yönelik *havuzlama* (pooling) katmanı [50] veya adım sayısının (stride) arttırılması önerilmiştir [39].

**(ii) Havuzlama:** Girdi görüntüsü üzerinden filtreler yardımıyla düşük seviyeli özniteliklerin (low-level features) çıkartılmasıyla öznitelik haritaları oluşmaktadır. Bu haritalar tensörlerin kanal boyutunda (üçüncü boyut) yer alırlar. Kanal sayısı artarken havuzlama işlemi, yerel filtre çıkışlarını kümeleştirerek ve maskenin adım sayısı  $s \geq 2$  seçilerek uzamsal boyutta daralmaya sebep olur. Bu sayede tensör boyutları kanal boyunca uzarken uzamsal olarak daralmaktadır. Yaygın olarak *En Yüksek Havuzlama* ve *Ortalama Havuzlama* ESA'lar için kullanılmaktadır. En yüksek havuzlama (2.6)

$$Y_{max}(u, v, k) = \max_{[-N_p/2] \leq i, j \leq [N_p/2]} \{X(u-i, v-j, k)\} \quad (2.6)$$

olarak tanımlanırken Ortalama havuzlama (2.7) ile ifade edilmektedir.

$$Y_{avg}(u, v, k) = \frac{1}{N_p^2} \times \sum_{[-N_p/2] \leq i, j \leq [N_p/2]} X(u-i, v-j, k) \quad (2.7)$$

En yüksek havuzlama filtre çıkışından elde edilen değerler üzerinden  $N_p \times N_p$  komşuluğundaki piksellerin içinden en yüksek değerli olanı seçmektedir. Bu sayede konumsal değişmezlik sağlanmaktadır. Havuzlama işleminin bir diğer katkısı ise *ezberleme* (overfitting) problemine karşı çözüm sunar. Ortalama Havuzlama katmanı ise yerel olarak ele aldığı nöronların ortalamasını bir sonraki katmana aktararak çalışmaktadır.

Bazı çalışmalarda havuzlama katmanı yerine alt örnekleme için evrişim operatörlerinin adım sayısı  $s_c \geq 2$  seçilmesi ile uzamsal daralma sağlanmaktadır [39, 51]. Havuzlama katmanlarının bilgiyi ciddi ölçüde kırdığı göz önüne alınarak böyle bir yol tercih edilmiştir.

**(iii) Yığın Normalizasyonu:** ESA'lar ardı ardına dizilmiş katmanlardan meydana gelir. Her katman bir önceki katmanın çıkışını girdi olarak kabul eder. Bilgi ağ içinde katmanlarda işlenerek ilerler ve güncelleme algoritması son katman olarak tanımlanan kayıp fonksiyonunu minimize edecek şekilde parametrelere yeni değerlerini atar. Katman parametrelerinin güncellemesi katmana gelen o anki girişe göre olmaktadır. Eğitim esnasında ilgili katmanın girişi her eğitim adımında değişmektedir. Bu ise ağın eğitimini olumsuz etkileyen *iç değişkenlerin kayması* (internal covariate shift) probleminin oluşmasına sebep olur. [52]'daki çalışmada probleme çözüm olarak, ara katmanlardaki aktivasyonları verimli bir eğitim rejiminde tutacak *yığın normalizasy-*

onu (batch normalization (BN)) isminde yeni bir katman önerilmiştir.

BN katmanını yaygın uygulamalarda evrişim ve aktivasyon katmanları arasına eklenmektedir. Aktivasyon öncesi yığın istatistiklerinin düzenlenmesini sağlayarak aktivasyonun doyuma ulaşmasını engeller. Bu sayede ara katmanlar, iç değişkenlerin kayması problemine karşı daha gürbüz olurlar ve ağların öğrenimi kararlı olur. Algoritma 1’de gösterildiği gibi BN katmanını ilk olarak yığını normalize eder. Ardından, öncesinde bulunduğu aktivasyonu eğitim esnasında doyuma varmasını engelleyecek şekilde yığın için uygun ortalama ( $\mu_{\mathcal{B}}$ ) ve değişim ( $\sigma_{\mathcal{B}}$ ) değerlerini öğrenir. Burada  $\mathcal{B}$  mini-yığını simgelemektedir. Böylece ağ içinde işlenen yığın, ara katmanlarda bulunan her aktivasyon için uygun bir noktaya haritalandırılır. Bu durum Şekil 2.3’de sunulmaktadır. Şekil 2.3(a)’da ilgili katmanın aktivasyonu öncesinde yığın tamamıyla negatif bölgededir. Bu durumda ReLU tüm örnekler için 0 değerini üretir. BN katmanını ise Şekil 2.3(b)’deki gibi ilk olarak yığını *beyazlaştırır* (whitening). Ardından, aktivasyonun başarıma en iyi katkı sağlayacağı yere yığının haritalandırılması için istatistikler yeniden düzenler (Şekil 2.3(c)).

BN, katmanının iç değişken istatistiklerinin ayarlamasının yanı sıra ağ için gerekli olan diğer üst-parametrelerin ayarlanmasındaki (hyperparameters) hassasiyetleri en aza indirmektedir. Ağlarda derinliğin artması ile birlikte çok sayıda sayının çarpılıp toplanmasından dolayı sayılar büyüyerek devam eder. Bu durumda taşmanın oluşmaması için küçük öğrenme oranları seçilmekte ve bu durum eğitimi yavaşlatmaktadır. BN katmanını yığın değerlerini normalize ettiği için öğrenme oranı yüksek seçilebilmektedir. Bu sayede ağın eğitimi çok daha hızlanır. Diğer bir katkısı ise ilk değer atamalarına karşı duyarlılığı minimize eder. Ayrıca derin ağları doyuma ve ezberleme problemlerine karşı korur [52, 53].

Tanımlanan bu katmanlar tez kapsamında sıklıkla kullanılan ve literatürde en yaygın

---

#### Algoritma 1 Yığın Normalizasyonu

---

**Girdi 1:**  $\mathcal{B} = \{x_1, x_2, \dots, x_m\}$

▷ Mini yığın kümesi

**Girdi 2:**  $\gamma, \beta$

▷ Öğrenilecek olan parametreler

**Çıktı**  $y_i = BN_{\gamma, \beta}[x_i]$

1: **Prosedür** BN

2:  $\mu_{\mathcal{B}} \leftarrow \mathbb{E}[x_i]$

3:  $\sigma_{\mathcal{B}}^2 \leftarrow \mathbb{E}[(x_i - \mu_{\mathcal{B}})^2]$

▷ sıfır ortalama

4:  $\hat{x}^i \leftarrow \frac{x_i - \mu_{\mathcal{B}}}{\sqrt{\sigma_{\mathcal{B}}^2 + \epsilon}}$

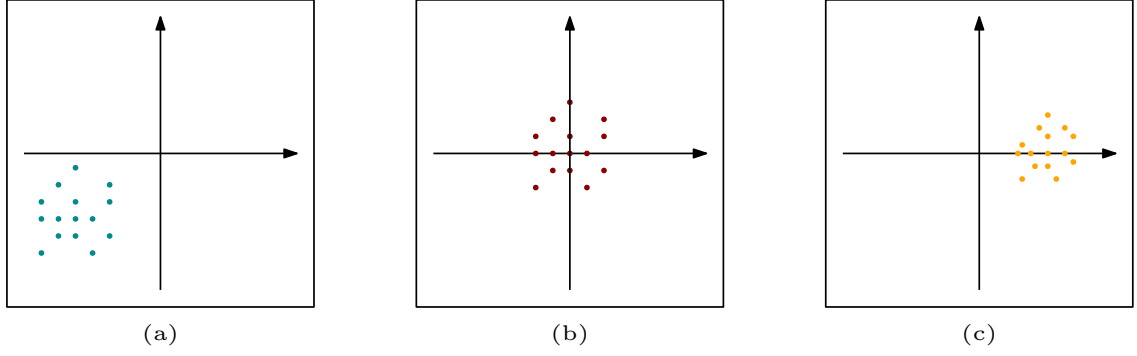
▷ normalizasyon

5:  $y_i \leftarrow \gamma \hat{x}^i + \beta$

▷ ölçekleme ve merkezleme

6: **çıktı**  $y_i$

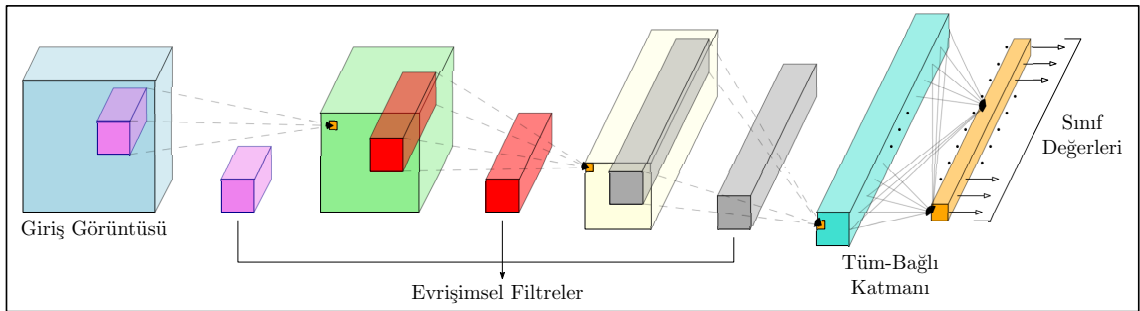
---



**Şekil 2.3** Yığın normalizasyonunun işlem sırasının iki boyutlu gösterimi:  
 (a) Aktivasyon öncesi yığın dağılımı (b) Yığın normalizasyonu (c) Yığın istatistiklerinin aktivasyona göre ayarlanması

tercih edilen katmanlardır. Bunların haricinde, ağı ezberlemeye karşı korumak için düzenleyen (regularization) *bırakma* (dropout) [35], katmanlar arasında değerlerin taşmasını engellemek için *yerel cevap normalizasyonu* (local response normalization) [35], daha etkin bir aktivasyon seçimi için gruplandırılmış kanallar arasında en yüksek değeri çıkış olarak veren *en yüksek çıktı* (maxout) [54], yüksek yığın değerlerinin kullanılmadığı durumlarda BN katmanının istatistiksel çıkarımları zorlandığı için geliştirilen *grup normalizasyonu* (group normalization) [55] gibi katmanlar da mevcuttur. Matematiksel katman tanımlarının sonrasında görüntü sınıflandırma problemi özelinde genel bir evrişimsel sinir ağı tasarımı Şekil 2.4’de verilmiştir. Ağ, giriş görüntüsünü uzamsal olarak daraltırken bilgiyi kanal boyunca toplamaktadır. Tüm-bağlantı katmanı, evrişimsel katmanlar üzerinden türetilmiş olan öznitelik vektörünü MLP gibi sınıflandırmaya çalışmaktadır. Ham veri üzerinden hem öznitelik vektörünün çıkartılması hem de onu sınıflandıracak uygun bir sınıflandırıcının aynı yapıda bulunması ESA’ların en önemli özelliklerinden birisidir.

Yapay görme (machine vision) alanında çığır açıcı gelişmeleri beraberinde getiren ESA’lar için birçok farklı bağlantı modeli önerilmiştir. Bu modeller 1000’i aşkın sınıfa ait 14 milyonun üzerinde resmi içeren geniş-ölçek veri seti ile IMAGENET [36] yarışmalarında test edilmişlerdir.



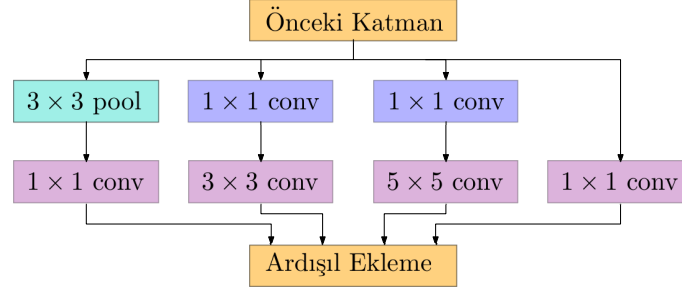
**Şekil 2.4** Görüntü sınıflandırma problemine yönelik tasarlanan örnek bir ESA yapısı

### 2.1.3 ESA için Örnek Ağ Yapıları

Yapay görme alanında nesne tespiti, bölütleme, hareketli nesnelerin takibi gibi problemlerin arasında çok sınıflı görüntülerin sınıflandırılması en temel problemler arasında yer almaktadır. Destek Vektör Makinesi (Support Vector Machine (SVM)), ortaya çıkışından ([56–58]) yaklaşık yirmi yıl boyunca analitik öznitelik çıkartma (handcrafted feature extraction) yöntemleri ile birlikte makine öğrenmesi ve yapay görme alanlarında hakimiyet kurmuştur. Özellikle SVM'in çekirdek fonksiyonları ile güçlendirilmiş versiyonlarının doğrusal olmayan dağılımdaki verilerin sınıflandırılmasında sergilediği üstün başarımlar ve teoremin arkasındaki köklü matematiksel temel, onu birçok çeşitli alanda uygulanabilir kılmıştır. Diğer taraftan, geniş-ölçekli verilerin sınıflandırılması ise SVM ve analitik öznitelik çıkartma algoritmalarının üstesinden gelmeye çalıştıkları önemli bir problem olarak kalmıştır. Özellikle uygulanacak olan problemde hangi özneliğin hangi sınıflandırıcı ile birlikte en iyi sonucu vereceği önceden bilinmemektedir. Bu ise farklı kombinasyonların problem üzerinde denenmesini gerektiren maliyetli bir iştir. ESA'ların bu iki olguyu bütüncül olarak bünyesinde barındırması ona büyük bir avantaj sağlamaktadır. Bu avantaj, 2012 yılında gerçekleştirilen IMAGENET yarışmasında AlexNet ([35]) ile dikkatlerin ESA'lar üzerine çekilmesine neden oldu. IMAGENET yarışmasında test görüntüsüne ait sınıf tahmininde doğru sınıfın ilk 5 tahmin içinde bulunuyorsa doğru, değilse yanlış tahmin olarak tanımlandığı *top-5* hata metriği bulunmaktadır. Alexnet, bir önceki yarışmayı %26'lık hata ile kazanan ESA'nın dışında geliştirilen modelin başarımlarını %15.3'e taşımıştır. Bu yüksek başarı sonrasında ESA ve DL üzerine birçok araştırmalar yapılmış ve 2012'den bu yana IMAGENET yarışmasının kazananları ESA tabanlı algoritmalar olmuştur [37–39, 51].

ESA'ların gelişiminde AlexNet dikkatleri üzerine çektikten sonra GoogleNet, VGG-16 ve ResNet gibi ağlar YSA alanında köşe taşı oluşturan kavramların ve katmanların öncüsü olmuşlardır. Bu nedenle ESA için örnek ağ yapıları çerçevesinde temel olarak bu üç ağ mimarisi incelenmiştir.

**(i) GoogleNet (Inception v1):** AlexNet'in tasarımında düşük-seviyeli özniteliklerin çıkartıldığı ilk evrimsel katmanında  $11 \times 11$  boyutlarında filtreler kullanılmıştır. Fakat girdi verisi için hangi boyutta filtrenin veya filtre yerine başka bir katmanın kullanımının daha iyi olduğu bilinmemektedir. Bu fikirden yola çıkarak [37] çalışması ile GoogleNet (diğer adı ile Inception v1) ortaya sürülmüştür. GoogleNet'in ana fikri, Şekil 2.5'de görüleceği üzere girdi tensörüne farklı boyutlarda filtrelerin uygulanması ve bu filtre çıktılarının üçüncü boyutta ardışıl eklenmesi (concatenate) üzerine kuruludur. Herbir evrim katmanının aktivasyonu ReLU'dur. Ardışıl ekleme işlemi ile gelen derinliğin hızlı artışını engellemek için tensör ilk olarak  $1 \times 1$  filtreler ile

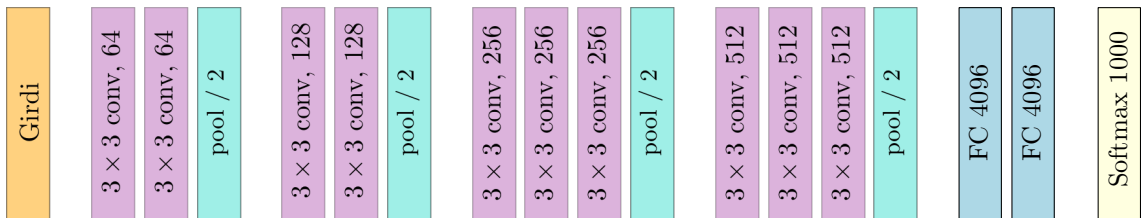


**Şekil 2.5** GoogleNet ile birlikte ileri sürülen *inception* modülüne dair bir örnek

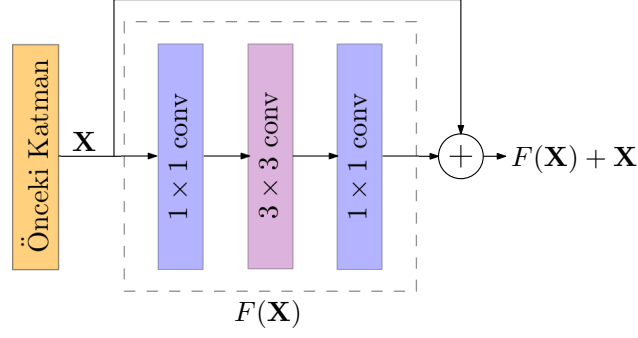
daha az derinliğe sahip alt-uzaya haritalandırılır. Bu alt-uzayda farklı boyutlardaki filtreler uygulandıktan sonra çıktı tensörleri, ana modülün çıktısını oluşturmak için ardışıl ekleme katmanına yönlendirilir. GoogleNet, BN katmanını kullanarak 2014 IMAGENET yarışmasında %6.67'lik *top-5* hata değerine ulaşmıştır.

(ii) **VGG-16:** *Visual Geometry Group* (VGG) tarafından 2014 IMAGENET yarışması için öne sürülen model ardışıl katmanların birleştirilmesi ile oluşturulmuştur [38]. AlexNet'in başarımını iyileştirmek için öne sürülen ağda sabit filtre boyutu olarak  $3 \times 3$  kullanılmıştır. Şekil 2.6'de *girdi* katmanından sonra düşük seviye özneteliklerin çıkartılması için her ne kadar  $3 \times 3$ 'lük filtreler küçük gelse de, arka arkaya  $3 \times 3$  boyutlarında filtre kullanımı  $5 \times 5$ 'lik bir kullanım etkisi oluşturmaktadır. Tasarımında 16 katman bulunduğu için *VGG-16* ismini alan model, 2014 IMAGENET yarışmasından %6.8'lik *top-5* hata oranına ulaşmıştır. VGG-16 oldukça düşük bir hata oranı elde etmesine rağmen GoogleNet'i hem başarımla hem de hesaplama maliyeti açısından geçememiştir. GoogleNet 7 milyon parametre içerirken VGG-16 138 milyon gibi çok büyük bir parametre kaynağına ihtiyaç duyar.

(iii) **ResNet:** ESA mimarilerinde ağ derinliğinin artması karmaşık problemlerin çözümünde önemli rol oynamaktadır. Özellikle geniş-ölçek verilerin sınıflandırılmasında ağın derinliği, aynı sınıfa ait çok çeşitli verilerin öznetelik uzayında birbirine yakın olarak haritalandırmasına doğrusal olmayan aktivasyonların yardımıyla imkân sağlar. Diğer taraftan ağın derinliğinin artması kaçınılmaz olarak eğitimin bozunumuna (degredation) sebep olmaktadır. BN katmanı doyum problemine karşı etkili bir çözüm olsa bile yüksek katman sayısına sahip ağların tasarımı için yeni bir



**Şekil 2.6** VGG-16 ağının blok şeması



**Şekil 2.7** ResNet-tipi ağlarda kullanılan atlamalı bağlantı mekanizmasının blok şeması.  $F(X)$  farklı tasarımlarda olabilmektedir

bağlantı yapısına ihtiyaç duyulmuştur. He ve ark. [39]'de ileri sürdükleri *rezidüal* bağlantı yapısı fazla sayıda katman içeren ESA'ların tasarımına imkan sunmuştur. Şekil 2.7, rezidüal bir bağlantıya örnek olarak verilmiştir.  $F(\cdot)$  olarak tanımlanan blok,  $X$  girdi tensörünü işlemektedir.  $1 \times 1$  evrişim katmanını ile  $X$  tensörü ilk olarak altı uzaya haritalandırılır. Sonrasında  $3 \times 3$  evrişim bloğu ile bilgi işlendikten sonra yine bir  $1 \times 1$  bloğu ile tensör  $X$  tensörü ile aynı derinliğe sahip olacak şekilde boyutlandırılır.  $F(\cdot)$  bloğunun çıktısı ile  $X$  tensörünün toplanması ile rezidüal çıktı elde edilir. Herbir evrişim bloğu BN ve ReLU katmanları da içermektedir. Rezidüal bloğun dikkate değer olan katkısı *atlamalı bağlantı* (skip connection) mekanizması ile ağın içinde kısa devre (shortcut) yollarının kurulmasıdır. Bu sayede ağın eğitiminde gereksiz olan katmanın ağırlıkları  $W, B \rightarrow 0$  olarak güncellenir. Böylece blok üzerinden bilgi tensörü  $F(X) + X \rightarrow X$  olarak akarak devam eder. Rezidüal bağlantı sayesinde geliştirilen *ResNet-152* 152 katmanlı yapısı ile 2015 IMAGENET yarışmasında %3.57'lik *top-5* hata değerini yakalamıştır.

AlexNet'in önerilmesinden sonra artan çalışmalar ile birlikte ESA'lar çok çeşitli problem türlerine uygulanarak alanlarında en üst başarıları yakalamışlardır. SVM'in çekirdek fonksiyonlarını kullanarak doğrusal olarak ayrıştırılamayan öznelikleri üst-uzayda doğrusal olarak ayrıştırması gibi farklı bağlantı yapıları çerçevesinde geliştirilen evrişimsel katmanlar, ham veri uzayından gelen girdileri doğrusal ayrıştırılabilir bir öznelik uzayına haritalandıran çekirdek fonksiyonları gibi davranmaktadır. Evrişim katmanlarının sonunda bulunan tüm-bağlı katman(lar)ı MLP gibi bir sınıflandırma yapmaktadır.

ESA'lar için geliştirilen bağlantı yapılarının yanı sıra, geniş-ölçek verilerin işlenmesinde yenilikçi güncelleme algoritmaları da başarıyı etkileyen önemli katkılar arasında yer almaktadır. Özellikle büyük verilerin işlenmesi için stokastik tabanlı algoritmalar önemli rol almışlardır.

#### 2.1.4 Güncelleme Algoritmaları

Geri yayılım algoritmasının temelini oluşturan gradyan inişi (gradient descent) kayıp fonksiyonunda oluşan hatayı ağ içinde öğrenebilir (learnable) parametreler üzerine türev yolu ile dağıtmaktadır. Böylece ilgili parametrenin kayıp fonksiyonunu minimize edecek olan türev yönelimi ve değerleri hesaplanır. MLP gibi ağlarda gradyan inişi için tüm veri setindeki örnekler üzerinden kayıp fonksiyonu ve türevleri hesaplanmakta ve bunların ortalaması güncelleme terimi olarak parametrelere ulaşmaktadır. Geniş-ölçekli veri setlerinde bu durum oldukça maliyetli bir iş yükü gerektirir.

**(i) Stokastik Gradyan İnişi:** Diğer yandan güncelleme için tüm veri seti yerine, mini yığınlar üzerinden türev yöneliminin kestirimi daha hızlı hesaplanabilmektedir. Tüm örneklerin kullanılmadığı bu yöntem *Stokastik Gradyan İnişi* (Stochastic Gradient Descent (SGD)) denilmektedir. SGD’de hatayı minimize edecek olan türev yönelimleri mini yığınlar üzerinden istatistiksel olarak kestirilerek ağırlıklar güncellenir.  $l$ . katmana ait  $t$ . andaki ağırlıklar  $\mathbf{w}_t^l$  olarak tarif edilirse, son katmandaki  $L(\cdot)$  kayıp fonksiyonunun ilgili ağırlıklar üzerine olan gradyan katkısı,

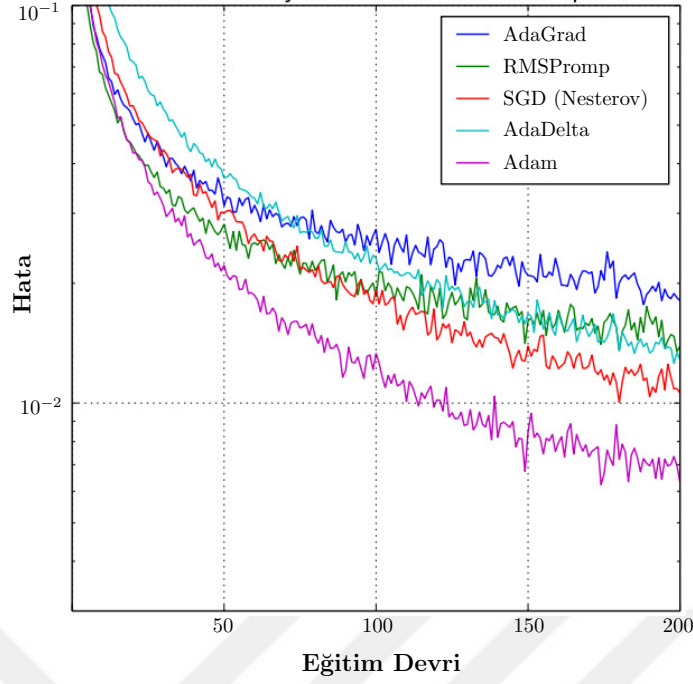
$$L_T(\mathbf{w}_t^l) = \frac{1}{N_B} \times \sum_{i=1}^{N_B} L_i(\mathbf{w}_t^l), \quad G_T(\mathbf{w}_t^l) = \frac{1}{N_B} \times \sum_{i=1}^{N_B} \nabla L_i(\mathbf{w}_t^l) \quad (2.8)$$

(2.8)’de verilmiştir.  $N_B$  mini yığındaki örnek sayısını ifade etmektedir. Eğitim örnekleri üzerinden hesaplanan gradyan inişlerinin ortalaması olan  $G_T(\mathbf{w}_t^l)$ , (2.9) kullanılarak  $\mathbf{w}^l$  ağırlıklarının güncellemesine yardımcı olur.

$$\mathbf{w}_t^l \leftarrow \mathbf{w}_t^l - \eta \times G_T(\mathbf{w}_t^l) \quad (2.9)$$

(2.9) ifadesinde geçen  $\eta$  öğrenme oranıdır.  $N_B$  kadar örnek üzerinden yapılan bir güncellemeye iterasyon denilirken tüm veri setinin taranması durumundaki adıma *eğitim devri* (epoch) denilmektedir. SGD algoritması hızlı olmasına rağmen en iyi nokta etrafında salınım değişintisi yüksektir. Bu nedenle Adam algoritması geliştirilmiştir.

**(ii) Adam Algoritması:** Adam algoritması AdaGrad [59] ve RMSProp [60] algoritmalarının güçlü yönlerinden esinlenerek [61] çalışmada sunulmuştur. RMSProp algoritmasının ikinci dereceden momentleri kullanması ve AdaGrad algoritmasının uyarlamalı öğrenme oranı kestirimi üzerine kurulmuş olan Adam, gradyan değerlerinin birinci ve ikinci momentlerini kullanarak uyarlamalı momentum tahmininde bulunmaktadır. Adam ve diğer algoritmaların karşılaştırılmaları



**Şekil 2.8** Adam algoritmasının diğer güncelleme algoritmaları ile karşılaştırmaları: Adam'ın uyarlamalı moment özelliği daha etkili bir öğrenme sağlamaktadır [61]

Şekil 2.8'de verilmiştir.

$L(\mathbf{w}^l)$ , 1. katmandaki ağırlıkların üzerine türevlenebilen bir kayıp fonksiyonu olsun. Adam algoritmasının amacı  $\mathbb{E}[\mathbf{w}^l]$ 'yi minimize etmektir.  $G_T(\mathbf{w}_t^l)$ , mini yığın içindeki her örnek için  $t$ . iterasyonda  $\mathbf{w}_t^l$  ağırlıkları üzerine olan toplam gradyanı göstermektedir. Algoritma gradyanın üssel hareketli ortalamaları ( $\mathbf{v}_t$ ) ile karesel gradyanları ( $\mathbf{s}_t$ )  $\beta_1$  ve  $\beta_2$  parametreleri ile kontrol altında tutmaktadır. Birinci ve ikinci momentlere göre hareketli ortalamaların güncellemesi (2.10)'da yapılmaktadır.

$$\begin{aligned}\mathbf{v}_t &\leftarrow \beta_1 \times \mathbf{v}_{t-1} - (1 - \beta_1) \times G_T(\mathbf{w}_t^l) \\ \mathbf{s}_t &\leftarrow \beta_2 \times \mathbf{s}_{t-1} - (1 - \beta_2) \times G_T^2(\mathbf{w}_t^l)\end{aligned}\tag{2.10}$$

$\beta_1$  ve  $\beta_2$  parametreleri [61] çalışmasında 0.9 ile 0.999 olarak verilmiştir. Bu değerlerde algoritma oldukça başarılı çalıştığı için birçok uygulamada bu şekilde kullanılmaktadır. Hareketli ortalamalara dair güncellemeler *yansız kestirim* (unbiased estimator) olarak çalışmaktadır. Daha doğru yönelimlerin tahmini için  $\mathbf{v}_t$  ve  $\mathbf{s}_t$  için (2.11)'de *yanlı düzeltme* (biased correction) yapılmaktadır.

$$\mathbf{v}_t \leftarrow \frac{\mathbf{v}_t}{1 - \beta_1^t}, \quad \mathbf{s}_t \leftarrow \frac{\mathbf{s}_t}{1 - \beta_2^t}\tag{2.11}$$

(2.11)'deki  $\beta$  ifadeleri üzerindeki  $t$  ifadesi üssel hesaplamayı simgeler. İterasyon sayısı artıkça yanlış düzeltmenin etkisi düşecektir. Algoritmanın ağırlık güncellemeleri,

$$\Delta \mathbf{w}_t \leftarrow -\eta \frac{\mathbf{v}_t}{\sqrt{s_t} + \epsilon}, \quad \mathbf{w}_{t+1} \leftarrow \mathbf{w}_t + \Delta \mathbf{w}_t \quad (2.12)$$

(2.12)'de yapılmaktadır. Ortalama yönelimini belirleyen  $\mathbf{v}_t$ 'nin moment,  $s_t$  ile kontrol edilmiştir. Böylece daha kararlı bir öğrenme sağlanır.  $\eta$  ve  $\epsilon$  öğrenme oranını ve çok küçük bir sayıyı temsil etmektedir. Bu değerin yaygın kullanımı 0.001 ve  $10^{-5}$ 'tir.

AdaGrad'ın moment yaklaşımını Rmsprop'un durağan olmayan ayarlamalara karşı iyi çalışması ile birleştiren Adam algoritması, stokastik olarak ilerleyen gradyan yönelimlerini birinci ve ikinci momentler üzerinden kontrol altına almaktadır. Bu sayede osilasyonlara sebep olacak moment değerlerinden kaçınılmış olur. Bu sebeple Adam, çeşitli problemler karşısında gürbüz olarak çalışabilmektedir.

## 2.2 İstatistiksel Analiz Metotları

İstatistiksel analiz teorisi, geniş-ölçek verilerin işlenmesinde önemli bir yere sahiptir. Öznitelik uzayı içinde örnekler arasındaki ilişkiyi modelleyebilme yeteneği etkili bir analiz çerçevesi sunmaktadır. İstatistiksel yöntemler güçlü bir matematiksel alt yapıya dayanırlar ve çok sayıda algoritmik araca sahiptirler. Bu sayede verileri etkili bir şekilde analiz edip işleyebilirler. İstatistiksel teoriler temelinde geliştirilen *Temel Bileşen Analizi* (Principal Component Analysis (PCA)) ve *Doğrusal Ayrım Analizi* (Linear Discriminant Analysis (LDA)) veri işlemede en yaygın kullanılan yöntemlerdendir. Ayrıca yine bu disiplin çerçevesinde kullanılan ve bilgi teorisinin dayandığı temeli oluşturan entropi kavramı, geniş-ölçek verileri anlamada önemli bir zemin oluşturmaktadır. Bu tez kapsamında, üçüncü bölümde geliştirilen ESA tabanlı modellerin kayıp fonksiyonu ve dördüncü bölümde öne sürülen istatistiksel özniteliklerin karşılaştırılması entropi temelinde geliştirilen yöntemler ile yapılmıştır.

### 2.2.1 Entropi

İstatistiksel mekanik çerçevesinde ortaya çıkan entropi, termodinamik bir sistemin büyük-ölçekli bir özelliğini ifade etmektedir. Sistemi karakterize eden makroskobik özellikler (basınç, sıcaklık vb.) ile mikroskobik konfigürasyonların (atom, molekül vb.) sayısı arasındaki ilişkiyi kurmaktadır. Bilgi teorisinde ise ilk defa Shannon tarafından ortaya atılan [62] entropi, bir kaynağın ürettiği bilginin miktarını ölçmede

kullanılmaktadır. Olasılıksal modeli olan bir sistem için entropi,

$$E = - \sum_{\langle i \rangle} p_i \times \log(p_i) \quad (2.13)$$

gibi tanımlanabilir.  $\log(\cdot)$  genellikle 2 tabanında kullanılmaktadır. (2.13)'de  $p_i \rightarrow 0$  ve  $p_i \rightarrow 1$  durumlarında  $E \rightarrow 0$  olmaktadır. Bu ise herhangi bir sistemde  $p_i$  olasılığına sahip olan bilginin üretilme olasılığının çok düşük veya çok yüksek olması durumunda gerçekleşmektedir. Dolayısıyla, her iki durumda da  $E \rightarrow 0$  olması, sistemin iki durum içinde benzer olan bir özelliğinin ölçütü olmalıdır. Bilginin belli durumlara öbeklenmesi veya oralarda bulunmaması bilginin *konsolide* olduğunun bir göstergesidir. Buradan yola çıkarak entropi değerinin fazla olması, durumlar arasında bilginin dağınıklığını başka bir deyişle *belirsizliğini* göstermektedir. Bu nedenle entropi için ayrıca *belirsizlik ölçütü* tanımında kullanılmaktadır.

### 2.2.2 Kayıp Fonksiyonları ve Dağılım Metrikleri

Fiziksel olguların sayısallaştırılması, onları matematiksel olarak analiz edilebilir kılmış ve birbirleri arasındaki ilişkilerin ortaya çıkartılarak doğa yasalarının keşfedilmesine yol açmıştır. Rasgele değişken kavramı ise, olaylar ile onların olasılıkları arasında bir bağıntı kurmaya yarayan matematiksel betimlemedir. Bu sayede olasılık uzayı bütününde tanımlanan olaylara ait *ortalama* ve *standart sapma* gibi bilgiler çıkartılabilmektedir. Bu değerler her ne kadar önemli bilgiler içerse de rasgele değişken için atanan değerlere bağlıdır. Örnek vermek gerekirse bir zarın yüzeylerine rasgele değişken olarak 1, 4, 5, 56, 123, 432 gibi değerler atansın. Böyle bir durumda çıkan ortalama ve standart sapma olayların olması sonucunda ortaya çıkacak bilgiyi betimlemede yetersiz kalır. Böyle bir uzaya ait daha genel bir ölçüt için, *dışarıdan atanan* rasgele değişkenlere bağlı olarak elde edilen ortalama ve standart sapma yerine, doğrudan olayların olasılıkları üzerinden tanımlanan entropi kullanılmaktadır. Buna benzer bir durum makine öğrenmesinde de ortaya çıkmaktadır.

Yapay zeka sistemlerinin öğrenme kalitesi, model için tanımlanan kayıp fonksiyonu ile doğrudan ilintilidir. Öğrenme için geliştirilen algoritmanın çıkışında elde edilen sınıflara ait değerlerin, gerçek sınıf etiketleri ile nasıl karşılaştırılıpda bir kayıp fonksiyonunun tanımlanacağı önemli bir problemidir. YSA'larda, ikili sınıflar için *dışarıdan* etiket atamaları genellikle  $C \in \{-1, 1\}$  veya  $C \in \{0, 1\}$  olarak yapılmaktadır. Bunun nedeni YSA çıkışında elde edilen sınıflara ait skor değerlerinin,  $\tanh(\cdot)$  veya  $\text{sigmoid}(\cdot)$  ile  $C$ 'nin tanımlandığı aralıklara haritalandırılabilmesidir. Bu sayede iki sınıflı ve sınıf skor değerleri  $s \in \mathfrak{R}^{2 \times 1}$  olan bir sisteme ait kayıp fonksiyonu  $\ell_2$ -norm

üzerinden,

$$L = 0.5 \times \|\mathbf{o} - \mathbf{t}\|^2 \quad (2.14)$$

olarak tanımlanabilir. (2.14) içinde  $\mathbf{o} \in \mathbb{R}^{2 \times 1}$ , sınıf skorlarının  $\tanh(\mathbf{s})$  veya  $\text{sigmoid}(\mathbf{s})$  kullanılarak haritalandırıldıkları değerleri temsil ederken  $\mathbf{t} \in \{0, 1\}^{2 \times 1} \vee \{-1, 1\}^{2 \times 1}$  ise sınıflara atanan etiketleri göstermektedir.  $\tanh(\cdot)$  ve  $\text{sigmoid}(\cdot)$  fonksiyonları  $|x| \rightarrow \infty$  durumunda türevleri 0 olmaktadır. Gradyan sönümü problemi olarak tanımlanan bu problem sebebiyle öğrenme çok yavaşlamakta veya yapılamamaktadır. Bu nedenle entropi tabanlı *çapraz-entropi kaybı* (ÇEK) (cross-entropy loss veya softmaxlog) kayıp fonksiyonu olarak ESA'lar için önerilmiştir.

Herhangi bir sınıflandırma modelinin ürettiği sınıf skor değerleri  $(-\infty, \infty)$  aralığında olabilmektedir. Bu değerlerin çok-sınıflı bir veri setinde *tek sınıf tabanlı ikili kodlama* (one-hot encoding) (OHE) ile üretilen etiketler ile karşılaştırılarak sisteme ait kayıp değerinin üretilmesi gerekmektedir. ÇEK dışarıdan atanan etiket bilgisi yerine sınıf skorlarının o anki çıktı değerleri üzerinden üretilen bağıl olasılık değerleri ile ilgilenmektedir.  $N_c$  sınıflı bir sisteme ait skor değerleri ( $\mathbf{s}$ ),

$$\mathbf{p}(i) = \frac{e^{s(i)}}{\sum_{i=1}^{N_c} e^{s(i)}} \quad (2.15)$$

ile olasılıksal bir gösterime dönüştürülmektedir. (2.15) fonksiyonu *softmax* olarak tanımlanmıştır [35]. Diğer taraftan belli bir sınıfa ( $c_o$ ) ait OHE etiketi  $\mathbf{q} = \delta(n - c_o)$  ( $n \in [1, N_c]$ ) olarak Delta-Dirac dağılımı olarak yorumlanabilir.  $\mathbf{p}$  ve  $\mathbf{q}$  olasılık kütle fonksiyonları arasındaki karşılaştırma,

$$D(q||p) = \sum_{i=1}^{N_c} \mathbf{q}(x) \times \log \left( \frac{\mathbf{q}(x)}{\mathbf{p}(x)} \right) \quad (2.16)$$

olarak tanımlanmaktadır. (2.16) ifade Kullback–Leibler Diverjansı (KLD) olarak isimlendirilmektedir.  $\mathbf{q} \in \{0, 1\}^{N_c \times 1}$  ifadesi ilgili sınıfa ait Delta-Dirac dağılımı olduğu için KLD,

$$J(\mathbf{p}) = -\log(\mathbf{p}(k)) \quad (2.17)$$

şeklinde sadeleştirilebilir. Kategorik ÇEK olarakda adlandırılan (2.17) için dikkate değer olan yer, modele ait kaybın o anki ilgili sınıfa ait skor değeri üzerinden üretilen olasılığa bağlı olmasıdır. Sınıf etiket değerinin kendisi yerine bağıl olasılığı olarak ifade

edilen bu kayıp fonksiyonu, öğrenmenin doğasına  $\ell_2$ -norma göre daha uygundur.

Bu bölümünde, Bölüm 4 çerçevesinde ele alınan yöntemlerin dayandığı kaynaklar anlatılmıştır. Bölüm 5'e ait geniş-ölçekli çakışık nesne probleminin temeli olan eğitici-siz kümeleme algoritmaları (2.3) kısmında detaylandırılmıştır.

## 2.3 Eğitici-siz Kümeleme Algoritmaları

Kümeleme algoritmaları, eğitici-siz makine öğrenmesi teknikleri kapsamında ele alınan ve verilerin öznitelik uzayında gruplandıran algoritmalarlardır. Benzer sınıflardaki örneklerin öznitelik uzayında birbirine daha yakın, farklı sınıflarda olanların ise daha uzak olması temeline dayanmaktadır. Algoritmalara örneklerin sınıf bilgisi verilmezken, dağılım içinde kaç sınıfın olduğu [27] veya noktaların belli bir komşuluktaki yoğunluğunun ne kadar olması gerektiği [63] gibi ek bilgilere ihtiyaç duyarlar.

### 2.3.1 K-Ortalama Algoritması

k-Ortalama algoritması birçok kümeleme algoritmasına [27] esin kaynağı olmuş köklü bir algoritmadır. Eğitici-siz algoritmalar bazında birçok alandan kullanılan k-Ortalama, parametre olarak sadece sınıf sayısına gerek duyar. Ardından rasgele atanan sınıf merkezleri etrafında algoritma veriyi kümelemeye çalışır. k-Ortalama ve benzeri algoritmaların her zaman yakınsak (converge) olması onların en önemli özelliklerinden biridir [31, 64]. Algoritmanın kayıp fonksiyonu,

$$J = \sum_{\mathbf{v}_i \in \mathcal{V}} \min_{k=1}^{N_c} \{\|\mathbf{v}_i - \mathbf{c}_k\|^2\} \quad (2.18)$$

olarak ifade edilmektedir. (2.18)'deki kayıp,  $\mathcal{V}$  kümesi içindeki herbir  $\mathbf{v}_i$  örneğin  $N_c$  tane sınıf merkezine ( $\mathbf{c}_k$ ) olan uzaklıkları arasından en düşük olanlarının toplamı üzerinden tanımlanmıştır. Algoritma ilk olarak  $\mathbf{c}_k$  merkezlerinin rastgele atanması ile ilklendirilmektedir. Algoritma 2'deki işleyiş ile  $\mathbf{c}_k$  merkezleri iteratif olarak uzayda belli noktalara yakınsarlar. Kayıp fonksiyonunun belli bir eşik değerinin ( $\theta$ ) altına düşmesi ile elde edilen  $\mathbf{c}_k$ 'lar,  $N_c$  tane sınıfın merkezleri olarak kabul edilir ve kümeleme işlemi bitirilir.

(2.18) iç-bükey (convex) bir problem değildir. Bu nedenle algoritma,  $\mathbf{c}_k$  değerlerinin ilk atamasına karşı oldukça hassastır. Farklı atamalar sonucunda farklı sınıflandırılmış kümelerin ortaya çıkması muhtemeldir. Ayrıca, algoritmaya küme içinde kaç tane sınıf

---

**Algoritma 2** k-Ortalama Algoritması

---

**Girdi 1:**  $\mathcal{V} = \{\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_m\}$ 

▷ Örnek Uzayı

**Girdi 2:**  $N_c$ 

▷ Sınıf Sayısı

**Çıktı:**  $\mathcal{U}_k$ 

▷ Sınıflandırılmış Kümeler

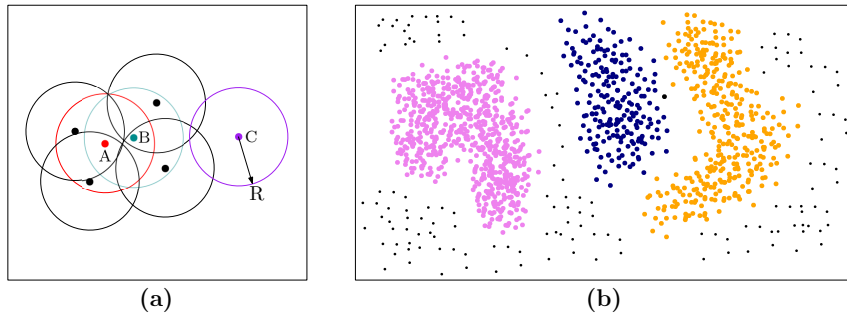
```
1: Prosedür K-ORTALAMA
2:   Rastgele  $\mathbf{c}_k$  değerlerinin atanması
3:    $\theta$ : Hata için alt sınır
4:    $J^* \leftarrow 2 \times \theta$ 
5:   while  $J^* \geq \theta$  do
6:      $\mathcal{U}_k \leftarrow \emptyset, \forall k$ 
7:     for  $\mathbf{v}_i \in \mathcal{V}$  do
8:        $\alpha \leftarrow \underset{k}{\operatorname{argmin}} \|\mathbf{v}_i - \mathbf{c}_k\|$ 
9:        $\mathcal{U}_\alpha \leftarrow \mathcal{U}_\alpha \cup \mathbf{v}_i$ 
10:     $\mathbf{c}_k \leftarrow \operatorname{ortalama}\{\mathcal{U}_k\}, \forall k$ 
11:     $J^* \leftarrow \sum_{\mathbf{v}_i \in \mathcal{V}} \min_{k=1}^{N_c} \{\|\mathbf{v}_i - \mathbf{c}_k\|^2\}$ 
çıktı:  $\mathcal{U}_k$ 
```

---

bilgisinin verilmesi gerekliliği başka bir dezavantajdır. Bu problemlerin üstesinden gelmek için uzaydaki örnekler arasındaki ilişkiyi, birbirleri arasındaki yerel mesafelere bakarak kestiren *Density-Based Spatial Clustering of Applications with Noise (DBSCAN)* algoritması geliştirilmiştir. Bölüm 2.3.2’de açıklanan DBSCAN, sınıf sayısından bağımsız olarak kümelemeyi, örneklerin tanımlanmış bir mesafe içindeki komşu sayısına göre yapmaktadır.

### 2.3.2 DBSCAN Algoritması

k-Ortalama algoritmasında kayıp fonksiyonu örnekler ve sınıf merkezleri arasındaki mesafenin  $\ell_2$ -normu üzerinden tanımlanmasından dolayı iç içe girmiş sınıfların ayrıştırılmasında yetersiz kalmaktadır. Buna yönelik, karmaşık kümeleme problemleri için yerel yoğunluklardan yola çıkarak yeni algoritmalar geliştirilmiştir [65, 66].



**Şekil 2.9** DBSCAN algoritmasının çalışmasına  $N_p = 4$  ve  $R = 2$  parametreleri için bir örnek. A ve B noktaları  $R$  yarıçapı içinde  $N_p$  komşuya sahip oldukları için küme noktası olurken C noktası gürültü olarak belirlenir.

---

**Algoritma 3** DBSCAN algoritması

---

**Girdi 1:**  $\mathcal{V} = \{v_1, v_2, \dots, v_m\}$  ▷ Örnek Uzayı  
**Girdi 2:**  $N_p$  ve  $R$  ▷ Minimum Komşu Sayısı ve Arama Yarıçapı  
**Çıktı:**  $s\{v_i\}$  ▷ Herbir  $v_i$  için sınıf ataması

1: **Prosedür** YOĞUNLUK TABANLI GRUPLAMA  
2:  $s\{v_i\} \leftarrow \emptyset, \forall i$  ▷  $v_i$  örneklerinin sınıf bilgileri "boş" olarak atanır.  
3:  $C \leftarrow 0$  ▷ Sınıf indisi  
4: **for**  $v_i \in \mathcal{V}$  **do**  
5:  $\mathcal{N} \leftarrow KomsuKume\{v_i, R\}$  ▷  $v_i$ 'nin  $R$  içindeki komşuları.  
6: **if**  $|\mathcal{N}| \leq N_p$  **then**  
7:  $s\{v_i\} \leftarrow Gurultu$   
8: **else**  
9:  $C++$  ve  $s\{v_i\} \leftarrow C$   
10:  $\mathcal{P} \leftarrow \mathcal{N} / \{v_i\}$   
11: **for**  $p_j \in \mathcal{P}$  **do**  
12: **if**  $s\{p_j\} = Gurultu$  **then**  
13:  $s\{p_j\} \leftarrow C$  ▷ Gürültü olan örnek sınır örneği olur.  
14:  $\mathcal{N}' \leftarrow KomsuKume\{p_j, R\}$   
15: **if**  $|\mathcal{N}'| \geq N_p$  **then**  
16:  $\mathcal{P} \leftarrow \mathcal{P} \cup \mathcal{N}'$   
**çıktı:**  $s\{v_i\} \in \mathcal{V}$

---

Bu algoritmaların en çok tercih edilenlerinden olan DBSCAN, ilk olarak [63] çalışmasında ortaya koyulmuştur. Algoritma, aynı küme içindeki örnekler için yerel bağlantısallığının (connectivity) farklı kümeler arasındaki örnekler için daha fazla olması özelliğinden faydalanarak gruplama yapar.

Algoritma 3, DBSCAN algoritmasının işleyişini göstermektedir. Dışarıdan  $R$  yakınlık yarıçapı ve  $N_p$  minimum komşu sayısının verilmesi gerekir. Algoritma ilk olarak uzay içindeki her örneği tarayarak  $R$  içindeki komşularını hesaplar. Mesafe metriği için  $\ell_2$ -normun yanı sıra  $\ell_1$ -norm veya Mahalonobis metriği de kullanılabilir. İlgili nokta, etrafında  $N_p$  ve daha üstü komşuya sahipse çekirdek (core), değilse gürültü olarak etiketlenmektedir. Bu durum  $R = 2$  ve  $N_p = 4$  için Şekil 2.9(a)'da gösterilmiştir. Çekirdek örneğinin tespitinden sonra kümeleme işlemi başlar. Eğer önceki döngülerde bu örneğe sınıf atanmışsa, örneğin komşuları bulunur ve aynı sınıfa dahil edilirler. Atanmamışsa yeni sınıf kümesi oluşturulur ve komşuları ile gruplandırılarak yeni kümeye dahil edilirler. Ardından kümelenmiş her bir komşu için bu sorgu yapılarak devam eder. Algoritmanın temsili bir çıktısı Şekil 2.9(b)'de verilmiştir. Örneklerin yakın olması durumunda etiket yayılımı kesintiye uğramadan devam eder ve sınıf içi yakınlığı güçlü olan dağılımlar kolaylıkla öbeklenebilmektedirler. Asıl kümeler arasındaki siyah noktalar ise verilen parametreler için gürültü olarak etiketlenmiştir.

Parametre tahmini her algoritmada olduğu gibi DBSCAN için de önemli bir problemdir.

Verinin yapısına uygun parametrelerin seçimi algoritmanın başarımını doğrudan etkilemektedir. Örnek olarak,  $R$  parametresi aynı sınıf içindeki noktalar arasındaki mesafeye nazaran küçük seçilirse etkin bir gruplandırma yapılamamaktadır. Bundan dolayı, sınıf-içi etiket yayılımı ya kesintiye uğrar ya da hiçbir etiketleme yapılamadan tüm küme gürültü olarak ele alınır. Diğer taraftan bu parametrenin yüksek seçilmesi durumunda sınıflar-arası boşluğu dolduran gürültüler (Şekil 2.9(b)) komşu değerlendirmesine girerler. Bu durumda, sınıf ataması yapılan bir kümenin etiketi başka kümeler üzerinden yayılarak farklı kümeler tek kümeymiş gibi sınıflandırılır.

## 2.4 Sonuç

Bölüm 2, tezin ana bölümleri olan Bölüm 3, 4 ve 5 için temel alınan matematiksel araçların toplamı ile oluşturulmuştur. Ana bölümlerde sadece tez kapsamında üretilen katkılar işlenmiştir. Yapılan bu düzenleme sayesinde bölüm-içi akıcılık korunurken gerekli matematiksel açıklamalar Bölüm 2’de detaylıca verilmiştir.

## Çok Sınıflı Geniş Ölçek Verilerin Az Sayıda Örneklem ile ESA Kullanılarak Modellenmesi ve Sınıflandırılması

---

Geniş-ölçek verilerin sınıflandırılması, işlem karmaşıklığı ve hesaplama maliyeti açısından analitik öznitelik çıkarma yöntemleri (handcrafted methods) ve sınıflandırma algoritmalarının üstesinden gelmeye çalıştıkları önemli bir problemdir. Diğer taraftan, veri miktarının fazla olması, DÖ tabanlı yöntemlerde ağların eğitimine katkı sunduğu için avantaj haline gelir. Literatürde ele alınan sınıflandırma problemlerinde yaygın olarak sınıf başına yüksek miktarda örneğin bulunduğu veri kümeleri kullanılmıştır. Ancak gerçek hayat problemlerinde (parmak izi, imza, yüz vb. biyometri tanıma) sınıf başına az bir eğitim örneği ile çok fazla sınıfın tanınması gibi durumlar da mevcuttur. Bu doğrultuda üçüncü bölüm kapsamında, sınıf başına az ( $\leq 10$ ) ve sınıf sayısının fazla ( $\geq 1k$ ) olduğu problemlerin çözümüne yönelik ESA tabanlı bir DÖ mimarisi tasarlanmış ve geliştirilmiştir. Ayrıca, tüm-bağlantı (fully-connected) katmanından alınan öznitelik vektörlerinin *Sınıf Merkezi Tabanlı Sınıflandırıcı* (SMTS) algoritması kullanılarak sınıflandırılmaları önerilmiştir.

Tezin bu bölümü üç alt bölümden oluşmaktadır. **3.1** bölümünde, önerilen ESA modeli ile birlikte SMTS algoritmasının detayları açıklanmıştır. **3.2**'de ESA ve SMTS'in uygulanabileceği bir alan olarak çevrim-dışı imza tanıma uygulaması seçilmiş ve önerilen sınıflandırıcının başarımlar değerlendirilmesi sunulmuştur. **3.3** bölümü ise yapılan çalışmanın sonuçlarını içermektedir.

### 3.1 ESA Tabanlı Ağ Yapısı ve SMTS Algoritması

Birçok araştırma, ESA'da ağ derinliğinin karmaşık öğrenme problemlerinin çözümünde önemli rol oynadığını göstermektedir. IMAGENET [36] yarışmasını kazanan ağlar incelendiğinde katman sayısının sürekli arttığı görülmüştür [35, 37, 38, 67]. Katman sayısının arttırılması hesaplama maliyetini de beraberinde getirir. Bu nedenle, probleme özgü yeni bir ağ tasarlamak, geniş-ölçekli sınıflandırma problemleri için önceden tasarlanmış ağları kullanmaktan başarımlar/maliyet ilişkisi açısından daha

uygundur. ESA yapıları belirli bir problem için geliştirildiği takdirde başarımları daha da artar [68]. Bu, katman sayısının, filtre boyutlarının veya katman türlerinin farklı veri setlerine ilişkin problemlerin çözümünde çeşitlilik göstermesi nedeniyle.

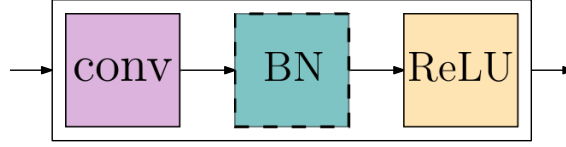
### 3.1.1 ESA Tabanlı Önerilen Model

ESA tabanlı geliştirilen modellerde veriler, ham olarak işlenebilmenin yanı sıra bir ön-işleme adımından sonra da ağı eğitiminde kullanılabilirler. Ön-işleme sayesinde veriler belli bir standart şablon içerisine haritalandırıldıkları için ağların eğitimi daha kararlı olmaktadır. Bu nedenle, çalışma kapsamında önerilen ağ modeli için, ağı uyarlanacağı problemin doğasına uygun bir ön-işleme adımının kullanılması planlanmıştır. Ön-işleme adımı sayesinde işlenecek olan görsel, ESA'da düşük-seviye (low-level) öznelilikleri çıkartan ilk katman filtre boyutları ile uyumlu hale getirilmektedir.

Bu tez kapsamında geliştirilen ESA modeli, Visual Geometry Group (VGG) tarafından öne sürülen VGG-tipi ESA modelleri baz alınarak ortaya koyulmuştur. Geliştirilen model için buradan sonra Geniş-ölçek ESA (G-ESA) ismi kullanılacaktır. G-ESA'nın ve VGG'nin literatürde sıklıkla kullanılan ağlarının parametreleri Tablo 3.1'de verilmiştir.  $N \times N$ , çekirdek boyutlarını, @, filtre numarasını,  $p$  ve  $s$ , sırasıyla sıfır-ekleme

**Tablo 3.1** Ağların parametre bilgileri | @: Kullanılan filtre sayısı  
 $p, s$ : Sıfır ekleme ve filtre kaydırma miktarları

| VGG-16                                                                                                                                     | VGG-M                                                                                                    | VGG-S                                                                                                    | G-ESA                                                                |
|--------------------------------------------------------------------------------------------------------------------------------------------|----------------------------------------------------------------------------------------------------------|----------------------------------------------------------------------------------------------------------|----------------------------------------------------------------------|
| Giriş Görüntü Boyutları: $224 \times 224 \times 1$                                                                                         |                                                                                                          |                                                                                                          |                                                                      |
| conv : $3 \times 3$ @64 p:1 s:1<br>conv : $3 \times 3$ @64 p:1 s:1                                                                         | conv : $7 \times 7$ @96 p:0 s:2                                                                          | conv : $7 \times 7$ @96 p:0 s:2                                                                          | conv : $5 \times 5$ @128 p:2 s:2                                     |
| maxpool : $2 \times 2$ p:0 s:2                                                                                                             | maxpool : $3 \times 3$ p:0 s:2                                                                           | maxpool : $3 \times 3$ p:2 s:3                                                                           | maxpool : $3 \times 3$ p:1 s:2                                       |
| conv : $3 \times 3$ @128 p:1 s:1<br>conv : $3 \times 3$ @128 p:1 s:1                                                                       | conv : $5 \times 5$ @256 p:0 s:1                                                                         | conv : $5 \times 5$ @256 p:0 s:1                                                                         | conv : $3 \times 3$ @256 p:1 s:2<br>conv : $3 \times 3$ @256 p:1 s:1 |
| maxpool : $2 \times 2$ p:0 s:2                                                                                                             | maxpool : $3 \times 3$ p:0 s:2                                                                           | maxpool : $2 \times 2$ p:1 s:2                                                                           | maxpool : $3 \times 3$ p:1 s:2                                       |
| conv : $3 \times 3$ @256 p:1 s:1<br>conv : $3 \times 3$ @256 p:1 s:1<br>conv : $3 \times 3$ @256 p:1 s:1                                   | conv : $3 \times 3$ @512 p:1 s:1<br>conv : $3 \times 3$ @512 p:1 s:1<br>conv : $3 \times 3$ @512 p:1 s:1 | conv : $3 \times 3$ @512 p:1 s:1<br>conv : $3 \times 3$ @512 p:1 s:1<br>conv : $3 \times 3$ @512 p:1 s:1 | conv : $3 \times 3$ @512 p:1 s:1<br>conv : $3 \times 3$ @512 p:1 s:1 |
| maxpool : $2 \times 2$ p:0 s:2                                                                                                             | maxpool : $3 \times 3$ p:0 s:2                                                                           | maxpool : $3 \times 3$ p:1 s:3                                                                           | maxpool : $3 \times 3$ p:1 s:2                                       |
| conv : $3 \times 3$ @512 p:1 s:1<br>conv : $3 \times 3$ @512 p:1 s:1<br>conv : $3 \times 3$ @512 p:1 s:1<br>maxpool : $2 \times 2$ p:0 s:2 | ↓                                                                                                        | ↓                                                                                                        | ↓                                                                    |
| conv : $3 \times 3$ @512 p:1 s:1<br>conv : $3 \times 3$ @512 p:1 s:1<br>conv : $3 \times 3$ @512 p:1 s:1<br>maxpool : $2 \times 2$ p:0 s:2 |                                                                                                          |                                                                                                          |                                                                      |
| conv : $7 \times 7$ @4096 p:0 s:1                                                                                                          |                                                                                                          |                                                                                                          |                                                                      |
| FC : 4096 akt. : ReLU<br>FC : 4096 akt. : ReLU                                                                                             | FC : 4096 akt. : ReLU<br>FC : 4096 akt. : ReLU                                                           | FC : 4096 akt. : ReLU<br>FC : 4096 akt. : ReLU                                                           | FC : 4096 akt. : ReLU                                                |
| etiket: $N_{cls}$ akt.: softmax                                                                                                            | etiket: $N_{cls}$ akt.: softmax                                                                          | etiket: $N_{cls}$ akt.: softmax                                                                          | etiket: $N_{cls}$ akt.: softmax                                      |



**Şekil 3.1** Tablo 3.1 içindeki her *conv* bloğunun iç yapısı: *conv* bloğunun başarımı BN bloğunun olup olmamasına göre değişkenlik gösterdiği için sınırı kesikli çizgilerle çizilmiştir

(padding) ve kayma (stride) parametrelerini temsil eder. *Kayma*, evrişim veya havuz çekirdeğinin hareketindeki adım miktarını ifade eder. *Sıfır-ekleme*, girdi tensörlerinin kenarlarına eklenen sıfır miktarıdır. *Tüm-bağlantı* (Fully-connected (FC)) katmanları, *Çok Katmanlı Perseptron* (Multi-Layer Perceptron (MLP)) benzeri bir bağlantı yapısı ile nöronlar arası tamamen bağlı katmanların tanımı için kullanılmaktadır. Tablo 3.1'deki her bir *evrişim* (conv) katmanının içeriği, Şekil 3.1'de gösterildiği üzere hem *Yığın Normalizasyonu* (Batch Normalization (BN)) hem de aktivasyon fonksiyonu olarak *Doğrultulmuş Doğrusal Birim* (Rectified Linear Unit (ReLU)) içerir.

İlk ağ olan VGG-16'da, bu çalışmada kullanılan en derin ağ, tüm mimari boyunca sabit çekirdek boyutu  $3 \times 3$  kullanılmıştır. İlk iki katman  $3 \times 3$ , @64,  $s = 1$  filtreleri ve kayma miktarından dolayı ağın ihtiyaç duyduğu parametre sayısı çok büyüktür. Bunun yanında VGG-16, ağın derin yapısı nedeniyle yüksek hesaplama maliyetine sahiptir. Diğer iki ağ, VGG-M ve VGG-S, parametre boyutlarına göre birbirine benzerdirler. İki ağ arasındaki fark, kullanılan havuz boyutları ve kayma parametreleridir.

G-ESA için ilk katmanda  $s = 2$  seçilerek 128 filtre kullanılmaktadır. Çok sayıda sınıf ve az miktarda eğitim verisinin işlenmesi nedeni ile düşük seviyedeki öznitelik çıkaran filtre sayısının fazla olması gerekliliği ön görülmüştür. Bu sayede filtre çeşitliliği artmakta ve verinin betimlenmesine daha fazla olanak sağlanmaktadır. Diğer taraftan filtre sayısına, işlem karmaşıklığı ve doyum (saturation) gibi etmenler bir üst sınır getirirler. G-ESA'ya yönelik empirik hesaplamalar doğrultusunda düşük-seviye öznitelik çıkartan filtre miktarı 128 olarak tespit edilmiştir. Artan filtre boyutundan kaynaklanan hesaplama maliyetini dengelemek için ilk katmandaki filtrelerin kaydırma parametresi  $s = 2$  olarak seçilmiştir. Bu sayede girdi görüntüsünün boyutları yarı yarıya düşmektedir. Ayrıca VGG-[M,S] ağlarının aksine, G-ESA'da birinci ve ikinci havuz katmanlarından sonra iki *conv* katmanı kullanılır. Ek olarak, VGG ağları ile G-ESA'nın arasındaki bir diğer fark, FC katmanlarının sayısıdır. VGG ağlarının iki FC katmanı varken G-ESA bir FC katmanı kullanılmaktadır.

Yapılan bu karşılaştırmalar sayesinde farklı derinlik, çekirdek büyüklüğü ve havuzlama katmanlarının etkileri araştırılmıştır. Ayrıca, FC katmanlarından elde edilen gömülü öznitelik vektörü (embedding feature vector (EFV)) kullanımının sınıflandırmaya olan

katkıları da analiz edilmiştir. Bununla birlikte EFV'lerin sınıflandırılması için özgün SMTS algoritması önerilmiştir. Kayıp fonksiyonu olarak tüm ağlar için ÇEK tabanlı *softmaxlog* tercih edilmiştir.

### 3.1.2 SMTS Algoritması

ESA'nın sonundaki FC katmanları, ReLU aktivasyon fonksiyonuna sahip MLP ağı olarak işlev görür. FC katmanları, evrimsel katmanlar tarafından oluşturulan öznitelik vektörünü doğrusal olarak sınıflandırmak için çabalarlar. Buna göre ESA, ağı başarımını artırmak için gelen görüntüleri doğrusal bir ayırt edici ile sınıflandırılabilir şekilde ağırlıkları günceller. Ayrıca bu süreç esnasında, FC katmanları doğrusal sınıflandırıcı özelliğine sahip üst-düzlem (hyper-plane) haline gelir. Bu şekilde döngüsel olarak birbirlerini iyileştiren mekanizma el iyilenmiş (optimum) denge noktasını bulur.

FC katmanından alınan EFV'ler doğrusal sınıflandırıcılar ile ayrıştırılabilir. Bu bağlamda *Destek Vektör Makineleri* (Support Vector Machine (SVM)) literatürde en yaygın kullanılan sınıflandırıcılardan biridir. Fakat sınıf ve sınıf başına düşen örnek sayılarının çarpımı  $N_c \times N_s$  olan eğitim kümesindeki toplam örnek sayısının fazla olması nedeniyle SVM,  $O(n^3)$  olan işlem karmaşıklığı nedeni ile böyle bir problem için elverişsizdir. Özellikle bire-karşı-bir (one-against-one) öğrenme stratejisi düşünüldüğünde  $N_c \geq 1k$  gibi sınıf sayısı için  $C(N_c, 2)$  tane, bire-karşı-hepsi (one-against-all) stratejisi için  $N_c$  tane üst-düzlem oluşturulmalıdır. Diğer taraftan

---

#### Algoritma 4 SMTS için eğitim ve test algoritması

---

|                                                                                               |                               |
|-----------------------------------------------------------------------------------------------|-------------------------------|
| <b>Girdi</b> $S_k$                                                                            | ▷ Sınıf $k$ İçin Örnek Kümesi |
| <b>Çıktı</b> $M \in \mathbb{R}^{d \times N_c}$                                                | ▷ Sınıf Merkez Matrisi        |
| 1: <b>Prosedür</b> EĞİTİM                                                                     |                               |
| 2: <b>for</b> $k = 1$ <b>to</b> $N_c$ <b>do</b>                                               |                               |
| 3: $\mathbf{v}_c \leftarrow \frac{1}{n} \sum_{\mathbf{v}_i \in S_k} \mathbf{v}_i$             |                               |
| 4: $\mathbf{M}_{:,k} \leftarrow \mathbf{v}_c$                                                 |                               |
| 5: <b>return</b> $M$                                                                          |                               |
| <hr/>                                                                                         |                               |
| <b>Girdi</b> $M, \mathbf{v}_t \in \mathbb{R}^{d \times 1}$                                    |                               |
| <b>Çıktı</b> $pr\_cls$                                                                        | ▷ Tahmini Sınıf               |
| 1: <b>Prosedür</b> TEST                                                                       |                               |
| 2: <b>for</b> $k = 1$ <b>to</b> $N_c$ <b>do</b>                                               |                               |
| 3: $\text{dist}(k) \leftarrow \ \mathbf{v}_t - \mathbf{M}_{:,k}\ $                            | ▷ $\ \cdot\  : \ell_2$ Norm   |
| 4: <b>sonuç:</b> $pr\_cls \leftarrow \arg \min\{\mathbf{dis} \in \mathbb{R}^{N_c \times 1}\}$ |                               |

---

yine sık kullanılan bir sınıflandırıcı olarak *k-En Yakın Komşu* (*k*-Nearest Neighbour (*k*-NN)) kullanılması halinde bir test örneği için  $N_c \times N_s$  sorgu yapılması gerekecektir.

Sınıf başına düşen eğitim örneği az miktarda olması herhangi bir sınıfın gömülü uzaydaki yayılımını sınırlandıran bir kısıttır. Buna dayanarak, sınıfın dar bir uzay bölgesinde yer tutması ve sınıfların çok sayıda olması nedeni ile *k*-NN için tüm eğitim örneklerini kullanmak yerine, sınıf merkezleri üzerinden 1-NN sınıflandırmayı temel alan Algoritma 4'deki SMTS algoritması önerilmiştir. Test örneği, sınıf merkezlerine olan  $\ell_2$  mesafesine göre sınıflandırılmıştır. Test örneği için hesaplama karmaşıklığı yalnızca EFV'nin boyutuna ve sınıf sayısı  $N_c$ 'ye bağlıdır.  $\mathbf{v}_i$ , herhangi bir *k* sınıfının eğitim için kullanılan EFV'si olsun ve  $\mathbb{S}_k = \{\mathbf{v}_i \mid \mathbf{v}_i \in \mathbb{R}^{d \times 1} \ i \in [1, n]\}$  kümesinde yer alsın. Bu bilgilere göre, eğitim ve test prosedürleri Algoritma 4 'de verilmektedir. *Eğitim* yordamında **M** matrisinin sütunları olarak atanan  $\mathbf{v}_c$  sınıf-ortalama vektörleri,  $\mathbf{v}_i$  EFV'leri üzerinden hesaplanır. *Test* prosedüründe ise test verisine ait EFV, **M** matrisindeki sütun vektörleri üzerinden 1-NN ile sınıflandırılır.

Önerilen ESA modeli ve SMTS algoritmasının denenmesi için pratikteki uygulama sahaları arasında en uygun olarak biyometrik işaretlerin sınıflandırılması problemi seçilmiştir. Bu problem, modellerin kişi başına az örnek ile çok fazla kişi üzerinden eğitilmesi gereksinimi nedeni ile önerilen ağ ve algoritmanın testi için uygun bir zemin teşkil eder. Bu çerçevede, öne sürülen güncel geniş-ölçek veri setlerinin olması nedeni ile çevrim-dışı imza tanıma problemi, model ve algoritmanın sınanması için tercih edilmiştir.

## 3.2 Uygulama Alanı: Çevrim-dışı İmza Tanıma

DÖ alanında biyometrik verilerin sınıflandırılması değerli bir konudur. Parmak izi tanıma, yüz tanıma, hareket üzerinden kişilerin tespit edilmesi gibi geniş bir uygulama alanı içinde barındırır. Biyometrik verileri sınıflandırma konusunu diğer sınıflandırma problemlerinden farklı kılan ise, sınıf başına çok fazla örnek almadan sınıflandırmanın yapılabilmesidir. Literatürde varolan çalışmalarda kişi başına 10 taneden fazla örnek alınarak modellerin eğitimi yapılmıştır. Buradan yola çıkarak, tez kapsamında önerilen yöntemin 4000 (*4k*) kişinin bulunduğu imza veri setine uygulaması yapılmıştır.

### 3.2.1 Literatür Özeti

İmza tanıma konusundaki son çalışmalar problemin çözümüne önemli katkılar sunsalar da, geniş ölçekli imza veri setleri kullanıldığında tanıma başarımları az örnek alınarak yapılan eğitimlerde hala yeterli değildir [69]. Literatürde,

imza özniteliklerinin çıkarılması ve tanınmasında çeşitli örüntü tanıma teknikleri uygulanmıştır. İmza tanıma için mevcut çoğu model, Histogramların Gradyanları (Histogram of Gradient (HoG)) [70, 71], Ölçek-Bağımsız Öznitelik Dönüşümü (Scale-Invariant Feature Transform (SIFT)) [72] ve grafometrik [73], geometrik [74] ve global [75] öznitelikleri kullanmışlardır. Bunların yanında sınıflandırma için, Saklı Markov Modelleri (Hidden Markov Models (HMM)) [76, 77], Şablon Eşleştirme (Template Matching) [78], Bulanık Mantık (Fuzzy Logic) [79] ve Destek Vektör Makineleri (Support Vector Machine (SVM)) [80] gibi makine öğrenme teknikleri ele alınmıştır. Bu sınıflandırma yöntemlerinin tümü analitik öznitelik çıkarımı yöntemlerine (handcrafted methods) ihtiyaç duyar ve başarımları sonuçları seçilen özniteliklerin imzaları sınıflar arası ayrımın yüksek olduğu bir uzayda temsil edilebilme kabiliyetine bağlıdır. Bu nedenle, öznitelik çıkarma adımı çok fazla çaba ve özen gerektirir.

Analitik özellik çıkarma ve sınıflandırma algoritmaları, işlem karmaşıklığı nedeni ile geniş-ölçek verilerin işlenmesinde yüksek hesaplama maliyetine ihtiyaç duyarlar. HoG, SIFT ve SURF gibi algoritmaların RAM gereksinimleri işlenecek görüntülerin fazla olması nedeniyle oldukça yüksektir. Benzer zorluklar sınıflandırma algoritmalarında da geçerlidir. Örnek olarak, 30k eğitim örneği içeren bir veri seti için, en yakın komşu ( $k$ -NN) algoritması her test durumu için 30k sorgu yapması gerekecektir. Ayrıca, SVM algoritması herhangi bir çekirdek fonksiyonu ile kullanıldığında  $O(n^3)$  hesaplama karmaşıklığına ve  $O(n^2)$  RAM miktarına gereksinimi vardır. Her ne kadar SVM için Primer Tahmini Alt Gradyan Çözücüsü (Primal Estimated sub-GrAdient SOLver for SVM (PEGASOS)) gibi Stokastik Gradyan İniş (Stochastic Gradient Descent (SGD)) tabanlı algoritmalar [81] 'da önerilmiş olsa da, çekirdek SVM + geniş-ölçek veri işleme hala hesaplama karmaşıklığı ve hafıza gereksinimleri bakımından önemli bir problemdir.

Baltzakisa ve ark. [74]'de, 115 kişi için bir çevrimdışı imza doğrulama sistemi geliştirmişlerdir. Global ızgara (grid) ve doku (texture) özniteliklerinin kullanıldığı çalışmada sınıflandırıcı olarak, iki evreli bir perseptron olan bir-sınıf-bir-ağ (One-Class-One-Network) seçilmiştir. Eğitim ve test için kullanılan veri oranları sırasıyla %75 ve %25'tir. Makalede doğrulama oranını olarak %93,60 bildirilmiştir. [82]'de, öznitelik çıkartmak için Ayrık Radon Dönüşümü (Discrete Radon Transform) , Temel Bileşen Analizi (Principal Component Analysis (PCA)) ve sınıflandırmada Olasılıksal Sinir Ağı (Probabilistic Neural Network) kullanılmıştır. Bu çalışmada MYCT imza veri seti tercih edilmiş ve 10 eğitim örneği kullanılarak test edilen sistemin Eşit Hata Oranı (Equal Error Rate) %9,87 olarak bulunmuştur. Hafemann ve ark. [83] ise, GPDS-960 veri setini kullandılar. Bu veri kümesi 881 imzacının verisini içermektedir. Ancak, tüm veri seti çalışmada kullanılmamıştır. Bu çalışmada, 160 imzacının verileri

GPDS 160 veri kümesi olarak ve 300 imzacının verileri GPDS 300 olarak bölünmüştür. Bu iki veri seti imza doğrulama için kullanılmıştır. Ayrıca çalışmada MCYT [84] (75 imzacı), CEDAR [85] (55 imzacı), Bengal (100 imzacı) ve Devanagari (160 imzacı) veri setleri kullanılmıştır. Eğitim ve test oranları %50'ye %50'dir. [86]'te Soleimani ve ark., GPDS Sentetik [87] veri setindeki 4000 imzacı verisinin yarısını kullanmış ve hata oranı %13,30 olarak bulmuşlardır. Jagtap ve ark. [88], GPDS Sentetik veri setinden 1.000 kişinin bir kısmını kullanarak bir çalışma öne sürmüşlerdir. Veri setinin %90'ı eğitim, %5'i validasyon ve %5'i test için ayrılan çalışmada %98'lik başarı elde ettiler. Fakat bu kadar yüksek miktarda eğitim verisi kullanmak pratik değildir. Benzer şekilde, [89]'da ise yazarlar GPDS Sentetik veri setindeki 560 imzacıya ait verinin %75'ini eğitim %25'ni test için bölmüşlerdir. Bu şartlar altında elde edilen %95'lik bir başarı gerçek hayat problemlerinin çözümünü yansıtan bir skor olduğu tartışmalıdır.

Son yıllarda imza tanıma için ESA ile yapılan çalışmalar da yaygınlaşmıştır. Khalajzadeh ve ark. [90]'de Farsça imza tanıma için yaptıkları çalışma, ESA'ların imza tanıma alanında kullanımına dair ilk yapılan çalışmalardan birini oluşturur. [91]'de ise çevrimdışı imza doğrulama için özgün bir ESA kullanılmış ve GPDS-160 için %1,72 hata oranı elde edilmiştir. Sounak ve ark. [92]'de, Siamese ağ yapısını kullanmışlardır. Literatürde ele alınan bu çalışmalar çoğunlukla test kümesinin eğitimden daha düşük veya aynı miktarda kullanılan çalışmalardır. Gerçek hayatta ise bir kişiden  $\geq 10$  imza toplamak pratik değildir. Özellikle kişi sayısının ( $\geq 1k$ ) olan uygulamalarda bu durum daha da zorlaşmaktadır. İmza tanıma alanındaki çoğu çalışma ise 1k veya daha az sayıda imzacıya sahip veri kümeleri ile ilgilidir [79, 80, 90, 91, 93–95]. Bu durumların üstesinden gelmek için modern biyometrik sistemlerle, kişi başına daha az veri toplayarak uygulanabilir metodların geliştirilmesi üzerine çalışılmaktadır [96–98]. Bu ihtiyaç doğrultusunda önerdiğimiz G-ESA ve SMTS algoritmalarını imza tanıma problemine uyarladık.

Çalışmada 4k kişiden ve kişi başı 24 (toplam 96k) gerçek (genuine) imzadan oluşan GPDS-Sentetik (bundan sonra GPDS-4k denilecektir) veri seti kullanılmıştır. Veri seti %25 eğitim ve %75 test (bundan sonra 25-75 olarak isimlendirilecek) ve %50 eğitim ve %50 test (bundan sonra 50-50 olarak isimlendirilecek) olarak bölünmüştür. Eğitim fazında kişi başına 25-75 için 6, 50-50'de ise 12 imza düşmektedir. Ayrıca, çalışma kapsamında dar-ölçek veri setlerinden MCYT ve CEDAR için de incelemeler yapılmıştır. Her veri seti için G-ESA, VGG tarafından önerilen VGG-[S,M,16] ağları ile karşılaştırılmıştır. VGG-16 ağı IMAGENET yarışmasında 2014 kazananıdır. Bunlarla birlikte, BN ve farklı aktivasyon fonsiyonlarının tanıma başarımına olan katkıları da irdelenmiştir.

### 3.2.2 Veri Seti ve Yöntemin Uyarlanması

Son yıllarda araştırmacılar, otomatik imza tanıma alanında geliştirilen algoritmaların karşılaştırılması için genel kullanıma açık imza veri setleri oluşturmuşlardır. Oluşturulan veri setleri imzaların toplanma şekline göre *çevrim-içi* veya *çevrim-dışı* olarak kategorilere ayrılırlar. Eğer bir veri setinde imzalar tablet, bilgisayar veya telefon aracılığı ile ekrana kalem ile çizilerek toplanmışsa *çevrim-içi*; bir kağıdın üstüne farklı özelliklerde (kalem, renk, kalınlık) çizilip resimleri çekilerek elde edilmiş ise *çevrim-dışı* olarak isimlendirilir. Çevrim-içi formattaki imzalar çevrim-dışı olanların aksine, zaman içinde konum, basınç, ivme gibi ek bilgiler de içerirler [99]. Çevrimdışı formattaki imzalar ise kağıt üstüne çizim işlemi bittikten sonra 300 dpi veya 600 dpi'de taranır ve ayrı imzalar olarak kaydedilir. Bu çalışmada, çevrim-dışı olarak kaydedilen CEDAR, MCYT ve GPDS-4k veri setlerini kullanılmıştır.

CEDAR imza veri seti [85] imza tanıma ve doğrulamada en sık kullanılan veri setlerinden birisidir. 55 kişiye ait imzaları içerir. Gerçek imzalar, 20 dakika arayla her imzalayandan 24 imza toplanarak oluşturulmuştur. CEDAR veri setindeki imza görüntüleri gri tonlama modunda ve *png* biçiminde bulunur. Ancak, tüm imza boyutları sabit bir şablonda ayarlanmamıştır.

MCYT veri kümesi [84] ise, 75 kişiden kişi başı 15 gerçek ve 15 sahte imza toplanarak oluşturulmuştur. MCYT'de imzalayanlar (gerçek veya sahte farketmez) aynı zeminde ve aynı kalemle imzalamışlardır. Tüm imza görüntüleri 600 dpi çözünürlükte bir tarayıcı ile dijitalleştirilmiştir. Bu görüntüler gri tonlama modunda ve *bmp* formatında mevcuttur. Diğer veri kümelerinin aksine imza görüntüleri, 850×360 çözünürlüğünde sabit boyutta kaydedilmiştir.

Çalışmada kullanılan diğer bir imza veri seti ise GPDS-4k'dır. Veri seti, [99] algoritması

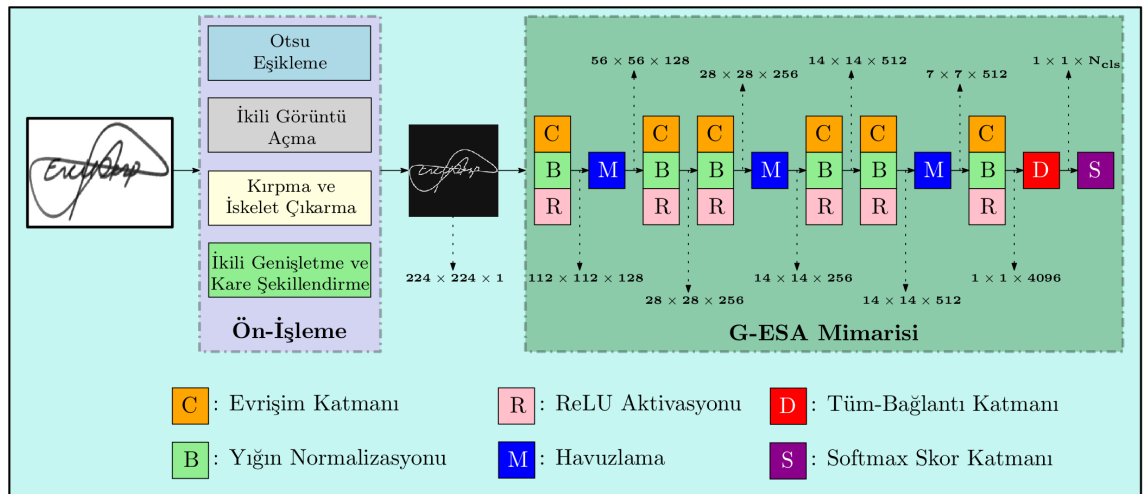


**Şekil 3.2** Kullanılan veri setlerinden alınan bazı örnekleri: İlk satır GPDS-4k'ya, ikinci satır CEDAR'a ve üçüncü satır ise MCYT'ye karşılık gelir

ile oluşturulmuş 4k sentetik kişinin Şekil 3.2’de gösterildiği gibi farklı kalem ve tonların simülasyon çıktısı olarak kişi başı 24 imzanın toplanması ile elde edilmiştir. İmzalar *jpg* biçiminde ve 600 dpi eşdeğer çözünürlüğe sahiptir. Tüm imza boyutları sabit bir şablonda ayarlanmamıştır. Bu durum, sınıf-içi değişkenliği arttırdığı için algoritmaların karşısına önemli bir problem olarak çıkar. Otomatik imza tanıma sistemlerinin başa çıkması gereken diğer bir zorluk ise eğitim sürecinde kullanılan örneklerin sayısıdır. Veri setlerine ait imza örnekleri Şekil 3.2 verilmiştir. GPDS-4k ve MCYT veri setlerinde imza örnekleri farklı piksel boyutlarına sahipken CEDAR’da ise örnekler belirli bir boyuta sabitlenmiştir. Diğer taraftan GPDS-4k veri setindeki imzalardaki kalınlık farklılıkları ESA tasarımı için göz önüne alınması gereken bir durumdur. Bu problemlerin aşılmasına yönelik, ham olarak alınan imza görüntüsü ESA’nın girişine uygun sabit bir standart görüntüye dönüştürülmesi için *ön-işleme* adımı uygulanmıştır.

İlk aşamada, gri seviye imza görüntüsü Otsu metodu [100] ile ikili görüntüye dönüştürülmüş ve ikili görüntüye morfolojik açma işlemi uygulanmıştır. Ardından imza haricindeki boşlukların giderilmesi için ikili görüntü üst, alt ve sağ, sol kısımlardan kırılmıştır. İmzaların farklı kalemler ile oluşturulması göz önünde bulundurulduğunda kalınlıklarının değişkenliği söz konusudur. Bu değişkenliği ortadan kaldırmak için öncelikle ikili imzanın iskelet yapısı (morphological skeleton) çıkartılmış ve sonrasında iskelet yapısı  $r = 3$  boyutunda yapısal disk elemanı ile genişletilmiştir (dilation). Böylece kalem genişliği ne olursa olsun ön işlemenin ardından görüntü sabit bir forma sahip olur. Ön-işleme adımı ile birlikte G-ESA’nın veri setlerine uyarlanması ile oluşturulan sınıflandırıcı modeli Şekil 3.3’de verilmiştir.

Diğer bir değişkenlik ise imzaların en boy oranlarıdır. Önerilen ESA giriş olarak



**Şekil 3.3** Ön-işleme adımı ve G-ESA’nın çevrim-dışı imza tanıma için uyarlanması ile elde edilen sınıflandırıcı modeli

$224 \times 244 \times 1$  boyutlarında görüntü kabul ettiği için ön işleme adımından elde edilen görüntülere yeniden boyutlandırma (resizing) işlemi uygulanması gerekir. En-boy oranının korunması için uzun olan kenar belirlendikten sonra kısa olan kenarın uzunluğu, uzun kenar ile aynı olacak şekilde ikili görüntüye simetrik olarak sıfır eklenir. Bu sayede ikili görüntü karesel bir yapıya dönüştürülür ve sonrasında yapılan boyutlandırma işleminde ölçek korunmuş olur. Örnek olarak  $I_{bw} \in \mathbb{R}^{50 \times 100}$  boyutlarında kırılmış bir ikili imza görüntü olsun.  $\mathbf{0}^{25 \times 100}$  ise  $[25, 100]$  boyutlarında sıfır matrisi olarak tanımlansın.  $I_{bw} \leftarrow [\mathbf{0}; I_{bw}; \mathbf{0}]$  sütun ekseninde birleştirme ile yeni  $I_{bw}$ 'nin şekli kare hale gelir.

Elde edilen bu modelin başarımlar değerlendirilmesi iki aşamadan oluşmaktadır. İlk aşamada MCYT ve CEDAR veri setleri üzerinden az sayıda sınıf ile ağlar test edilmiştir. Ardından GPDS-4k veri seti kullanılarak kapsamlı bir değerlendirme ile ağların başarımlar ölçütleri farklı senaryolar için çıkartılmıştır.

### 3.2.3 Başarım Değerlendirmesi

*Başarım Değerlendirmesi* bölümü iki alt bölümü içermektedir. *Test İş Akışı* adlı ilk bölüm, ayrıntılı eğitim ve test süreçlerinin nasıl organize edildiği ile ilgilidir. İkinci bölüm *Sonuçların Değerlendirilmesi*'de ise ilk aşamada MCYT ve CEDAR veri setlerinden elde edilen sonuçlar üzerinden başarımlar değerlendirilmesi yapılmış; ikinci aşamada ise GPDS-4k veri seti üzerinden ağların kıyaslanmasına geçilmiştir. Bütün ağlar, MATLAB için geliştirilen MatConvNet [101] aracı ile eğitilmiş ve test edilmiştir.

**(i) Test İş Akışı:** Ağların eğitimi esnasında ön-işleme adımı her mini-yığın (mini-batch) için tekrar tekrar yapılmasının gerekliliği, özellikle GPDS-4k veri seti kullanıldığında fazladan bir iş yükü getirmektedir. Bu nedenle ağların eğitimi ve testinden önce tüm resimler Şekil 3.3'de gösterilen ön-işleme adımları uygulanarak kaydedilmiştir. Bu sayede mini-yığınların tekrarlı ön-işleme adımına girmesinin önüne geçilmiştir. Ardından, beş eğitim ve test alt kümesi rastgele 25-75 ve 50-50 oranları için ayrı ayrı oluşturulmuş ve görüntülerin indeks değerleri kaydedilmiştir. Ağların eşit şartlarda eğitilip test edilebilmesi için tüm ağlarda aynı indis değerleri kullanılmıştır. Ağırlıkların ilklendirilmesi için her ağda rasgele atama kullanılmış, VGG ağlarının IMAGENET yarışmasında öğrendiği ağırlıklar adil bir karşılaştırma için tercih edilmemiştir. Eğitim fazında model, her iterasyon sonrasında test verisinin %50'si ile test edilmiştir. Bu veri sadece ağ modelinin öğrenme sürecini görselleştirmek amacıyla kullanılmıştır. 50 eğitim devri sonunda eğitilen modelin başarımları, tüm test verisi üzerinden hesaplanmıştır. Ağların eğitiminde kayıp fonksiyonu olarak çapraz-entropi tercih edilmiş ve veriler sıfır-merkezli hale getirilerek kullanılmıştır.

Kapsamlı bir testin yapıldığı GPDS-4k veri setinde ana test iş akışı dört aşamadan oluşmaktadır. İlk aşamada, ağlar BN kullanılmadan [38, 102]’de ifade edilen parametreler ile test edilmiştir. İkinci aşamada, BN’nin tanıma başarımı üzerindeki etkisinin gözlemlenmesi için ağlarda BN katmanı kullanılmıştır. Üçüncü aşamada, FC katmanlarından elde edilen EFV’ler STMS algoritması kullanılarak sınıflandırılmıştır. Ayrıca sızdıran ReLU’nun (leaky ReLU) G-ESA üzerindeki etkileri son aşamada incelendi. Tablo 3.5 içindeki başarımların değerleri, rasgele oluşturulmuş beş eğitim ve test alt kümesiyle ağların test edilerek elde edilen ortalama değerlerdir. Başarımların kriterleri olarak *Doğruluk* (Accuracy), *Ortalama Hassasiyet* (Average Precision) ve *makro F-ölçüm* metrikleri kullanılmıştır. Bu değerler  $\mathbf{K} \in \mathbb{R}^{N_c \times N_c}$  ( $N_c$ : sınıf sayısı) karışıklık matrisinden elde edilmiştir. Doğruluk ölçütü (4.4) ile  $\mathbf{K}$  üzerinden hesaplanabilir.

$$ACC = \frac{\sum_i \mathbf{K}_{ii}}{\sum_{i,j} \mathbf{K}_{ij}}. \quad (3.1)$$

*Doğruluk* sınıflandırma problemlerinde tek başına yeterli bir kriter değildir. Bu sebeple her sınıf için *hassasiyet* (precision) ve *duyarlılık* (recall) değerleri (4.5) kullanılarak hesaplanır.

$$PRE_i = \frac{\mathbf{K}_{ii}}{\sum_i \mathbf{K}_{ij}}, \quad RCL_i = \frac{\mathbf{K}_{ii}}{\sum_j \mathbf{K}_{ij}} \quad (3.2)$$

(4.4) ve (3.3) kullanılarak hesaplanan değerler üzerinden toplam başarımların ölçümü olarak, *Ortalama Değer* ve *makro F-ölçümü* (3.3) ile bulunur.

$$Ort. PRE = \frac{1}{N_c} \sum_{i=1}^{N_c} PRE_i \quad (3.3)$$

$$F_i = \frac{2 \times PRE_i \times RCL_i}{PRE_i + RCL_i}, \quad Makro F1 = \frac{1}{N_c} \sum_{i=1}^{N_c} F_i \quad (3.4)$$

Ağların eğitimi için kullanılan eğitim parametrelerinin listesi Tablo 3.2’de verilmiştir. *Eğitim oranı* (learning rate),  $[10^{-2}, 10^{-3}]$  aralığında iterasyona göre üssel olarak azalmaktadır. Bir mini-yığına kaç resmin işlendiği ise *Mini-Yığın* parametresi ile belirtilmiştir. Bu değer, GTX 1080 Ti ekran kartında VGG-16 için çalıştırılabilecek seviyede seçilmiştir. *Ağırlık ilklendirme* için [103] çalışmasında önerilen atama algoritması kullanılmıştır. En iyileme için ise Nesterov ile İvmelendirilmiş Gradyan (Nesterov Accelerated Gradient (NAG)) tercih edilmiştir. *Eğitim Devri* ve *momentum* değerleri sırasıyla 50 ve 0.9 olarak seçilmiştir. Eğitim ve test işlemlerinin yürütülmesi

**Tablo 3.2** Ağların eğitimi için kullanılan üst-parametre listesi

| Üst-Parametreler               |                                    |
|--------------------------------|------------------------------------|
| Öğrenme Oranı                  | $[10^{-2} \ 10^{-3}]$ üssel azalım |
| Mini Yığın Miktarı             | 30                                 |
| Ağırlık İklendirmesi           | Xiavier                            |
| Güncelleme Alg.                | NAG                                |
| Momentum                       | 0,9                                |
| Ağırlık $\ell_2$ Düzenleyicisi | $10^{-3}$ veya $10^{-4}$           |
| Eğitim Devri                   | 50                                 |

için donanım olarak Intel Xeon E5-2630 12 çekirdek 2.60 GHz CPU, 48 GB RAM ve 11 GB GTX 1080 Ti GPU kullanılmıştır.

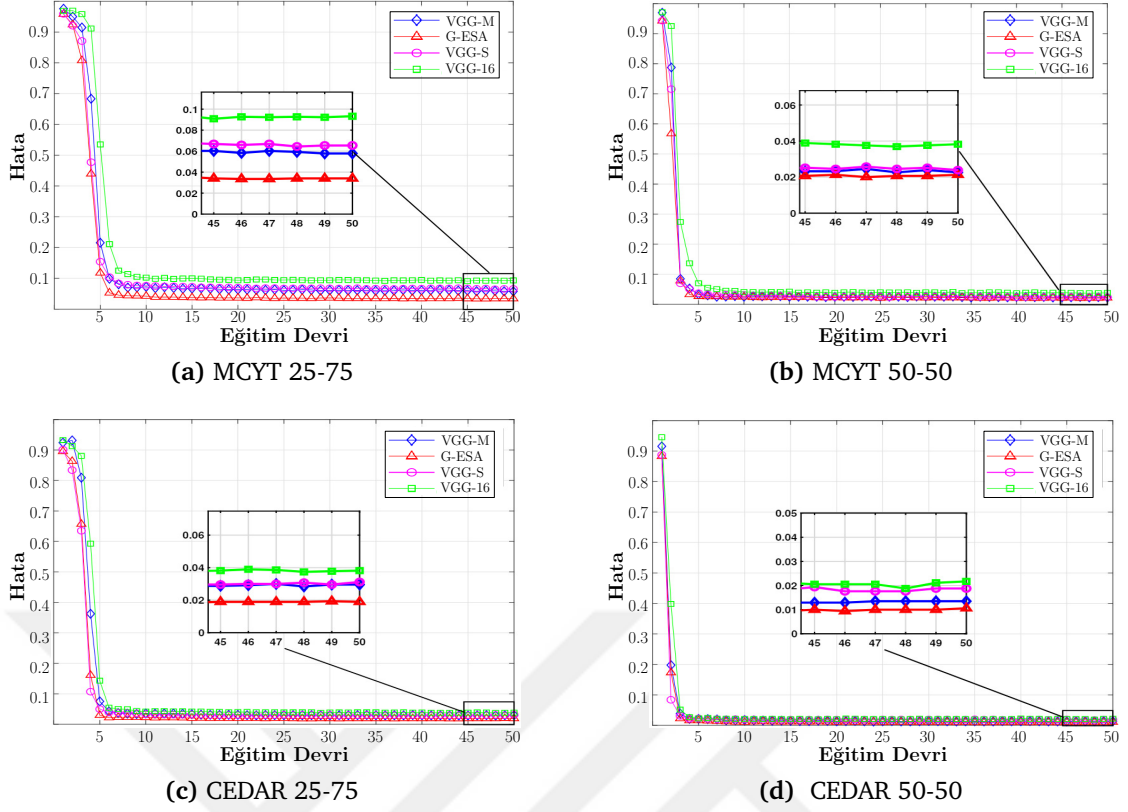
**(ii) Başarım Analizi ve Ağların Karşılaştırılması:** Test aşaması farklı veri setleri üzerinde kurgulanmıştır. MCYT ve CEDAR veri setleri, daha az sayıda kişi barındırmasına rağmen gerçek kişilerden oluşturulduğu ve literatürde çoğunlukla kullanıldığı için gerçekçi bir senaryo olması nedeni ile ilk başarım değerlendirilmesinde tercih edilmiştir. İkinci etapta ağlar, GPDS-4k veri seti kullanılarak az sayıda örnek ile 4k kişinin tanınmasına yönelik geniş-ölçek problemi doğrultusunda test edilmişlerdir.

MCYT ve CEDAR veri setlerine ait sonuçlar Tablo 3.3’de verilmiştir. Bu karşılaştırmada BN kullanımı ve aktivasyon fonksiyonu ReLU’nun sızıntı parametresi 0 olarak seçilmiştir. Eğitim fazında 25-75 oranı MCYT için kişi başına 4 ve CEDAR’da ise 6 eğitim örneğine karşılık gelmektedir. Sonuçlardan görüleceği üzere az örnek ile yapılan testlerde MCYT veri setinde %96,41 ve CEDAR’da %98,30 doğruluk başarımı ile G-ESA önemli bir üstünlük seğilemiştir. CEDAR veri seti 55 kişi, MCYT ise 75 kişiden oluştuğu ve 25-75 oranında CEDAR’da kişi başına 6 eğitim örneği düştüğü için CEDAR veri setinde MCYT’ye göre başarım daha iyi çıkmıştır. 50-50 oranında ise eğitim örneklerinin artması, ağların tamamının öğrenmesine önemli bir katkı sağlar. Böylece tüm ağlar etkili öğrenme ile yakın sonuçlar ortaya koymaktadırlar.

**Tablo 3.3** MCYT ve CEDAR veri setleri üzerinden rasgele oluşturulan beş farklı alt kümenin kullanılması ile elde edilen sonuçların ortalamaları

| Model  | MCYT          |               |               |               | CEDAR         |               |               |               |
|--------|---------------|---------------|---------------|---------------|---------------|---------------|---------------|---------------|
|        | 25-75         |               | 50-50         |               | 25-75         |               | 50-50         |               |
|        | Doğruluk      | Makro F1      | Doğruluk      | Makro F1      | Doğruluk      | Makro F1      | Doğruluk      | Makro F1      |
| VGG-S  | 93,47%        | 93,34%        | 97,53%        | 97,48%        | 97,13%        | 97,07%        | 98,15%        | 98,12%        |
| VGG-M  | 93,96%        | 93,87%        | 97,73%        | 97,71%        | 97,35%        | 97,30%        | 98,75%        | 98,74%        |
| VGG-16 | 90,95%        | 90,85%        | 95,93%        | 95,84%        | 96,59%        | 96,55%        | 98,15%        | 98,12%        |
| G-ESA  | <b>96,41%</b> | <b>96,36%</b> | <b>98,13%</b> | <b>98,11%</b> | <b>98,30%</b> | <b>98,29%</b> | <b>98,88%</b> | <b>98,87%</b> |

Şekil 3.4, her devirde eğitilen ağları test ederek elde edilen hata grafiklerini



**Şekil 3.4** MCYT ve CEDAR veri setlerinin eğitim fazındaki öğrenme eğrileri

göstermektedir. Bu noktada, eğitim aşamasının gelişim süreci görsel olarak sunulmuştur. İlgili grafikte ağlarda kullanılan BN katmanının etkisi, eğitimin hızlı bir şekilde gerçekleşmesi ile görülmektedir. 10. devirde neredeyse eğitim tamamlanmıştır. Toplamda 50 devir boyunca eğitimin sürdürülmesinin nedeni, GPDS-4k için 50 devir kullanılmasıdır. Şekil 3.4(a) ve Şekil 3.4(c) sırasıyla MCYT ve CEDAR veri setleri için 25-75 oranı için elde edilen tanıma hatasını gösteren grafikleridir. Son 5 devir incelendiğinde, VGG-S ve VGG-M ağlarının yakın başarımlar değerlerine sahip olduğu görülmektedir. VGG-16 ağı en fazla parametre içermesine rağmen en yüksek hata oranını vermektedir. G-ESA ise en az hataya sahiptir. Şekil 3.4(b) ve 3.4(d), MCYT ve CEDAR veri setleri için 50-50 oranında ağların eğitim safhalarını irdeleyen grafiklerdir. Ağların hata seviyeleri nispeten yakın seviyelerde olsalar bile önerilen G-ESA yine en düşük hata oranına sahiptir.

Bir veri kümesi oluşturmak için imza örnekleri toplandığında, aynı kişilerden iki veya daha fazla seansta imzalar toplanabilmektedir. Bu örnekler aynı ortamda alınmasına rağmen, numunelerin zamansal bileşenleri arasında farklılıklar olabilir. Bu tanıma performansını etkileyen önemli bir faktördür. MCYT ve CEDAR veri kümeleri için daha gerçekçi bir test çerçevesi oluşturmak için yeni bir kayıt politikası olarak veri kümelerinin ilk  $\theta \in \{2, 4, 6, 8, 10\}$  imzalarının alınarak eğitimde kullanılmasını da önerilmiştir. Bu şartlar altında ek olarak, MCYT ve CEDAR veri setlerinin

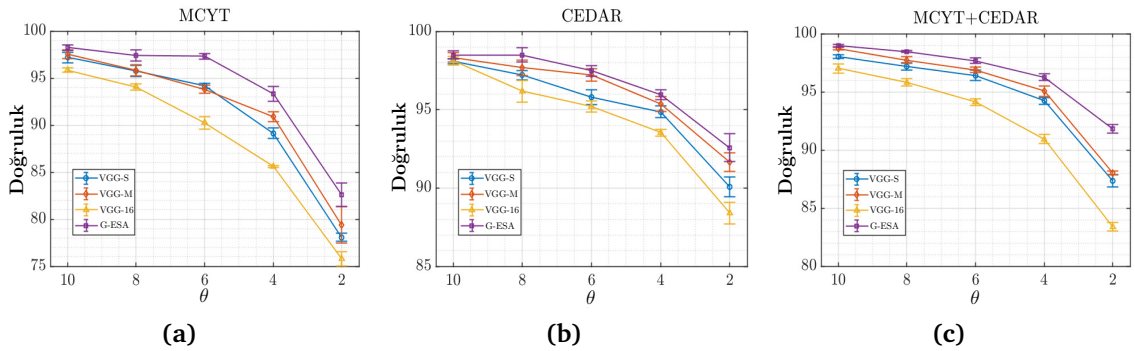
**Tablo 3.4** MCYT, CEDAR ve MCYT+CEDAR veri setlerinde bulunan kişilerden ilk  $\theta \in \{2, 4, 6, 8, 10\}$  imzanın eğitim için alınarak yapılan testlerin sonuçları

| Veri Seti  | Ağlar  | Doğruluk      |               |               |               |               | Makro F1      |               |               |               |               |
|------------|--------|---------------|---------------|---------------|---------------|---------------|---------------|---------------|---------------|---------------|---------------|
|            |        | 2             | 4             | 6             | 8             | 10            | 2             | 4             | 6             | 8             | 10            |
| MCYT       | VGG-S  | 78,07%        | 89,14%        | 94,19%        | 95,73%        | 97,17%        | 77,34%        | 89,01%        | 94,08%        | 95,72%        | 97,17%        |
|            | VGG-M  | 79,41%        | 90,91%        | 93,81%        | 95,81%        | 97,55%        | 78,90%        | 90,85%        | 93,72%        | 95,85%        | 97,55%        |
|            | VGG-16 | 75,82%        | 85,60%        | 90,25%        | 94,06%        | 95,84%        | 75,28%        | 85,53%        | 90,09%        | 94,05%        | 95,77%        |
|            | G-ESA  | <b>82,58%</b> | <b>93,33%</b> | <b>97,30%</b> | <b>97,41%</b> | <b>98,24%</b> | <b>82,17%</b> | <b>93,26%</b> | <b>97,29%</b> | <b>97,40%</b> | <b>98,24%</b> |
| CEDAR      | VGG-S  | 90,08%        | 94,85%        | 95,80%        | 97,20%        | 98,08%        | 89,79%        | 94,79%        | 95,77%        | 97,17%        | 98,05%        |
|            | VGG-M  | 91,65%        | 95,36%        | 97,21%        | 97,68%        | 98,34%        | 91,32%        | 95,29%        | 97,21%        | 97,67%        | 98,33%        |
|            | VGG-16 | 88,40%        | 93,53%        | 95,19%        | 96,18%        | 98,13%        | 88,09%        | 93,47%        | 95,12%        | 96,13%        | 98,12%        |
|            | G-ESA  | <b>92,58%</b> | <b>95,96%</b> | <b>97,47%</b> | <b>98,50%</b> | <b>98,49%</b> | <b>92,40%</b> | <b>95,94%</b> | <b>97,43%</b> | <b>98,47%</b> | <b>98,47%</b> |
| MCYT+CEDAR | VGG-S  | 87,34%        | 94,21%        | 96,42%        | 97,18%        | 98,03%        | 85,33%        | 93,38%        | 95,94%        | 96,94%        | 97,72%        |
|            | VGG-M  | 88,01%        | 95,05%        | 96,89%        | 97,74%        | 98,72%        | 85,78%        | 94,36%        | 96,39%        | 97,41%        | 98,50%        |
|            | VGG-16 | 83,40%        | 90,93%        | 94,11%        | 95,81%        | 97,01%        | 81,06%        | 89,53%        | 93,11%        | 95,15%        | 96,20%        |
|            | G-ESA  | <b>91,82%</b> | <b>96,26%</b> | <b>97,68%</b> | <b>98,46%</b> | <b>98,97%</b> | <b>91,00%</b> | <b>95,82%</b> | <b>97,35%</b> | <b>98,34%</b> | <b>98,77%</b> |

harmanlanması ile oluşturulan 2445 imza örneği ve 130 gerçek kişiden içeren yeni bir hibrit veri setini de test ortamına dahil edilmiştir.

Ağların ağırlık ilklendirilmesi [103] yöntemine göre rastgele olduğundan dolayı ağlar, her eğitiminde başka bir yerel minimuma göre sonuç çıkartmaktadırlar. Bu nedenle modeller yine 5 kere çalıştırıldıktan sonra çıkan başarımların ortalaması alınarak Tablo 3.4 oluşturulmuştur. Standart sapmalarda Şekil 3.5’de gösterilmiştir. Tablo 3.4’te verilen değerlerden görüleceği üzere G-ESA , diğer ağlara nazaran önemli bir üstünlük elde etmiştir. Özellikle eğitim için sadece 2 imza alınması ile birlikte CEDAR ve CEDAR+MCYT için %90’ın üzerinde bir doğruluk elde edilmiştir. Şekil 3.5’de eğitim için alınan imza sayısı arttıkça bütün ağların başarımların yüzdeleri, büyük bir artış göstermektedir.

Ağların gerçek kişiler tarafından oluşturulan veri setlerinin testinin ardından geniş-ölçek skalasında GPDS-4k ile testleri gerçekleştirilmiş ve kapsamlı bir analiz



**Şekil 3.5** MCYT, CEDAR ve MCYT+CEDAR veri setleri üzerinden eğitim aşamasında kişi başına ilk  $\theta \in \{2, 4, 6, 8, 10\}$  imzanın alınması ile eğitilen modellerin test başarımlarının kıyaslanması

ardından elde edilen sonuçlar Tablo 3.5’de sunulmuştur. Tablo 3.5, dört ana gruptan oluşmaktadır. Her grup kendi içinde dört ağın kıyaslamasını içermektedir. Tablo içindeki her  $\Delta\%0,1$  değişim 25-75 oranı için 75, 50-50 için ise 48 imzaya denk gelmektedir. İlk grupta BN kullanılmayarak ağların başarımları sergilenmiştir. İkinci grupta ise BN kullanımının (+ **BN**) tanıma doğruluklarına olan katkıları irdelenmiş, üçüncü grupta da *SMTS* algoritmasının (+ **SMTS**) başarımları iyileştirmesi gösterilmiştir. Son grupta, *G-ESA + BN + SMTS* yapısının adı *G-ESA\_v2* olarak isimlendirilmiş ve ReLU için sızıntı parametresinin analizi yapılmıştır.



**Tablo 3.5** GPDS-4k veri seti için ağların detaylı karşılaştırılması. Başarım metriklerindeki % $\Delta$  0,1’lik değişim, 25-75 oranlarında 75 imza ve 50-50 oranları için 48 imzaya karşılık gelir. *Eğitim* sütununda [sn] cinsinden bulunan değerler, ağların iki farklı oran için bir eğitim devrinde tükettikleri zaman miktarlarıdır. *Test* sütunu ise ağların 5k örnek ile test edilmesi için gereken süreyi ifade eder. *Ağ [MB]* parametresi, ağın toplam parametre büyüklüğünü verirken, *Veri [MB]* ise eğitim fazında ağın ihtiyaç duyduğu veri işleme alanı ile ilgilidir

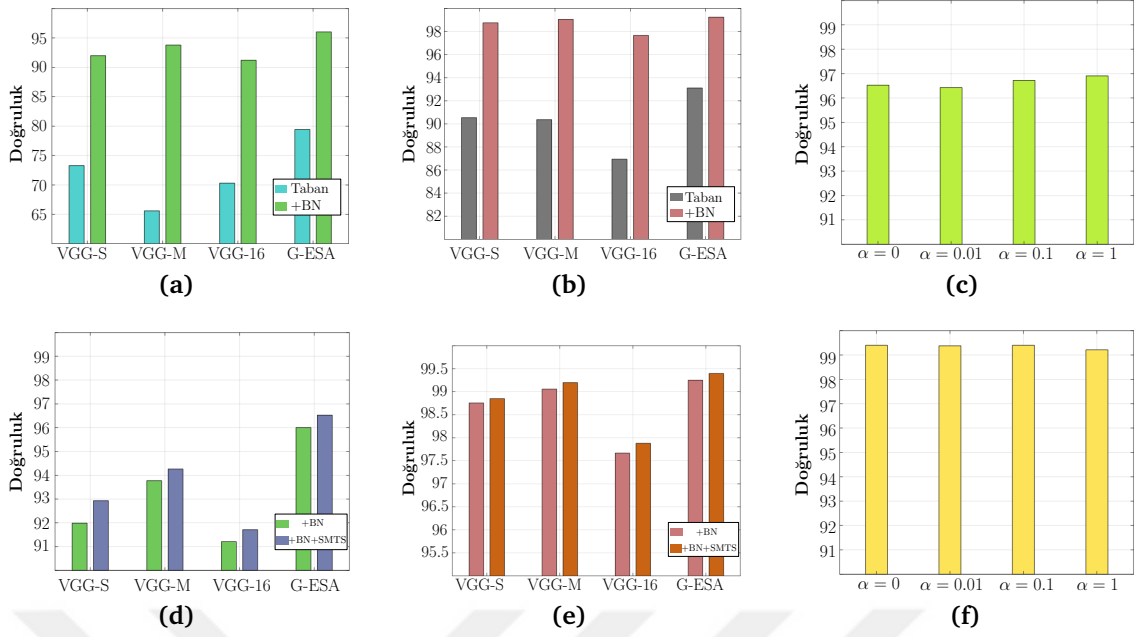
| Model                       | Gereksinimler |              |           |               |            | 25-75         |               |               | 50-50         |               |               |
|-----------------------------|---------------|--------------|-----------|---------------|------------|---------------|---------------|---------------|---------------|---------------|---------------|
|                             | Eğitim [sn]   |              | Test [sn] | Ağ [MB]       | Veri [MB]  | Doğruluk      | Ort. Kesinlik | Makro F1      | Doğruluk      | Ort. Kesinlik | Makro F1      |
| VGG-S                       | 144           | 285          | 32        | 431           | 840        | 73,29%        | 74,58%        | 72,64%        | 90,54%        | 91,30%        | 90,32%        |
| VGG-M                       | 110           | 221          | 26        | 439           | 570        | 65,59%        | 67,44%        | 65,01%        | 90,37%        | 91,43%        | 90,21%        |
| VGG-16                      | 706           | 1366         | 57        | 575           | 3270       | 70,30%        | 71,74%        | 69,57%        | 86,94%        | 87,88%        | 86,53%        |
| G-ESA                       | 108           | 218          | 27        | 471           | 780        | <b>79,42%</b> | <b>80,92%</b> | <b>78,85%</b> | <b>93,11%</b> | <b>93,73%</b> | <b>92,98%</b> |
| VGG-S + BN                  | 168           | 341          | 32        | 432           | 900        | 91,46%        | 91,99%        | 91,84%        | 98,88%        | 98,76%        | 98,75%        |
| VGG-M + BN                  | 131           | 264          | 26        | 440           | 600        | 93,77%        | 94,15%        | 93,66%        | 99,06%        | 99,16%        | 99,05%        |
| VGG-16 + BN                 | 1115          | 2231         | 57        | 576           | 4830       | 91,21%        | 91,73%        | 91,01%        | 97,67%        | 97,86%        | 97,64%        |
| G-ESA + BN                  | 142           | 285          | 27        | 472           | 810        | <b>96,01%</b> | <b>96,28%</b> | <b>95,96%</b> | <b>99,25%</b> | <b>99,34%</b> | <b>99,25%</b> |
| VGG-S + BN + SMTS           | $\Delta 172$  | $\Delta 300$ | 41        | $\Delta 62,5$ | aynı (+BN) | 92,93%        | 93,38%        | 92,83%        | 98,97%        | 98,85%        | 98,84%        |
| VGG-M + BN + SMTS           | $\Delta 127$  | $\Delta 235$ | 34        |               |            | 94,26%        | 94,59%        | 94,16%        | 99,20%        | 99,28%        | 99,18%        |
| VGG-16 + BN + SMTS          | $\Delta 226$  | $\Delta 430$ | 59        |               |            | 91,71%        | 92,20%        | 91,54%        | 97,88%        | 98,08%        | 97,88%        |
| G-ESA + BN + SMTS           | $\Delta 133$  | $\Delta 239$ | 35        |               |            | <b>96,53%</b> | <b>96,79%</b> | <b>96,50%</b> | <b>99,40%</b> | <b>99,46%</b> | <b>99,39%</b> |
| G-ESA _v2 - $\alpha : 0$    | $\Delta 133$  | $\Delta 239$ | 35        | $\Delta 62,5$ | aynı (+BN) | 96,53%        | 96,79%        | 96,50%        | <b>99,40%</b> | <b>99,46%</b> | <b>99,39%</b> |
| G-ESA _v2 - $\alpha : 0,01$ | $\Delta 142$  | $\Delta 260$ | 38        |               |            | 96,43%        | 96,70%        | 96,40%        | 99,38%        | 99,45%        | 99,38%        |
| G-ESA _v2 - $\alpha : 0,1$  | $\Delta 142$  | $\Delta 260$ | 38        |               |            | 96,72%        | 96,97%        | 96,69%        | 99,40%        | 99,76%        | 99,39%        |
| G-ESA _v2 - $\alpha : 1$    | $\Delta 142$  | $\Delta 260$ | 38        |               |            | <b>96,91%</b> | <b>97,13%</b> | <b>96,88%</b> | 99,21%        | 99,30%        | 99,20%        |

Ayrıca Tablo 3.5'te ağların eğitim ve test maliyetleri *Gereksinimler* bölümünde listelenmiştir. Eğitim fazında ağlar tarafından bir eğitim devri boyunca harcanan zaman miktarı *Eğitim [sn]*'de verilmiştir. *Test [sn]* sütunu ise, ağların 5k örnek miktarı ile test edildiklerinde sarfedilen zaman tüketimlerinden oluşur. *Net [MB]* sütunu, ağların parametre boyutlarıyla ilgilidir ve *Veri [MB]*, eğitim aşamasının 30 mini-yığın için ihtiyaç duyduğu parametre boyutunu belirtir. Üçüncü ve dördüncü dördümlü grup içinde bulunan  $\Delta$ , BN eklenmiş ağın eğitilmesinden sonra SMTS algoritmasında sınıf merkezlerinin hesaplanması için gereken ek süreyi ifade eder. Benzer şekilde, *Gereksinimler* içinde bulunan  $\Delta$  ise SMTS için ek bellek yükünü temsil eder. 62.5 MB, sınıf merkez vektörlerinin depolandığı **M** matrisinin bellek miktarıdır.

İlk grupta BN kullanılmaması nedeni ile ağların başarımları %80'in altındadır. Bu grup içindeki ağlara taban ağlar denilmektedir. VGG-16 her ne kadar en fazla parametreye sahip olsada, taban ağlar arasında başarımları %70, 30'da kalmıştır. Bu durum parametre miktarının arttırılmasının her zaman başarımları iyileştiremeyeceğinin bir göstergesidir. Taban ağlar arasında en düşük başarımları VGG-M göstermiştir.

İkinci grupta ise, ağlara BN katmanının eklenmesiyle başarımlar değerlerinde gözle görülür bir artış oluşmuştur. BN, *conv* ve *ReLU* katmanları arasına eklenir ve kendisine gelen mini-yığını sıfır-merkezli ve birim varyanslı yaptıktan sonra aynı mini-yığını, kendisinin akabinde gelen *ReLU* aktivasyon fonksiyonunun doyum bölgesine uygun ve sınıflandırmayı iyileştirecek olan ( $\gamma$ ) ile çarpar ve ( $\mu$ ) değerini ekler. BN'nin istatistiklerde yaptığı bu düzenleme sayesinde çapraz-kanal normalizasyonu (cross channel normalization) [35] veya nöron indirgeme (dropout) [104] katmanlarının kullanımına gerek kalmaz. Ayrıca BN kullanımı, ağırlık ilklendirmesinden kaynaklanan problemlerin giderilmesinde de yardımcı olur. Şekil 3.6'de, BN'nin etkisi daha net gözlemlenmektedir. Özellikle VGG-M, taban ağlar arasında  $\Delta$ %28,18 artış ile en önemli gelişmeye sahiptir. G-ESA ise hem 25-75'te hem de 50-50'de %96,01 ve %99,25 ile en iyi tanıma başarımlarını elde eden ağıdır. BN için diğer bir karşılaştırma ise Şekil 3.7'de verilmiştir. BN eklenmiş ağlarda bir eğitim adımı için zaman tüketimi artsa da öğrenme hızındaki artış bu durumu dengelemektedir. Eğitim fazında sadece 5-10 eğitim adımı geçtikten sonra kararlı bir öğrenme seviyesine gelen ağlar bunun güzel bir kanıtıdır.

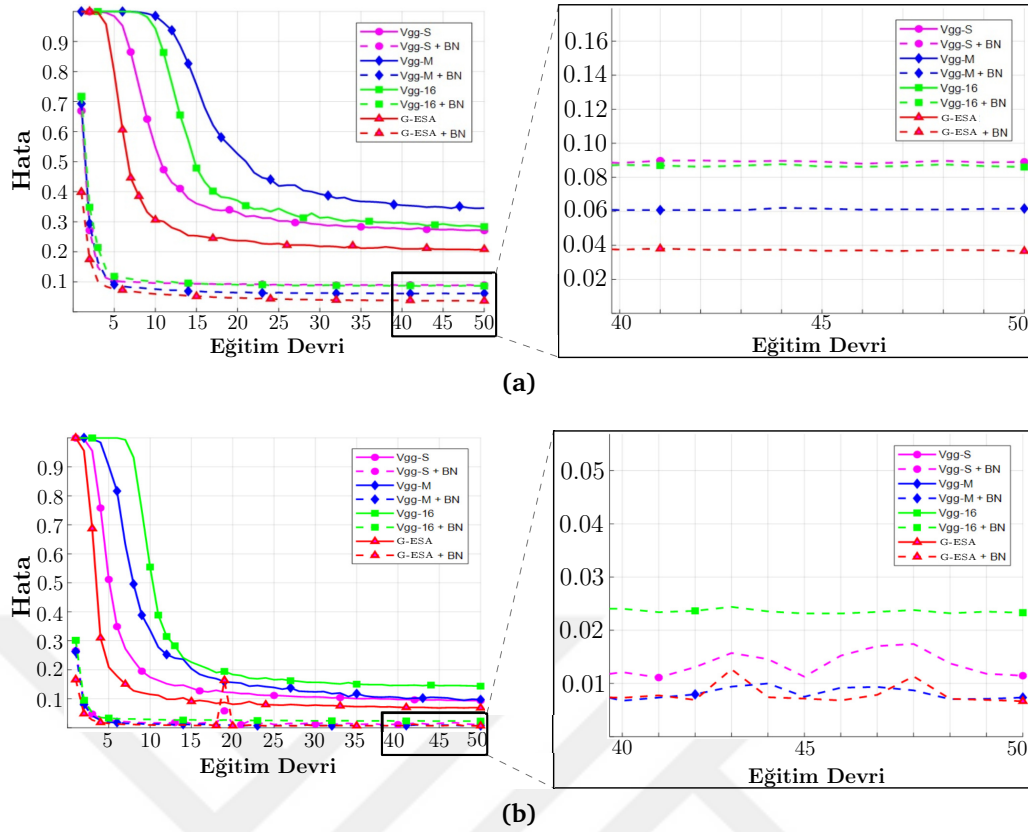
Üçüncü grupta, SMTS algoritmasının katkıları göz önüne serilmiştir. SMTS sınıflandırmada, ESA tarafından karar verme amacı ile oluşturulan üst-düzlemlere alternatif olarak test için elde edilen EFV'nin sınıf merkezlerine olan uzaklığını baz alır. Tablo 3.5'ten görüleceği üzere SMTS ile birlikte bütün ağların başarımlar değerleri artmıştır. Ayrıca bu artış Şekil 3.6(d) ve Şekil 3.6(e)'de de gözlemlenmektedir. Tüm ağlar için ortalama artış 25-75 için %0,5, 50-50 için %0,17 seviyelerindedir.  $\Delta$ %0,1



**Şekil 3.6** Taban, BN ekli ve SMTS uygulanan ağların başarımlarını karşılaştırmaları: (a) ve (d): sırasıyla 25-75 için Taban - BN ekli ağlar ile BN ekli - SMTS uygulanmış ağların sonuçları. (b) ve (e) ise, 50-50 için Taban - BN ekli ağlar ile BN ekli - SMTS uygulanmış ağların sonuçlarını ifade eder. Son olarak (c) ve (f):  $\alpha$  değerlerine göre 25-75 ve 50-50 için G-ESA\_v2 sonuçlarıdır

farkın, 25-75 için 75 imzaya ve 50-50 için 48 imzaya karşılık gelir. *G-ESA + BN + SMTS* bu dörtü içinde 25-75 için en yüksek %96,53 tanıma doğruluğuna erişmiştir. SMTS algoritmasının getirdiği maliyetler incelendiğinde, ikinci dörtlülğe göre eğitim için gerekli olan zaman artışı olduğu not edilmelidir. Bu Tablo 3.5'te Eğitim [sn] sütununda  $\Delta$  olarak ifade edilmiştir. Diğer taraftan ağ parametresi gereksinimi için sınıf merkezlerinin saklandığı matris eklendiğinden dolayı  $\Delta 62.5$  MB bir ek alana ihtiyaç duyulmuştur.

SMTS algoritmasının eklenmesiyle birlikte, varsayılan ağ *G-ESA + BN + SMTS* model yapısı *G-ESA\_v2* olarak adlandırılmıştır. Son grupta bulunan analizler ise sızdıran ReLU'nun *G-ESA\_v2* ağı üzerindeki etkilerine dayanmaktadır. Tablo 3.5'te  $\alpha \in \{0, 0.01, 0.1, 1\}$  değerleri sızıntı oranı için kullanılmıştır. Sızdıran ReLU, sırasıyla  $\alpha = 0$  ve  $\alpha = 1$  olduğunda, ReLU ve doğrusal aktivasyon gibi davranır. Bu nedenle,  $\alpha = 0$  olduğunda, *G-ESA\_v2* başarımlar ölçütleri ölçütleri üçüncü bölümde bulunan *G-ESA + BN + SMTS* ile aynıdır. Model için çalışma kapsamındaki en yüksek değer, Şekil 3.6(c)'den görüleceği üzere 25-75 için %96,91, olarak  $\alpha = 1$  değeri için kaydedilmiştir. Sızıntı faktörü değişimi, düşük miktarda eğitim örneği ile yapılan test senaryosundaki başarımlarını  $\Delta \%0,38$  arttırmıştır. 50-50 oranı için ise en yüksek değer  $\alpha = 0$  için %99,40 olarak hesaplanmıştır.



**Şekil 3.7** Taban ve BN ekli ağların eğitim hızlarının karşılaştırılması: **(a)** ve **(b)**: sırasıyla 25–75 ve 50–50 oranları için yakınsaklık hızı analizidir. BN eklenen ağların bir eğitim devrindeki zaman tüketimi BN eklenmemiş ağlara nazaran daha fazla zaman tüketimi olmasada (Tablo 3.5), tüm eğitim sürecindeki öğrenme hızları, eklenmemiş ağlardan çok daha yüksektir

### 3.3 Sonuç

Tezin bu bölümünde, çok sınıflı geniş-ölçek verilerin az sayıda örneklem ile ESA tabanlı bir model üzerinden literatüre katkı sağlamıştır. Geliştirilen ESA modeli, gerçek kişilerden oluşan MCYT, CEDAR ve MCYT+CEDAR veri setleri üzerinden test edilmiştir. Bu veri setleri üzerinden kişi başına  $\theta \in \{2, 4, 6, 8, 10\}$  düşük miktarda eğitim örneği kullanılarak eğitim ve testler yapılmıştır. Ortaya çıkan sonuçlar G-ESA modelinin literatürde kullanılan benzer ağlara göre daha iyi sonuçlar verdiğini göstermiştir. Geniş-ölçek kapsamında ise GPDS-4k veri seti kullanılmış ve önerilen ESA modelinin ve SMTS algoritmasının katkıları sunulmuştur. SMTS algoritması EFV'leri kullanarak sadece sınıf merkezlerini baz alarak bir sınıflandırma yapmaktadır. Bu sayede çalışmada kullanılan bütün ağlarının başarımlarının arttığı gözlemlenmiştir. Bir başka katkı olarak model içinde aktivasyon fonksiyonundaki sızdırma değerinin en iyilemesi yapılarak 25-75 için G-ESA\_v2, en yüksek tanıma başarımı olan %96, 91'i elde etmiştir. Ayrıca en derin ve en büyük parametre büyüklüğüne sahip ağların her zaman başarılı olmadığı gösterilmiştir. Veri setleri için özel olarak tasarlanmış ağlar

daha iyi sonuçlar verdiđi ön plana sürölmüştür. Bu çalışma kapsamında oluşturulan bir makale Neurocomputing dergisinde kabul almıştır (10.1016/j.neucom.2019.03.027). Ayrıca geliştirilen yöntem ile ilgili 25.Sinyal İşleme ve Uygulamaları Kurultayı (SIU'17) konferansında bir bildiri sunulmuştur (10.1109/SIU.2017.7960454).



## Az Sınıflı Yüksek Çözünürlüklü Görsel Verilerin İstatistiksel Modele Dayalı Sınıflandırılması

---

Görsellerin sınıflandırılması çoğu zaman örüntü tabanlı (kenar, köşe, desen) öznitelik çıkartma yöntemleri ile yapılmaktadır. Sınıflar arasındaki farklılıkların geometrik öznitelikler üzerinden tanımlandığı problemler için başarılı olan bu yöntemler yüksek çözünürlüklü ve örüntü bağımsız problemler (hücre dağılımları, hiperspektral imzalar, renk uzayında yapılan bölütleme) için işlem yükü ve çözüm sistematigi açısından elverişli değildir. Diğer taraftan özellikle renkli görüntülerin bölütlenmesinde kullanılan istatistiksel yöntemler, veri çokluğunu bir avantaja dönüştürerek sınıflar arası ayırt edici istatistiksel modellerin geliştirilmesine olanak verirler. Bu yöntemlerin başka bir katkısı ise örüntü bağımsız problemlere karşı etkili bir çözüm ortaya koymakla birlikte, işlem yükü bakımından da önemli kazanımlar sağlarlar. Bu çerçevede Bölüm 4 dahilinde, az sınıflı ve yüksek çözünürlüğe sahip görsellerin sınıflandırılmasına yönelik özgün bir metodun geliştirilmesi ele alınmıştır.

### 4.1 Literatür Özeti

İstatistiksel öznitelikler birçok alanda kullanılan, veriyi dağılım bilgisi üzerinden betimleyen önemli yöntemlerdir. Hesaplama karmaşıklığı açısından kolay elde edilmesi ve probleme özgü geliştirmelere açık olması nedeniyle geniş bir alanda kullanılmaktadır. Prieto ve ark. [105]'de, rulman bozulmaları için elde ettikleri titreşim işaretleri üzerinden standart sapma, ortalama, en düşük ve yüksek değer, kurtosis gibi 15 tane istatistiksel özniteliği çıkartmışlardır. Ardından eğrisel bileşen analizi ile (curvilinear component analysis) özniteliklerin iz düşümü alınmıştır. Doğrusal olmayan bu yöntem sayesinde sınıf içi değişinti minimize edilmiştir. İz düşüm sonrasında ortaya çıkan öznitelik hiyerarşik MLP ile sınıflandırılmıştır. Latif ve ark. [106] ise beyin tümörüne ait manyetik rezonans görüntülerinin sınıflandırılmasında dalgacık dönüşümü ile elde ettikleri 100 özniteliğe ek olarak birinci ve ikinci dereceden momentleri barındıran 52 tane istatistiksel özniteliği

kullanmışlardır. Dalgacık dönüşümü sayesinde örüntüsel bilgiler çıkartılırken istatistiksel öznitelikler dağılım tabanlı bilgi sağlamaktadır. Toplam 152 öz niteliğin kullanıldığı çalışmada MLP sınıflandırıcısı ile %96,72 doğruluk değerine ulaşmışlardır. [107]'teki çalışmada ise ses işaretlerinde konuşma olan yerlerin tespiti için kısa-zaman sıfır-geçiş oranı, kepsral tepe sayısı, enerji ölçümü gibi frekans bölgesine ait istatistiksel öznitelikler çıkartılmıştır. Bu öznitelikler, 10 ms ses parçaları üzerinden dalgalanma frekansının (pitch frequency) tespitinde kullanılmışlardır. Algoritmanın testi için 10 ms 50k konuşma çerçevesine sahip TIMIT veri seti kullanılmıştır. Dong ve ark. [108] sentetik açıklık radarı (syntetic apperture radar (SAR)) üzerinden elde edilen görüntülerin sınıflandırılması için polarimetrik öznitelikler üzerinden istatistiksel bir model öne sürmüşlerdir. Geliştirdikleri yöntem, çok-değişkenli gerçek değerlere sahip değişinti matrisi (covariance matrix) üzerinden istatistiksel model olarak alfa-durağan dağılımı kullanmaktadır. Kara, su, tarım alanı ve binalara ait örüntü bağımsız SAR görüntüleri üzerinden elde edilen modeller simetrik KL diverjansı ile karşılaştırılarak sınıflandırılmıştır. Yöntemin başarımlar analizi ise RADARSAT-2 ile çekilen  $1600 \times 1200$  piksellik görüntüler aracılığı ile yapılmış ve yöntem 11 sınıf üzerinden %83,25'lik genel başarımlar ortalamasını yakalamıştır. Yazarlar ayrıca ortaya koydukları çalışmanın SAR haricinde LIDAR ve hiperspektral görüntülere uygulanabileceğini iddia etmişlerdir. İstatistiksel özniteliklerin kullanıldığı bir diğer alan ise iris görüntülerinin sınıflandırılmasıdır. Basal ve ark. [108]'teki çalışmada gözden ayrıştırılan iris görüntülerini radyal dönüşümle iki boyutlu yassı bir görüntü hale getirmişlerdir. Önerdikleri yöntemde satır ve sütun değerleri üzerinden ayrı ayrı ortalama, standart sapma, kurtosis gibi 6 istatistiksel öznitelik çıkartılmıştır. Hamming mesafe ölçütünün kullanıldığı sınıflandırma aşamasında başarımlar ölçütü için CASIA-Iris-V4 [109] veri seti tercih edilmiştir. [110]'daki çalışmada ise döküman içindeki metin ve resimlerin bölütlenmesine yönelik bulanık c-ortalama (Fuzzy c-Means (FCM)) tabanlı bir algoritma önerilmiştir. İlk olarak döküman RGB renk uzayından dönüşüm matrisi ile YCbCr renk uzayına haritalandırılmıştır. Bu uzayda döküman üzerinden alınan yerel ortalama ve standart sapma gibi istatistiksel öznitelikler metin ve görüntü için ayırt edici niteliktedir. Öznitelikler üzerine FCM uygulandığında sınıflar kolaylıkla tespit edilebilmektedir. Bu çalışmada da önemli olan sınıflandırma için ele alınan problemin örüntü bağımsız bir yapıda olmasıdır.

Bu çalışmalar göz önünde bulundurulduğunda, yüksek çözünürlüğe sahip görüntülerin analizi ve sınıflandırılması üzerinde çalışılması gereken değerli bir alan olarak tespit edilmiştir. Böylesi görüntülerde örüntü tabanlı yöntemler fazla işlem yükü gereksinimleri nedeniyle gerçek hayat problemleri için pratik değildirler. Dolayısıyla, geniş-ölçek yoğun veri içeren görüntülere yönelik istatistiksel bir sınıflandırıcı tasarımı için özgün bir metot geliştirilmiştir.

## 4.2 Yerel Histograma Dayalı Simetrik Kullback-Leibler Diverjansı

Görüntüler üzerinden elde edilen histogram verisi önemli istatistiksel bilgiler içerir. Problemin özelliğine göre genel olarak elde edilebildiği gibi yerel olarakta kullanılmaktadır. Yerel histogram kullanımı ile görüntüde bölgesel analiz imkanı hale gelir ve daha gürbüz öznelilik çıkartılmasının önü açılır. Yerel histogram temelli algoritmalara örnek olarak HoG verilebilir. İşlenecek olan resmi bloklara böldükten sonra her blok içindeki yönelimlerin açisal histogramını hesaplayarak çalışan bu algoritma, ayırt edici öznelilikler sağlar. Bu çalışma kapsamında, herhangi bir büyüklük üzerinden elde edilen yerel histogramların sınıflandırılmasına yönelik algoritma tasarlanmıştır.

Yerel histogramlar, problemin doğasına uygun olarak pikseller üzerinden veya görüntünün dönüştürüldüğü farklı bir uzay içerisinde kullanılarak çıkartılabilir. Piksel tabanlı bir yöntem ele alındığında simetrik KL diverjansı,

$$D_{KL}(\mathbf{q}_i, \mathbf{p}_i) = \frac{1}{2} \left[ \sum_{x=0}^{255} \mathbf{q}_i(x) \log \frac{\mathbf{q}_i(x)}{\mathbf{p}_i(x)} + \mathbf{p}_i(x) \log \frac{\mathbf{p}_i(x)}{\mathbf{q}_i(x)} \right] \quad (4.1)$$

olarak tanımlanır. (4.1)'te bulunan  $\mathbf{q}$  ve  $\mathbf{p}$  vektörleri  $i$ . bölge için karşılaştırması yapılan olasılık kütle fonskiyonlarıdır (OKF). İki görüntünün benzerliğinin ölçülmesi ise Algoritma 5'de verilmiştir. Farklı görüntülere ait,  $N_b$  tane bölgeden elde edilen histogramların tutuldukları  $\mathcal{Q}$  ve  $\mathcal{P}$  kümeleri (4.1) ile karşılaştırılır. Her bölgeden elde edilen OKF'ler üzerinden uzaklık hesaplanır. Ardından bulunan uzaklıkların ortalaması alınarak görüntüler arası benzerlik bulunur.

Elde edilen bu bilgiler kullanılarak birçok farklı sınıflandırma yöntemi geliştirilebilir. Tez planında metodun geliştirilmesi ise veri tabanı içindeki örnekler ile test örneği arasında hesaplanan uzaklık değerlerinin  $k$ -NN tabanlı bir algoritma ile sınıflandırılması yönünde olmuştur. Herhangi bir görüntünün eğitim veri seti içindeki  $N_{trn}$  örnek ile karşılaşmasından  $\mathbf{d}_{KL} \in \mathbb{R}^{N_{trn} \times 1}$  benzerlik vektörü çıkartılır. Mesafe

---

### Algoritma 5 Farklı iki görüntünün yerel histogramlar üzerinden benzerlik hesabı

---

**Girdi**  $\{\mathbf{q}_1, \mathbf{q}_2, \mathbf{q}_3, \dots, \mathbf{q}_{N_b}\} \in \mathcal{Q}$ ,  $\{\mathbf{p}_1, \mathbf{p}_2, \mathbf{p}_3, \dots, \mathbf{p}_{N_b}\} \in \mathcal{P}$      $\triangleright$  Bölgesel Histogramlar  
**Çıktı**  $d_{kl}$

1: **Prosedür** OLASILIKSAL BENZERLİK METRİĞİ

2:      $\phi \leftarrow 0$

3:     **for**  $i = 1 : N_b$  **do**

4:          $\phi = D_{KL}(\mathbf{q}_i, \mathbf{p}_i) + \phi$

5:

6:     **dönüş**  $d_{kl} = \phi / N_b$

---

---

**Algoritma 6 Ağırlıklı -  $k$ -NN ( $wkNN$ )**

---

**Girdi**  $\mathbf{d}_{kl} \in \mathbb{R}^{N_{trn} \times 1}$ ,  $\mathbf{g} \in \mathbb{R}^{N_{trn} \times 1}$ ,  $k$

**Çıktı**  $\hat{y}$

▷ Tahmini Sınıf

1: **Prosedür** SINIF AĞIRLIKLARININ BULUNMASI

2:  $\mathbf{s} \leftarrow \min\_secim\{\mathbf{d}, k + 1\}$

▷ Artan Sıralama

3:  $\mathbf{c} \leftarrow \mathbf{g}$  üzerinden  $\mathbf{s}$  ile ilgili sınıflar

4:  $\mathbf{s}_n = \mathbf{s}(1 : k) / \mathbf{s}(k + 1)$

▷ Normalizasyon

5:  $\mathbf{s}_w = f(\mathbf{s})$

▷  $f(\cdot)$ : Çekirdek fonk.

6:  $\hat{y} = \max_r \{ \sum_{i=1}^k \mathbf{s}_w(i) \cdot I(\mathbf{c}(i) = r) \}$

▷  $I(\cdot)$ : İndikatör

7: **Dönüş:**  $\hat{y}$

---

tabanlı elde edile değerler için  $k$ -NN tabanlı algoritmalar oldukça kullanışlıdır. Bu çalışmada  $k$ -NN algoritmasının bir türevi olan *ağırlıklandırılmış  $k$ -NN* (weighthed  $k$ -NN ( $wkNN$ )) algoritması sınıflandırmasında kullanılmaktadır.

Algoritma 6'da verilen  $wkNN$ ,  $k$ -NN üzerinden türetilmiş örnek tabanlı ağırlıklandırmaya sahip bir sınıflandırma algoritmasıdır. Benzerlik kriteri minimum mesafe ölçütü üzerinden yapılmaktadır.  $wkNN$  test örneğinin eğitim örnekleri arasında hesaplanan mesafeleri arasından bir çekirdek (kernel) fonksiyonu ile örnek tabanlı ağırlıklandırma yapmaktadır. [111] çalışmasında  $wkNN$  için birçok çekirdek fonksiyonu önerilmiştir. Bu çalışma kapsamında  $f(x) = 1 - |x|$  üçgen çekirdeği skor atama için kullanılmıştır.  $f(x)$  için hesaplanan  $x$  mesafesinin  $[0, 1]$  arasında olması önemlidir. Bunun için algoritmada  $k + 1$  tane minimum mesafe değeri sıralı olarak alınmaktadır. Ardından seçilen mesafeler arasında en büyük değere sahip olan  $k + 1$ . mesafesi ile diğer örnekler normalize edileren  $x$  değerleri istenilen aralığa getirilmiş olurlar. Normalizasyon işleminden sonra  $f(x)$ ,  $x$  mesafesi ile ilgili olan sınıf için skor atamaktadır. Atanan bu skor değerleri sınıflar üzerinden toplanarak en yüksek skora sahip sınıf tahmin edilen olarak çıktı olarak verilmektedir.  $wkNN$  girdi olarak  $\mathbf{d}_{kl} \in \mathbb{R}^{N_{trn} \times 1}$  test görüntüsünün veri setindeki örneklere olan istatistiksel uzaklığı ile  $\mathbf{g} \in \mathbb{R}^{N_{trn} \times 1}$  eğitim örnekleri ile ilişkili olan etiket bilgisi ve  $k$  komşu sayısını alır.

### 4.3 Uygulama Alanı: Serviks Kanserinde Öncü Lezyonların Sınıflandırılması

Dünya Sağlık Örgütü'nün (World Health Organization) bir kuruluşu olan Küresel Kanser Gözlemevinin (Global Cancer Observatory) yayınladığı güncel istatistiklere göre yaygın görülen kanser türleri arasında sekizinci sırada yer alan serviks kanserinde yaklaşık 570.000 vakadan 311.000'i 2018 yılında ölümle sonuçlanmıştır. Düşük İnsani Gelişme Endeksi (Low Human Development Index) düzenlemesinde ise kadınlar arasında serviks, meme kanserinden sonra en sık görülen ikinci ölüm sebebidir

[112]. Pek çok araştırmada [113–115] gösterildiği üzere neredeyse tüm serviks kanseri vakalarının sebebi İnsan Papilloma Virüsünün (Human Papilloma Virus (HPV)) epitel doku içinde yer alan bazal tabakaya ulaşması ve buradaki hücreleri deforme etmesinden kaynaklıdır. HPV servikste epitel hücrelerle temas ettiğinde, hücrelerin olgunlaşması kaybolur ve bazal tabakada hızlı proliferasyon (hücre çoğalması) görülür. Bunlara ek olarak, hücrelerin çekirdeklerindeki düzensizlikler artar ve çekirdek genişlemesi ile birlikte hiperkromazi ortaya çıkar [116, 117]. Hücre proliferasyonuna paralel olarak, mitoz sayısında da artış gözlenir. Bu nedenle, hücre yoğunlaşması bazal tabakadan (segment için katman (layer) çağrışımı yaptığı için tabaka kullanıldı.) üst tabakaya doğru dağılım gösterir. Hücre iskeletindeki proteinleri etkileyen HPV, çekirdeğin etrafında büyük, hiperkromatik olan ve koilosit adı verilen açık bir alanın oluşmasına neden olur. Bu anormal değişimlerin epitelin alt, orta ve üst kısımlarda görülüp görülmediğine göre dokular derecelendirilir [118, 119].

HPV'nin etki ettiği hastalıklı dokulara lezyon denilmektedir. Lezyon içersinde hastalığın gelişim evrelerinin anlaşılması üzerine tıp literatüründe iki derecelendirme sistemi kullanılmaktadır. Bunlar *Squamous Intraepithelial Lesion* (SIL-iki aşamalı derecelendirme sistemi) ve *Cervical Intraepithelial Neoplasia* (CIN-üç aşamalı derecelendirme sistemi). Histopatolojik tanı, hücrelerin displastik değişikliklerinin düzeyine göre belirlenir. CIN sisteminde derecelendirme, anormal hücrelerin üst, orta ve bazal tabakada bulunma yoğunluğuna göre yapılmaktadır [118–120].

Histopatolojide, yüksek çözünürlüklü kameraların ve mevcut görüntü işleme tekniklerinin tanı sürecine dahil edilmesiyle yeni bir dönem başladı. Daha nesnel kararlar vermek ve gözlemci-ici ve gözlemciler arası değişkenliğin üstesinden gelmek için Bilgisayar Destekli Tanı (Computed Aid Diagnosis (CAD)) sistemlerinin etkili bir şekilde kullanılması yaygınlaştı. Bu sayede patalogların karar vermede zorlandıkları derecelendirmelerde, CAD sistemlerinden destek alarak daha isabetli tanıların verilebilme imkanı doğmuştur. Özellikle CIN derecelendirme sisteminde CIN-2 tanısı için yüksek seviyede gözlemci-ici ve gözlemciler arası değişkenlik oranı vardır. Çünkü patoloji rutininin çoğunda, CIN-2 lezyonunu CIN-1 veya CIN-3 lezyonlarından Hemotoksilen ve Eosin (H&E) ile boyanmış preparatlar [121, 122] ile ayırt etmek zor olabilmektedir. Bu problemlerin aşılabılmesinde hücre segmentasyonu ve doku sınıflandırması için bir takım ön çalışmalar yapılmıştır [123–125].

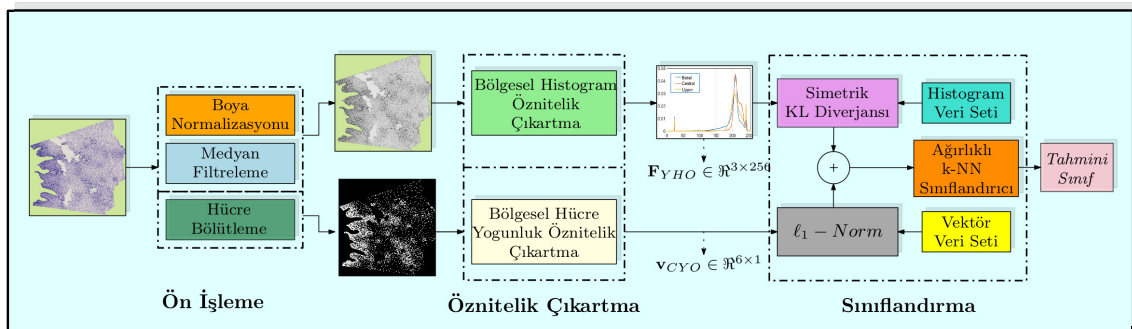
Dijital patolojide, tedavi sürecini yönlendiren servikal kanser-öncesi lezyonların derecelendirilmesi önemlidir. Bununla birlikte tanı, dokuların eğriliği ve epitel altı doku kıvrımlarının lezyon içersinde farklı bölgelerden geçmesi ile teğetsel enine kesit olarak papilla adı verilen kısımların oluşması nedeni ile karmaşık hale gelebilmektedir. Epitel içersinde papillanın olduğu bölgeler bazal tabaka içersinde

kabul edilir. Çoğu hücre tabanlı morfometrik algoritmalar [126–128] papillayı göz ardı etmektedirler. Bu durum, algoritmaları gerçek hayatta sıkça karşılaşılan girift yapıdaki epitellerin derecelendirilmesinde yetersiz kılmaktadır. Ayrıca epitel içerisinde bulunan hücrelerin öbekler halinde bulunması nedeniyle morfometrik algoritmaların çakışık hücre (overlapped cell) probleminde çözümleri gerekmektedir. Çakışık hücrelerin ayrıştırılması ise  $\times 20$  ve  $\times 40$  optik yakınlaştırmalar ile elde edilen yüksek çözünürlüklü görüntüler için oldukça maliyetli çözümlemelere gereksinim duyar. Bu bağlamda, hem papilla durumunu göz önünde bulunduran hem de çakışık hücre problemi yerine hücre dağılımları üzerinden derecelendirme yapan istatistiksel tabanlı yöntem öneriyoruz. Sınıflandırma sisteminin blok şeması, Şekil 4.1’de gösterilmiştir.

#### 4.3.1 Literatür Özeti

H&E tekniği, düşük maliyetli ve güvenilir bir yöntem olduğundan dolayı servikal lezyonların sınıflandırılmasında (derecelendirilmesinde) kullanılır. Ancak kesin tanıya sadece H&E preparatları ile ulaşmak zor olabilir. Bu gibi durumlarda, patoloğlar yardımcı bir yöntem olarak P16, Ki67 gibi immünohistokimyası kullanırlar. Bütün bu yöntemler patoloğların yorumlarına dayandığından, gözlemci-içi ve gözlemciler arası oluşan farklılıklar, üstesinden gelinmesi gereken temel problemlerdir. Bu doğrultuda hücreye ait morfolojik özelliklerin dijital verilere dönüştürülmesi ve bilgilerin CAD sistemleri ile incelenmesi patoloğların karar vermelerine yardımcı olmaktadır. Bu yöntemlerden bazıları hücreleri tekil olarak ele alırken bazıları ise hücre kümelerini analiz edip bölgesel istatistiklere göre sonuç çıkartmaktadır.

Castellanos *ve ark.* [129] çalışmasında, 46 Normal, 33 CIN-1 ve 29 CIN-3 sınıflarına ait toplam 108 histopatolojik slayt kullanmışlardır. H&E boyalı lezyonlar alt (bazal), orta ve üst tabakalara ayrıldıktan sonra her tabaka floresan altında incelenmiş ve lezyonlar, tabakaların ortalama floresans yoğunluğuna göre sınıflandırılmıştır. Ji *ve ark.* [130],



**Şekil 4.1** Servikal lezyonların sınıflandırılması için önerilen istatistiksel özniteliklere dayalı algoritmanın blok diyagramı

Servikal lezyonları kolposkopik görüntülerden sınıflandırmak için istatistiksel bir doku analiz yöntemi kullandı. Bu yöntem lezyonları, bileşke entropisi, açt entropisi, açt tepe yoğunluğu dahil 13 istatistiki özelliğı kullanarak en kısa mesafe sınıflandırıcısı ile sınıflandırır. [131] çalışmasında ise epitel doku bazal ve üst tabaka arasında 10 katmana ayrılmış ve her katmanın floreson ömrünün ortalamasına bakılmıştır. Veri seti olarak 32 H&E slaytı kullanılmıştır. Floresan ömür ortalamasının yanında hesaplanan standart sapma değeri öznitelik olarak kullanılmış ve Aşırı Öğrenme Makinesi (Extreme Learning Machine) kullanılarak lezyonlar sınıflandırılmıştır.

Zhang ve ark. [132]'daki çalışmada, servikal hücrelerin sınıflandırılması için *Deep-Pap* adında ESA sunmuşlardır. Bu çalışmada hem pap smear hem de H&E boyalı veri setleri ağı performansı test etmek için kullanılmıştır. Konstandinou ve ark. [133], servikal dokuların CIN sistemine göre derecelendirilmesi için hücrelerin istatistiksel ve örüntüye dayalı özniteliklerin kullanılmasını önermişlerdir. Bu yöntemde ilk olarak, RGB formatındaki H&E görüntüleri gri ölçekli hale dönüştürölür ve daha sonra hücreler kenar belirleme operatörleri kullanılarak bölütlenirler. Açısıl ikinci moment, entropi, ortalama, çarpıklık, histogramın kurtosisi gibi toplam 63 istatistiksel özellik hücrelerden elde edilir. [134]'da servikal lezyonların H&E görüntüsü üzerinden hücrelere ait morfolojik olan 27 farklı öznitelik çıkartılmıştır. Ortalama çekirdek alanı, sitoplazma oranı, çekirdek alanı/arka plan oranı gibi öznitelikler, LDA + SVM sınıflandırıcısı üzerinden dokunun derecelendirilmesinde tercih edilmiştir. [135]'de, biyopsi dokularının HSIL olup olmadığını belirlemek için *tek saçılım gösteren elastik ışıık* kullanılması önerilmiştir.

#### 4.3.2 Katkıları

Servikal SIL derecelendirme çalışmalarının çoğı hücre bölütlemesi üzerinedir ve sınıflandırma için yapılan çalışmalar bölütleme için yapılanlara nazaran daha az sayıda kalmaktadır. Ayrıca, epitel üzerinde tanjansiyel kesit olarak karşımıza çıkan papilla durumu sınıflandırma problemlerinde göz ardı edilmiştir. Ek olarak, hücrelerin tekil morfolojik özniteliklerinin çıkartılması, yüksek çözünürlüklü görüntülerde hesaplama yükünü arttırır. Bu çalışmanın ana amacı, papilla durumlarında içeren ve epitelin yayılımından bağımsız tekil hücreler ile ilgilenmeyen istatistiksel tabanlı bir sınıflandırıcı tasarımının geliştirilmesidir. Çalışmanın katkıları şöyle sıralanabilir:

- Epitelin bazal ( $S_B$ ), orta ( $S_M$ ) ve üst ( $S_U$ ) tabakalarından elde edilen ve *Yerel Histogram Özniteliğı (YHÖ)* olarak adlandırılan yerel histogramlar istatistiksel öznitelik olarak kullanılır. YHÖ ilgili bölge içirisine dahil olan pikseller üzerinden hesaplanır. Bu bilgi tabakanın bileşenlerine (hücre çekirdeğı,

sitoplazma, beyaz bölgeler) ait yoğunluk bilgisini vermektedir. Ayrıca hücre dağılımının tabakalara ait yoğunluğu için ise ikinci istatistiksel öznitelik olarak SegNet derin öğrenme ağı yardımıyla *Çekirdek Yoğunluğu Haritasını (ÇYH)* çıkarılır. ÇYH üzerinden *Çekirdek-Yoğunluk Özniteliği (ÇYÖ)* elde edilir.

- YHÖ ve ÇYÖ istatistiksel olarak çıkarıldıklarından, hücreleri tekil olarak ele almazlar. Bu nedenle, morfometrik yöntemlerin aşaması gereken çakışık hücre problemi ortadan kalkmış olur.
- YHÖ ile ÇYÖ epiteli farklı uzaylarda temsil etmektedirler. Bu iki farklı özniteliğin birleştirilmesi için YHÖ'ler arasında *simetrik Kullback-Leibler Diverjansı* ve ÇYÖ içi  $\ell_1$ -norm kullanılarak özgün bir mesafe metriği geliştirilmiştir.
- Mesafe metriği olarak elde edilen benzerlik ölçütü üzerinden sınıflandırma için *k-NN* ile birlikte ağırlıklandırılmış (weighted) *k-NN wkNN* kullanılmıştır.

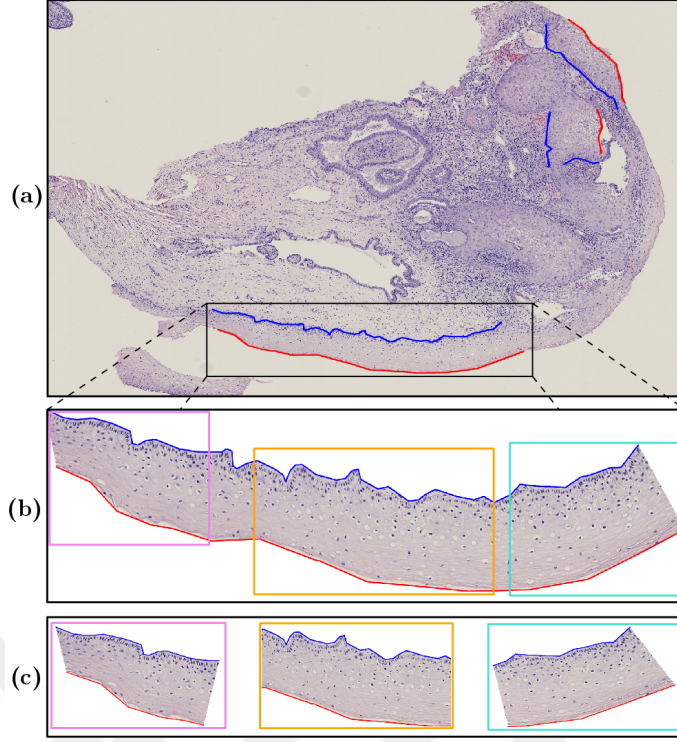
#### 4.4 Veri Seti ve Yöntemin Uyarlanması

Bu bölümde, İstanbul Medipol Üniversitesi Hastanesi Patoloji Laboratuvarında toplanan yeni bir veri seti açıklanmaktadır. Veri seti detaylandırılırken, *slayt*, *epitel* ve *küçük epitel parçası* (KEP) adlı literatüre özgü üç terim açıklanmış ve veri setinin özellikleri verilmiştir. Yöntem bölümünde, ön işleme adımları, istatistiksel olarak epitelleri temsil eden YHÖ ve ÇYÖ'nün çıkartılması ve wkNN sınıflandırıcısı ile dokuların sınıflandırılması ele alınmıştır.

##### 4.4.1 Veri Seti

Veri seti [136]'de bahsedildiği üzere 54 hastadan 127 H&E boyalı histo-patolojik görüntülerin  $\times 20$  optik yakınlaştırma ile taranması ile elde edilmiştir. Bir örneği Şekil 4.2(a)'da gösterilen herbir slaytın büyüklüğü yaklaşık 0,5 ile 2 GB arasındadır.

Slayt görüntülerinde ilgilenilen bölgenin çizimi, bazal (mavi çizgi) ve epitel yüzey (kırmızı çizgi) koordinatları uzman patologlar tarafından yapılmıştır. İlgilenilen bu bölge Şekil 4.2(b)'de gösterildiği gibi *epitel* olarak adlandırılır. Bir epitel dokusu farklı boyutlarda ve şekillerde bulunabilir. Çok kavisli ve dar veya çok düzgün ve büyük olabilir. Ayrıca aynı epitelde farklı CIN derecelerine sahip alt bölgelerde bulunabilmektedir. Bu vakaları ele almak için epitel dokuları, Şekil 4.2(c)'de görüleceği üzere KEP'lere ayrıştırılmıştır. Herbir KEP farklı epitelden türetildiği için renk ve şekil varyasyonu oldukça yüksektir. Bu durum analizleri çok daha zorlu kılmaktadır.

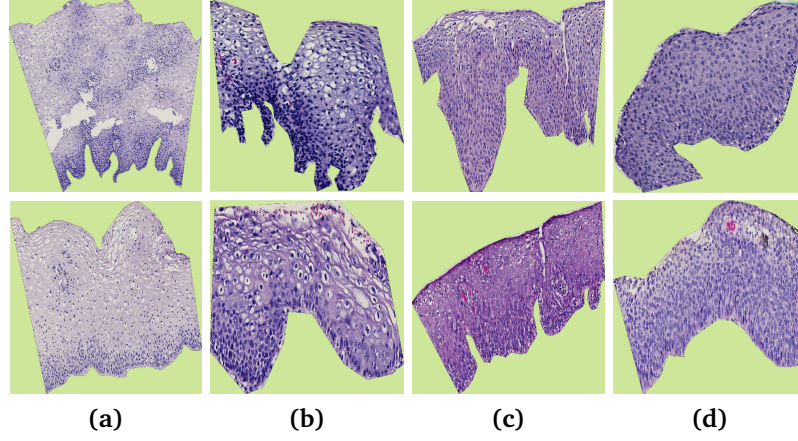


**Şekil 4.2** Veri seti içinde kullanılan kavramların açıklanması: **(a)** Yüksek çözünürlüklü H&E slayt görüntüsü, **(b)** Tanı için incelenecek olan epitel doku, **(c)** Küçük Epitel Parçaları (KEP)

İki uzman patolog epitellerin bütünü ve KEP'ler için ayrı ayrı tanı koymuşlardır. Bazı dokuların tanılarında gözlemciler arasında farklı değerlendirmelerden dolayı ihtilaflı kararlar oluşabilmektedir. Bu nedenle, patologlar bazı dokular için farklı CIN seviyeleri tanımlamışlardır. Bu durumda, kesin tanı için KEP'ler hakkında ek bilgi gerekir. Patologlar ortak bir tanı üzerinde anlaşmak için epitelin tamamına ait Ki-67 ve P16 görüntülerini incelerler. Bu nedenle, veri setinde her bir epitel veya bir KEP için üç tanı toplanır: Patolog A, Patolog B ve Kesin Tanı. Bu bağlamda, Normal, CIN-1, CIN-2 ve CIN-3 dereceleri için toplam 471, 240, 107, 139 KEP sınıflandırılmıştır. Şekil 4.3'de gösterildiği gibi elde edilen KEP'ler birbirlerinden farklı geometride ve renk skalasında yer almaktadırlar. Bu durumda gürbüz bir sınıflandırma işlemi için ön işleme adımı önemli hale gelir.

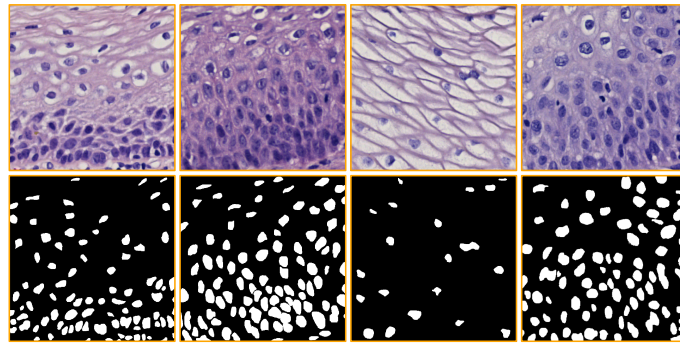
#### 4.4.2 Ham Servikal KEP Görüntüsü için Ön İşleme Süreci

H&E görüntüleri için önemli sorunlardan biri, dokuların farklı zamanlarda boyandıkları için standart bir renklendirmenin olmamasıdır. Diğer bir sorun ise görüntüleme sistemlerinin ölçüm esnasında görüntülere dahil ettiği gürültüler ve bozulmalardır. Bu etkilerin en aza indirilmesi için ön işleme adımı, Doku Normalizasyonu (DN) ile birlikte ortanca (medyan) filtreleme teknikleri kullanılmıştır.

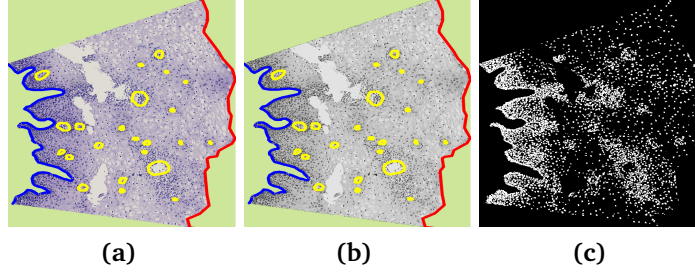


**Şekil 4.3** Veri setindeki dokuların sınıflara göre dağılımı: **(a)** 471 Normal, **(b)** 240 CIN-1, **(c)** 107 CIN-2, and **(d)** 139 CIN-3

DN, H & E görüntülerinde hücresel yapıların koyu tonlarda baskın olarak yer aldığı Hemotoksilen bandını ayırtmak için renk boyutunda dekonvolüsyon işlemi yapılmaktadır. Dekonvolüsyon işlemi sayesinde H bandını, arka plan boyasının bulunduğu E bandından ayırmak mümkün olur. DN [137–139] gibi çalışmalarda ele alınan histopatolojik görüntülerin sınıflandırılmasında aşılması gereken önemli bir safhadır. Bu çalışma kapsamında [139]’de önerilen,  $2 \times 3$  boyutlarında sabit bir dekonvolüsyon maskesi kullanılarak dokuların H bandına ayrıştırılması algoritması hızlı çalışması nedeni ile tercih edilmiştir. Ayrıştırılmış olan gri seviyedeki H bandı görüntüsü, hem tarama sisteminden hem de dönüşüm esnasında oluşan gürültülere sahiptir. Bu gürültülerin giderilmesi için  $11 \times 11$  boyutlarında bir şablon kullanılarak ortanca filtreleme yapılmıştır. Ortanca filtrenin boyutu  $\times 20$  optik yakınlaştırma oranında bir hücrenin yaklaşık boyutları göz önüne alınarak empirik olarak seçilmiştir. Bu iki adımın ardından Şekil 4.5(a) ham KEP görüntüsü Şekil 4.5(b)’deki gri seviye görüntüye dönüştürülmüştür. Şekil 4.5(b), istatistiksel özniteliklerin çıkartılması için hazırdır.



**Şekil 4.4** 4-derinlikli SegNet ağının eğitimi için üretilen veri setinden örnekler: 100 tane  $512 \times 512$  büyüklüğündeki görüntüler ile eğitilen SegNet, sadece hücreleri bulmak için kullanılmaktadır



**Şekil 4.5** Ön-işleme adımı sonrası çıktılar: Doku Normalizasyonu ve ortanca filtre kullanılarak *RGB KEP görüntüsü* (a), *Gri-seviye görüntüye* dönüştürülmüştür (b). SegNet kullanılarak *Çekirdek Yoğunluk Haritası* (c) elde edilmiştir. Mavi: bazal membran, Kırmızı: Epitel yüzey, and Sarı: Papilla sınırları

Diğer bir başka ön işleme adımı ise hücre yoğunluğunun tespit edilmesi için sadece hücreleri bölütleyen önceden eğitilmiş 4-derinliğe sahip SegNet [140] ESA'nın kullanılmasıdır. SegNet'in eğitimi için ilk olarak epitel dokulardan  $512 \times 512$  boyutlarında 100 adet parça kesilmiştir (Şekil 4.4). Bu parçalarda hücre olan yerler ikili görüntü olacak şekilde etiketlenmiş ve ardından SegNet eğitilmiştir. Eğitilmiş olan SegNet ağının girişi  $512 \times 512$  boyutlarında parçalar kabul etmektedir. Bunun için test edilecek olan gri seviye görüntü Şekil 4.5(a)  $512 \times 512$ 'lik bloklara bölünmüş ve her bir blokta hücre yoğunlukları tespit edilmiştir. Ağın çıkışında elde edilen görüntü *eşikleme + ikili açma* işleminin ardından *Çekirdek Yoğunluğu Haritası* (ÇYH) olarak adlandırılan ve Şekil 4.5(c)'de görülen ikili görüntüye dönüştürülür. ÇYH, Çekirdek-Yoğunluk Özniteliği (ÇYÖ) çıkartmak için kullanılmaktadır. Ön işleme adımları sonrasında, istatistiksel özniteliklerin elde edilmesi için gereken gri seviye görüntü ile ÇYH oluşturulmuştur. Öznitelik çıkartma algoritmaları ayrıca Şekil 4.5(a) (b)'de sarı çizgiler olarak görülen papilla sınırlarını da göz önünde bulundurarak işlem yaparlar. Papilla bölgeleri bazal katmana ait bir bölge olduğu için histogram özniteliğini önemli ölçüde etkiler.

#### 4.4.3 Dokuların İstatistiksel Öznitelikleri

Sınıflandırma sistemlerinde öznitelik çıkartma, tanıma başarımını doğrudan etkileyen temel bir adımdır. Her ne kadar sınıflandırma algoritması etkili olsada, çıkarılan özellikler ayırt ediciliğe sahip değilse, sonuçları tatmin edici olmayacaktır. Diğer bir önemli faktör, özniteliklerin makul bir süre içinde hesaplanması gerekliliğidir. Bu şartlar altında, hücre bazlı morfometrik özniteliklere alternatif olarak servikal dokuların sınıflandırılması için YHÖ bu tez kapsamında önerilmiştir. YHÖ, dokudaki hücre ve stoplazma dağılımının bölgesel yoğunluğuna göre farklılık kazanmaktadır. YHÖ, hücreleri ayrı ayrı incelemediğinden, çakışık çekirdek problemini çözömlemesine gerek yoktur. Bu nedenle özniteliklerin çıkarılması

---

**Algoritma 7 Yerel Histogram Algoritması (YHA)**

---

**Girdi:**  $\mathbf{I}_{gry} \in \mathbb{R}^{N_x \times M_y}$ ,  $N_w = 11$ ,  $\mathbf{C}_b \in \mathbb{R}^{N_b \times 2}$ ,  $\mathbf{C}_u \in \mathbb{R}^{N_u \times 2}$ ,  $\mathbf{C}_p \in \mathbb{R}^{N_p \times 2}$ .  
**Çıktı**  $\mathbf{p}_i$ ,  $i \in \{1, 2, 3\}$ .

- 1: **Prosedür** BÖLGESEL HISTOGRAMIN ÇIKARTILMASI
- 2:  $\mathcal{T}_i \leftarrow \emptyset$
- 3: **Dokuyu Katmanlara Ayırma**
- 4:  $\mathcal{G} = \{\mathbf{P}_j \in \mathbb{R}^{N_w \times N_w} \mid \mathbf{I}_{gry} \text{ ait parçalar}\}$
- 5: **for**  $\forall \mathbf{P}_j \in \mathcal{G}$  **do**
- 6:     **if** Merkez  $\mathbf{P}_j \in \mathbb{S}_B$  **then**
- 7:          $\mathcal{T}_1 \leftarrow \mathcal{T}_1 \cup \mathbf{P}_j$
- 8:     **else if** Merkez  $\mathbf{P}_j \in \mathbb{S}_C$  **then**
- 9:          $\mathcal{T}_2 \leftarrow \mathcal{T}_2 \cup \mathbf{P}_j$
- 10:    **else if** Merkez  $\mathbf{P}_j \in \mathbb{S}_U$  **then**
- 11:          $\mathcal{T}_3 \leftarrow \mathcal{T}_3 \cup \mathbf{P}_j$
- 12: **Bölgesel Histogram Çıkartımı**
- 13:  $\mathbf{p}_i \leftarrow \text{Normalize\_Histogram}\{\mathcal{T}_i\}$

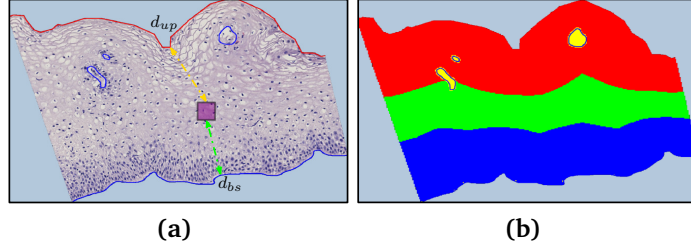
---

morfometrik yöntemlere nazaran hızlıdır.

**(i) Yerel Histogram Öz niteliği:** YHÖ, bazal membran ile epitel üzey arasında yerel hücre yoğunluğuna dayalı bir algoritmadır. HPV virüsü epitel dokuya etki ettiğinde bazal bölgedeki hücreler hızla çoğalarak orta ve üst katmanlara kadar çıkmaktadır. Bu durumda katmanlara ait renk bilgisi değişir ve katman histogramı doğrudan etkiler. Katmanlar içinde hücrelerin yoğun olduğu bölgelerde histogram koyu tonlarda baskınken, hücrelerin az, stoplazma ve koilosit adı verilen beyaz yapıların fazla olduğu yerlerde ise açık tonlar baskındır. Ayrıca hücrelerin boyutları da histogramı etkileyen diğer bir faktördür. Epitel içinde bulunan hücre, sitoplazma ve koilosit yapılarının renk bilgisi üzerinde oluşturduğu bu değişiklikler dokuların YHÖ'ler kullanılarak sınıflandırılabilceği fikrinin temelini oluşturmaktadır. YHÖ'nün diğer bir avantajı ise morfometrik öz niteliklerin çıkartılmasında karşılaşılan ve algoritmaların çalışma süresini uzatan çakışık hücre problemini, hücre popülasyonu ile ilgilenmesi nedeniyle ortadan kaldırmasıdır. YHÖ'nün çıkartılması için *Yerel Histogram Algoritması* Algoritma 7'de verilmiştir.

YHA giriş olarak üç argümana ihtiyaç duyar: 1.  $\mathbf{I}_{gry}$ , gri-seviye lezyon görüntüsü, 2.  $\mathbf{C}_b$ ,  $\mathbf{C}_u$ ,  $\mathbf{C}_p$ , bazal, üst ve eğer varsa papilla kordinatları, 3.  $N_w$ , ızgara boyutu. Ön işleme aşamasından elde edilen gri seviyeli hematoksilin görüntüsü YHÖ çıkarmak için kullanılmaktadır. Algoritma bu üç argümanı aldıktan sonra, *Dokuyu Katmanlara Ayırma* ve *Bölgesel Histogram Çıkartımı* işlemlerini yapmaktadır.

*Dokuyu Katmanlara Ayırma* aşamasında,  $\mathbf{I}_{gry}$  görüntüsü  $N_w \times N_w$  boyutlarındaki blok  $\mathbf{P}_j$  görüntülerine parçalanır. Herbir  $\mathbf{P}_j$  görüntüsünün arkaplan veya epitel dokuya ait olup

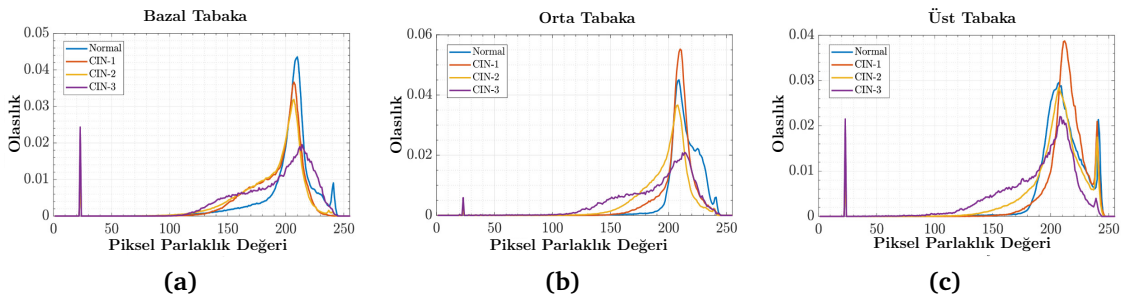


**Şekil 4.6**  $\mathbf{P}_j \in \mathbb{R}^{N_w \times N_w}$  ızgara görüntüsü (magenta) dokuya ait katmanlar üzerinden elde edilir. Bazal sınıra olan uzaklık  $d_{bs}$ , yeşil ok ile epitel yüzeye olan sınır  $d_{up}$  ise turuncu ok ile işaretlenmiştir.  $r$  değerinin aralığına göre  $([0, 1/4], [1/4, 3/4], (3/4, 1])$   $\mathbf{P}_j$ , *basal* ( $\mathbb{S}_B$ ), *orta* ( $\mathbb{S}_C$ ) or *üst* ( $\mathbb{S}_U$ ) katman olarak sınıflandırılır

olmadığı kararı verildikten sonra epitel dokuya ait olan  $\mathbf{P}_j$  blok görüntüsünün merkez noktasından bazal membrana ve epitel yüzeye olan uzaklıkları hesaplanır. Eğer  $\mathbf{P}_j$  dokusuna yakın bir papilla kordinatı varsa bazal membran olarak o sınır ele alınır.

$d_{bs}$  ve  $d_{up}$ ,  $\mathbf{P}_j$ 'nin merkez kordinatının bazal ve epitel yüzeye olan en kısa mesafeleri olsun. Buradan konum oranı  $r = \frac{d_{bs}}{d_{bs} + d_{up}}$  olarak hesaplanmaktadır.  $\mathbf{P}_j$  ızgara görüntüsünün konumu  $r$ 'nin  $[0, 1/4]$ ,  $[1/4, 3/4]$ ,  $(3/4, 1]$  durumlarına göre *basal*, *orta* or *üst* katmana ait olarak tanımlanır. Şekil 4.6'de açıklanan  $\mathbf{P}_j$  parçalarının katmanları belirlendikten sonra  $\mathcal{T}_i$  kümeleri altında toplanır. YHA'nın *Bölgesel Histogram Çıkartımı* kısmında ise  $\mathcal{T}_i$  kümeleri üzerinden bölgesel histogramlar ( $h_i$ )  $[0-255]$  değerleri arasında çıkarılmaktadırlar.  $h_i$ 'nin OKF'si olarak yorumlanabilmesi için  $\mathbf{p}_i = \frac{h_i}{\sum h_i}$  dönüşümü yapılmıştır. Sonuç olarak üç katmana ait elde edilen  $\mathbf{p}_1, \mathbf{p}_2, \mathbf{p}_3$  OKF'ler dokuyu istatistiksel olarak betimleyen öznelilikler (YHÖ) olmuşlardır.

YHÖ'ler KEP dokuları için ayırt edici oldukları Şekil 4.7'da gösterilmektedir. Şekil 4.7(a) (b) (c) sırasıyla bazal, orta ve üst katmana ait farklı derecelerdeki dokulara ait histogram eğrileridir. Hücre yoğunluğu tüm dokularda bazal katmanda fazla olduğu için bu katmana ait  $[100 - 200]$  arasındaki değerler, diğer katmanlara nazaran daha yüksek çıkmıştır. Normal dokularda orta ve üst katmana ait hücre yoğunluğu azaldığından dolayı  $\mathbf{p}_2$  ve  $\mathbf{p}_3$ 'teki koyu bölgenin değeri azalmıştır. Diğer



**Şekil 4.7** CIN seviyelerine göre katmanlardan elde edilen histogram grafikleri

tarafından CIN-3 derecesi göz önüne alındığında hücreler dokunun neredeyse tamamını kapladığı için koyu tonlara ait piksel değerleri tüm okf'lerde etkindir.

**(ii) Çekirdek Yoğunluk Haritası Tabanlı Yardımcı Öznitelik:** YHÖ, histogram aracılığıyla katmana dair doku özellikleri (çekirdek, sitoplazma koilosit ve beyaz boşluklar) hakkında bilgi verir. Bu bilgi HPV virusu nedeniyle çoğalan hücrelerin katman içerisinde yer tuttuğu alan miktarını vermemektedir. Bu bilgi ise patoloji rutini açısından tanılarının koyulmasında başlıca önem arz etmektedir. Bundan dolayı ÇYH'nin çıkartılması, sınıflandırıcının katmanlar içindeki hücre yoğunluğu bilgisinin edinmesi açısından oldukça önemlidir. Bir ÇYH, gri-seviyedeki H bandından bölütlenerek elde edilen ve sadece hücre yapısının bulunduğu ikili görüntüdür. Ön işleme adımı bahsedilen SegNet, gürbüz hücre bölütleme için kullanılmıştır. Bu sayede  $\mathbf{I}_{bw} \in \{0, 1\}^{N_x \times M_y}$  ÇYH elde edilir.

Katman içi ve katmanlar arasındaki bağıl çekirdek yoğunluklarının bulunması için ÇYH kullanılır. Çekirdek yoğunluğu YHÖ'lerin çıkarılma esnasında Algoritma 7'de hesaplanmaktadır.  $\mathbf{P}_j$  ızgara görüntüleri  $\mathbf{I}_{gry}$  için kullanılırken aynı zamanda ÇYÖ'lerinde çıkartılması için  $\mathbf{I}_{bw}$  üzerinde gezdirilirler.  $i$ . katman için  $W_i$  beyaz ve  $T_i$  ise tüm alan değerleri olsun. Katman içi ve katmanlar arası bağıl ÇYÖ'ler sırasıyla  $\mathbf{x}, \mathbf{y} \in \mathbb{R}^{3 \times 1}$ ,

$$\mathbf{x} = \left[ \frac{W_1}{T_1}, \frac{W_2}{T_2}, \frac{W_3}{T_3} \right]^T, \mathbf{y} = \left[ \frac{W_1}{S_w}, \frac{W_2}{S_w}, \frac{W_3}{S_w} \right]^T, S_w = \sum_{i=1}^3 W_i \quad (4.2)$$

olarak tanımlanmıştır. (4.2)'de  $\mathbf{x}$ , katman-içi çekirdek yoğunluğunu ifade ederken  $\mathbf{y}$  ise, katmanlar arasındaki çekirdek dağılımının birbirlerine olan oranını vermektedir. Genel bir ÇYÖ olarak  $\mathbf{v}_y$  ise şu şekilde tanımlanır:  $\mathbf{v}_y = [\mathbf{x}^T, \mathbf{y}^T]^T$ .

#### 4.4.4 Kullback – Leibler Diverjansı Tabanlı Ağırlıklı $k$ -NN Algoritmasının Uyarlanması

YHÖ ile ÇYÖ birbirlerinden nitelik olarak bağımsız olduğundan dolayı iki dokunun benzerliğinin ölçütünde birlikte kullanılması için yeni bir metrik öne sürülmüştür. (4.1)'de bulunan uzaklık metriği olasılıksal tabanlıdır. Her katmandan elde edilen OKF'ler üzerinden uzaklık hesaplanmasının yanı sıra ÇYÖ'ler için  $\ell_1$  veya  $\ell_2$ -norm kullanılarak çekirdek yoğunluğu bakımından karşılaştırma da yapılır.

Herhangi bir KEP görüntüsünün eğitim setindeki örneklerle olan benzerliği Algoritma 8 tarafından bulunmaktadır.  $N_{trn}$  eğitim setindeki örnek sayısı olsun. Test için KEP üzerinden alınan OKF'ler, daha önceden eğitim seti üzerinden elde edilen

---

**Algoritma 8** *Test dokusunun veri seti ile YHÖ'ler üzerinden benzerlik hesabı*

---

Girdi  $\mathbf{Q} \in \mathbb{R}^{3 \times 256}$ ,  $\mathbf{P} \in \mathbb{R}^{3 \times 256 \times N_{trn}}$   
Çıktı  $\mathbf{d}_{kl} \in \mathbb{R}^{N_{trn} \times 1}$

```
1: Prosedür OLASILIKSAL BENZERLİK METRİĞİ
2:   for  $i = 1 : N_{trn}$  do
3:      $\phi \leftarrow 0$ 
4:     for  $j = 1 : 3$  do
5:        $q_j \leftarrow \mathbf{Q}(j, :)$ ,  $p_j \leftarrow \mathbf{P}(j, :, i)$ 
6:        $\phi = D_{KL}(q_j, p_j) + \phi$ 
7:      $\mathbf{d}_{kl}(i) = \phi / 3$ 
8:   return  $\mathbf{d}_{kl}$ 
```

---

$\mathbf{P} \in \mathbb{R}^{3 \times 256 \times N_{trn}}$  tensörü üzerinden benzerlik hesaplamasına tutulurlar. Bu sayede herbir eğitim örneği için  $\mathbf{d}_{kl} \in \mathbb{R}^{N_{trn} \times 1}$  olasılıksal mesafe vektörü elde edilir. Bu vektör Algoritma 6'te verilen  $wkNN$  algoritması ile KEP görüntülerinin sınıflandırmasında kullanılmaktadır.

İki görüntünün ÇYÖ vektörleri arasındaki benzerlik  $\ell_1$ -norm kullanılarak belirlenebilir.  $\mathbf{V}_{trn} \in \mathbb{R}^{6 \times N_{trn}}$  eğitim görüntülerinden elde edilmiş eğitim vektörlerinin tutulduğu matris olsun.  $\mathbf{v}_{tst} \in \mathbb{R}^{6 \times 1}$  test vektörünün  $\mathbf{V}$  matrisi sütunları ile  $\ell_1$ -norm uzaklığının tutulduğu mesafe vektörü  $\mathbf{d}_y \in \mathbb{R}^{N_{trn} \times 1}$  olarak ifade edilir.  $\mathbf{d}_{kl}$  ve  $\mathbf{d}_y$  vektörleri farklı türlerde olsa bile iki metrik vektörü aynı boyutlarda olduğu için toplanabilir hale gelmiştir. Toplanan bu vektör  $\mathbf{d}_{gl} \in \mathbb{R}^{N_{trn} \times 1}$  (4.3)'de hesaplanmaktadır.

$$\mathbf{d}_{gl} = \mathbf{d}_{kl} + \gamma \times \mathbf{d}_y \quad (4.3)$$

$\mathbf{d}_{gl}$ , hem parlaklık değerlerine ait istatistiksel bilgiyi hem de çekirdek yoğunluk bilgisini içermektedir. Farklı özniteliklere ait bilgilerin kullanılması algoritmayı daha etkili kılmaktadır. (4.3) ile hesaplanan metrik  $wkNN$  ile sınıflandırılarak dokunun derecesine karar verilir.  $\gamma$ ,  $\mathbf{d}_y$  için ölçekleme faktörüdür.

Önerilen yöntemin serviks dokuları için uygulanışı esnasında takip edilen gelişmelerin katkıları, *Algoritmaların Analizi ve Başarım Karşılaştırmaları* başlığı altında irdelenmiştir. Herbir aşama ayrı ayrı test edilerek yöntem içindeki ağırlıklı katkısı gözler önüne serilmiştir. Bu sayede yöntem içine eklenen adımların gerekliliği analiz edilmiştir.

## 4.5 Algoritmaların Analizi ve Başarım Karşılaştırmaları

YHÖ ve ÇYÖ'ler toplam 957 KEP görüntüsünden elde edilmişlerdir. Bu özneteliklerin karşılaştırılması mesafeye dayalı olduğu için sınıflandırıcı olarak  $k$ -NN ve  $wkNN$  kullanılmıştır. Bu kombinasyonlar arasında en iyi öznetelik ve sınıflandırıcı çiftinin belirlenmesi için parametre en iyilemesinin yapılması gerekmektedir. Bu amaç doğrultusunda *Önerilen Algoritmanın Sınanması* bölümünde öne sürülen yaklaşımların aşama aşama KEP'lerin derecelendirmesindeki başarımına olan etkileri ele alınmıştır. Bu kısımda veri setini en iyi sınıflandıran algoritmanın parametreleri kestirilmiştir. *DeneySEL Sonuçlar* kısmında ise önerilen algoritmanın morfometrik özelliklerle karşılaştırılması ve patolojlara göre analizlerinin yapılması ile ilgilidir.

Veri seti 471 normal, 240 CIN-1, 103 CIN-2 ve 132 CIN-3 KEP görüntüsü içerdiği için dengesiz bir veri setidir. Dengesiz veri setlerinin analizinde *Doğruluk* metriği tanıma sisteminin başarımı hakkında yeterli bir bilgi sağlamaz. Bu nedenle (4.4)'te tanımlanan *Doğruluk* metriğinin yanı sıra *F-Skor* (4.6) değerlerindeki hesaplanması gerekmektedir.

$$ACC = \frac{\sum_i K_{ii}}{\sum_{i,j} K_{ij}} \quad (4.4)$$

$K \in \mathbb{R}^{4 \times 4}$  sınıflar arası karışım matrisidir. Çok sınıflı kümelerde *Makro F-Skor* değerinin hesaplanması için ilk olarak precision  $i$ . sınıf için ( $PRE_i$ ) and recall ( $RCL_i$ ) değerlerinin hesaplanması gerekmektedir (4.5). Tüm sınıflar için *F-skor* hesaplandıktan sonra (4.6) üzerinden *Makro F-Skor* hesaplanır.

$$PRE_i = \frac{K_{ii}}{\sum_i K_{ij}} , \quad RCL_i = \frac{K_{ii}}{\sum_j K_{ij}} \quad (4.5)$$

$$F_i = \frac{2 \times PRE_i \times RCL_i}{PRE_i + RCL_i} , \quad Macro\ F1 = \frac{1}{N_c} \sum_{i=1}^{N_c} F_i \quad (4.6)$$

Algoritmaların test edilmesi sırasında veri seti, rastgele örneklem ile belli bir test oranına göre eğitim ve test alt gruplarına bölünmüştür. Tablo ve şekillerdeki tüm başarım değerlerinin, 300 rastgele denemenin ortalama ve standart sapması hesaplanarak elde edildirmiştir. Algoritmalar hem CIN hem de SIL tabanlı derecelendirme sistemleri tarafından değerlendirildiğinden, ilk olarak patolojların tanısını incelemek önemlidir. Tablo 4.1, kesin tanı ve patolojlar arasında CIN ve SIL temelli derecelendirmeler için yüzdelik olarak karışıklık matrislerini sunar. CIN tabanlı sistemde CIN-1 ve CIN-2'nin sınıflandırılması zor bir problemdir. SIL tabanlı sistem incelendiğinde, patolojların performansı, CIN-2 ve CIN-3'ü HSIL sınıfları olarak

**Tablo 4.1** Patologların CIN ve SIL tabanlı derecelendirme sistemlerine göre karışıklık matrisi: Gözlemciler arası değişkenlik kesin tanıya göre yüzde olarak gösterilmiştir

| CIN Derecelendirme |                    | Patolog 1 (ACC: 80,36%) |       |       |       | Patolog 2 (ACC: 86,83%) |       |       |       |
|--------------------|--------------------|-------------------------|-------|-------|-------|-------------------------|-------|-------|-------|
|                    |                    | Normal                  | CIN-1 | CIN-2 | CIN-3 | Normal                  | CIN-1 | CIN-2 | CIN-3 |
| Kesin Tanı         | Normal             | 80,04                   | 15,71 | 3,83  | 0,42  | 90,02                   | 9,34  | 0,64  | 0     |
|                    | CIN-1              | 5,42                    | 73,33 | 15,42 | 5,83  | 10,83                   | 80,42 | 6,67  | 2,08  |
|                    | CIN-2              | 0                       | 1,87  | 81,31 | 16,82 | 7,78                    | 15,89 | 73,83 | 2,81  |
|                    | CIN-3              | 0,72                    | 0     | 6,47  | 92,81 | 0,72                    | 0     | 2,16  | 97,12 |
|                    | SIL Derecelendirme | Patolog 1 (ACC: 83,18%) |       |       |       | Patolog 2 (ACC: 87,46%) |       |       |       |
|                    |                    | Normal                  | HSIL  | LSIL  |       | Normal                  | HSIL  | LSIL  |       |
|                    | Normal             | 80,04                   | 15,71 | 4,25  |       | 90,02                   | 9,34  | 0,64  |       |
|                    | HSIL               | 5,42                    | 73,33 | 21,25 |       | 10,83                   | 80,42 | 8,75  |       |
|                    | LSIL               | 0,41                    | 0,81  | 98,78 |       | 3,66                    | 6,91  | 89,43 |       |

içerdiğinden dolayı artmıştır.

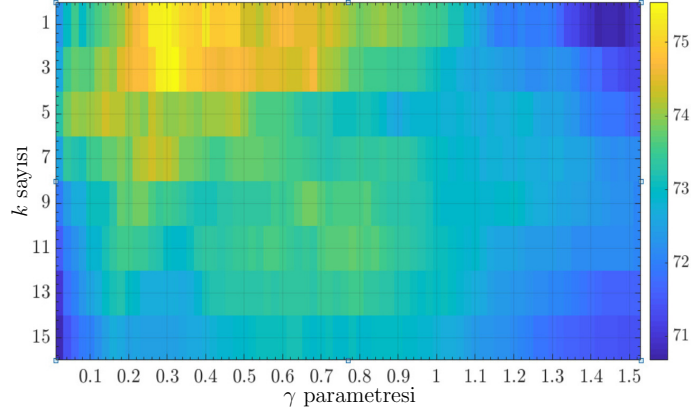
#### 4.5.1 Önerilen Algoritmaların Başarım Analizi

Bu bölümde, YHÖ ve ÇYÖ'nin katkılarını tartışılmış ve ayrıca  $k$ -NN ve  $wkNN$  sınıflandırıcılarının başarıma olan etkileri analiz edilmiştir. Bunun yanında, öznelilik vektörleri çıkartılırken ÇYÖ'ler için kullanılan ölçeklendirme parametresi ( $\gamma$ ) ve  $wkNN$  için en iyi ( $k$ ) değerlerinin tespiti de incelenmiştir. Araştırmalar doğrultusunda, YHÖ +  $\ell_1$  -norm ÇYÖ ve  $wkNN$  servikal prekanseröz lezyon sınıflandırması için başarılı ve faydalı bir kombinasyon olduğu görülmüştür.

Tablo 4.2 YHÖ'nün ve  $wkNN$  sınıflandırmaya olan katkıları ön plana çıkarılmıştır. YHÖ ve  $k$ -NN'nin doğruluğu %69,90'da kalırken, YHÖ +  $\ell_2$ -norm ÇYÖ ve  $wkNN$  % 75,21'lik bir doğruluk ve %71,46'lık F-skor değerine ulaşmıştır. Izgara tabanlı arama yapılarak  $k \in \{1, 3, 5, 7, 9, 11, 13\}$  ve  $\gamma \in \{0, 0,02, 0,04, \dots, 1, 5\}$  parametreleri üzerinden algoritmanın en iyilemesi yapılmıştır. Ayrıca  $\ell_1$  ve  $\ell_2$ -norm karşılaştırmaları doğrultusunda tanıma sistemi için uygun olan norm seçilmiştir. Veri seti %90 eğitim ve %10 test için ayrılarak her parametre kombinasyonu için 300 defa rastgele çalıştırılmış 300 denemenin ortalaması olarak bakıldığında norm:  $\ell_1$ ,  $\gamma = 0,3$ ,  $k = 3$  değerleri ile %75,55 $\pm$ 4,02 doğruluk ve %72,05 $\pm$ 4,88 makro F1-skor değeri ile Şekil 4.8'de en iyi

**Tablo 4.2** Önerilen algoritmaların  $\gamma = 0,5$   $k = 3$  parametreleri için başarım analizi

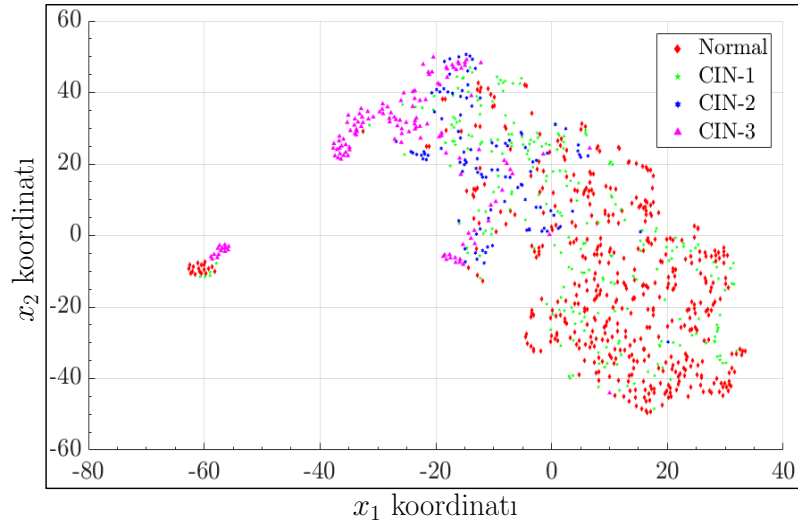
| Başarım Kıstasları:    | Doğruluk (%)     |                  | Makro F-Skor (%) |                  |
|------------------------|------------------|------------------|------------------|------------------|
| Sınıflandırıcılar:     | $k$ -NN          | $wkNN$           | $k$ -NN          | $wkNN$           |
| YHÖ                    | 69,90 $\pm$ 4,59 | 72,64 $\pm$ 4,69 | 65,06 $\pm$ 5,55 | 68,27 $\pm$ 5,44 |
| YHÖ + ÇYÖ ( $\ell_1$ ) | 72,05 $\pm$ 4,38 | 74,82 $\pm$ 4,39 | 68,86 $\pm$ 5,20 | 71,51 $\pm$ 5,03 |
| YHÖ + ÇYÖ ( $\ell_2$ ) | 71,76 $\pm$ 4,14 | 75,21 $\pm$ 4,15 | 68,68 $\pm$ 4,96 | 71,46 $\pm$ 4,85 |



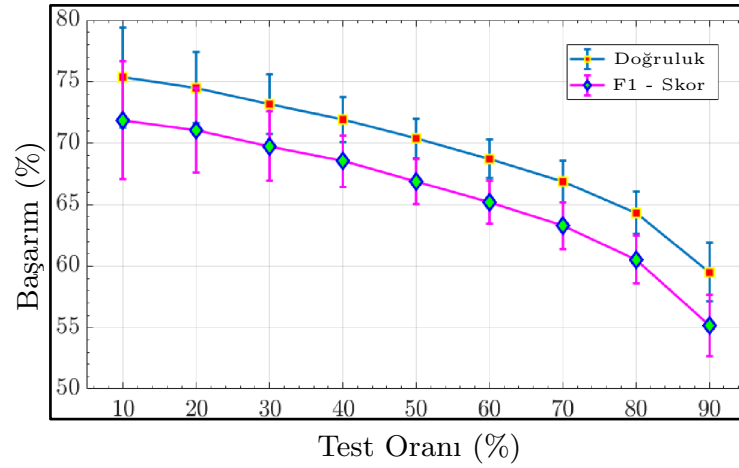
**Şekil 4.8** Parametre uzayında ızgara tabanlı en iyileme sonuçları:  $k = \{1, 3, 5, 7, 9, 11, 13, 15\}$  ve  $\gamma = \{0, 0, 02, \dots, 1, 50\}$  değerleri arasında en iyi doğruluk başarımı  $75,55 \pm 4,02$ , norm:  $\ell_1$   $\gamma$ -Parameter: 0,3 ve  $k$ -NN: 3 konfigürasyonu ile elde edilmiştir

sonuç elde edilmiştir. Bu parametrelerin kullanılması ile elde edilen sınıflandırıcıya *İstatistiksel Özniteliklerin Füzyonu ile Sınıflandırma* (İÖFS) denilmiştir. Yaklaşık % 75'lik bir doğruluk oranı patologların %80,36 ve %86.83 doğrulukla çalıştığı zorlu bir veri seti için oldukça iyi bir değerdir.

Bu çalışmada ayrıca 957 KEP görüntüsü için öne sürülen KL diverjans +  $\ell_1$ -norm metrik uzayındaki veri dağılımı t-SNE algoritması [141] kullanılarak görsel hale getirilmiştir. Şekil 4.9'da görüleceği gibi CIN-1 sınıfına ait dokular normal dokular ile iç içe geçmiştir. Ayrıca CIN-3 sınıfına ait örneklerin uzayın farklı bölgelerinde öbeklendiğide gözlemlenmektedir. Bu kadar zor bir veri seti için oluşturulan metrik uzayında yakınlık temelli  $wkNN$  algoritmasının elde ettiği başarı t-SNE görselinden daha iyi anlaşılmaktadır.



**Şekil 4.9** Veri setinin önerilen metrik uzayından t-SNE algoritması ile iki boyutlu düzleme haritalandırılarak görselleştirilmesi



**Şekil 4.10** İÖFY algoritmasının farklı test oranları altındaki başarımların analizi

Bir diğer önemli analiz ise algoritmanın değişen eğitim ve test seti büyüklüğüne göre başarımın ne derecede değiştiğinin ölçülmesidir. Bu nedenle, İÖFS'nin doğruluğunu ve F-skorunu ölçmek için test oranı %10'dan %90'a kadar değiştirildi. Şekil 4.10'da görüldüğü gibi, İÖFS'nin başarımı dramatik olarak düşmemektedir. %90 test oranına rağmen, %60 civarında bir doğruluk elde edebilmiştir.

#### 4.5.2 İÖFY Algoritmasının Morfometrik Algoritmalar ile Kıyaslanması

H&E ile boyanmış preparatların değerlendirilmesinde morfometrik öznitelikler dokuların sınıflandırılmasında çoğunlukla kullanılmaktadır [121, 142–145]. [136] 'de hücre morfolojisi olarak, ortalama hücre alanı, çevresi, eksantrikliği, bazal membranın açısı, vb. özellikler öznitelik vektöründe kullanılmıştır. Bu özelliklere ek olarak, bu değerlerin bazal, orta ve üst katmanlar arasındaki oranları da ayrıca öznitelik vektörüne dahil edilmiştir. Elde edilen özniteliklerin tümü bir KEP için 54 uzunluğunda *Morfometrik Öznitelik Vektörü (MÖV)* oluşturmaktadır.

MÖV'nin sınıflandırılması için bu çalışmada *wkNN* ve çekirdek SVM (RBF ve polinom) kullanıldı. Çekirdek parametresinde SVM için en iyi değerlerin bulunmasında  $\sigma \in \{1, 3, 5, 10, 15, 20, 50\}$  arasında arama yapılmıştır.  $C = 1$  değeri olarak seçildi. Arama sonucunda  $\sigma = 10$  RBF ve polinom çekirdeği en iyi başarımları vermiştir. Tablo 4.3'de MÖV ile İÖFS sonuçları karşılaştırılmıştır.

Sınıflandırıcıların analizinde toplam ve sınıflar-arası başarımlar yüzdelik olarak incelenmektedir. Sınıflandırıcılar veri setinin 300 kere rastgele ve eğitim ve test için %90'a %10 ayırma oranlarına ayrıştırılması ile elde edilen kümeler üzerinden test edilmiştir. Adil bir karşılaştırma olması için her bir sınıflandırıcının testi aynı ayrıştırılmış veri kümeleri üzerinden yapılmıştır. 300 denemenin başarımlar doğruluk ve F-skor değerlerinin ortalaması ve standart sapması Tablo 4.3'de verilmiştir. İÖFY

**Tablo 4.3** İÖFY'nin farklı sınıflandırıcılar ile sınıflandırılan morfometrik öznitelikler ile karşılaştırılması

| Sınıflandırıcılar                 | Doğruluk (%) |       |       |       |              | Makro F-Skor (%) |       |       |       |              |
|-----------------------------------|--------------|-------|-------|-------|--------------|------------------|-------|-------|-------|--------------|
|                                   | Normal       | CIN-1 | CIN-2 | CIN-3 | Total        | Normal           | CIN-1 | CIN-2 | CIN-3 | Total        |
| MÖV- <i>wkNN</i>                  | 73,94        | 31,86 | 26,94 | 62,33 | 56,34 ± 4,21 | 70,54            | 33,86 | 27,61 | 63,04 | 48,32 ± 6,83 |
| MÖV-SVM <sub>poly</sub>           | 86,31        | 30,10 | 45,70 | 71,79 | 65,49 ± 3,98 | 78,47            | 37,14 | 43,15 | 74,52 | 58,32 ± 5,42 |
| MÖV-SVM <sub>r<sub>bf</sub></sub> | 87,83        | 27,22 | 46,70 | 69,07 | 64,84 ± 4,14 | 77,99            | 34,64 | 44,05 | 72,59 | 57,32 ± 5,44 |
| İÖFY                              | 83,70        | 58,13 | 62,18 | 85,07 | 75,04 ± 3,94 | 84,07            | 58,86 | 58,07 | 85,81 | 71,70 ± 4,59 |

%75,04 ± 3,94 'lik tanıma performansı ile zorlu bir veri seti için kayda değer bir başarı sağladığı açıktır. Sınıf bazında başarımlar ele alındığında MÖV-SVM<sub>poly</sub> ve MÖV-SVM<sub>r<sub>bf</sub></sub>'nin Normal sınıflı dokuların sınıflandırılmasında daha başarılı olduğu görülmektedir. Bunun nedeni SVM'in eğitiminde 471 tane Normal dokunun diğer lezyonlardan oldukça fazla olması ve modeli domine etmesidir. Fakat diğer derecelerin tespitinde SVM tabanlı sınıflandırıcılar yetersiz kalmışlardır. Diğer taraftan İÖFY, önerdiği metrik uzayının ayırma kabiliyeti ve *wkNN* sınıflandırıcısının yerel eğitim noktaları üzerinden sınıflandırma yapması nedeni ile diğer sınıflarının tespitinde üstün bir başarı elde etmiştir. Tablo 4.4, İÖFY ve SVM-tabanlı sınıflandırıcılar arasında en iyi sonucu veren MÖV-SVM<sub>poly</sub> algoritmalarının karşılaştırılması ile oluşmuştur. Normal dokular için MÖV-SVM<sub>poly</sub>, %86,31 doğruluk bulurken İÖFY %83,70'lik bir başarı elde etmiştir. Diğer taraftan MÖV-SVM<sub>poly</sub>, CIN-1 dokusuna ait %56,74'lük bölümü Normal olarak sınıflandırmıştır.

Bir başka önemli karşılaştırma ise SIL derecelendirmesi kullanılarak yapılmıştır. SIL derecelendirme sisteminde dokular Normal, Düşük-SIL (Low-SIL) ve Yüksek-SIL (High-SIL) olarak sınıflandırılmaktadır. CIN-2 ve CIN-3 bu derecelendirme sisteminde Yüksek-SIL olarak tek sınıfta ele alınmıştır. CIN için oluşturulan 300 eğitim ve test kümelerinin haricinde SIL için yeni kümeler oluşturulmuş ve elde edilen sonuçlar Tablo 4.4'de verilmiştir. SIL sisteminde, CIN-2 CIN-3 arasında meydana gelen

**Tablo 4.4** İÖFY ve morfometrik öznitelikler arasında en iyi sonucu veren MÖV-SVM<sub>poly</sub> algoritmalarının karışım matrisleri

| CIN Derecelendirme | İÖFY (ACC: 75,04%) |       |       |                                       |       | MÖV-SVM <sub>poly</sub> (ACC: 65,49%) |       |       |       |
|--------------------|--------------------|-------|-------|---------------------------------------|-------|---------------------------------------|-------|-------|-------|
|                    | Normal             | CIN-1 | CIN-2 | CIN-3                                 |       | Normal                                | CIN-1 | CIN-2 | CIN-3 |
| Kesin Tanı         | Normal             | 83,70 | 13,20 | 2,26                                  | 0,85  | 86,32                                 | 7,46  | 5,16  | 1,06  |
|                    | CIN-1              | 26,69 | 58,13 | 13,31                                 | 1,88  | 56,74                                 | 30,10 | 9,92  | 3,24  |
|                    | CIN-2              | 4,82  | 24,03 | 62,18                                 | 8,97  | 15,76                                 | 24,06 | 45,70 | 14,48 |
|                    | CIN-3              | 1,83  | 3,33  | 9,75                                  | 85,07 | 3,21                                  | 8,62  | 16,38 | 71,79 |
| SIL Derecelendirme | İÖFY (ACC: 78,20%) |       |       | MÖV-SVM <sub>poly</sub> (ACC: 70,38%) |       |                                       |       |       |       |
|                    | Normal             | HSIL  | LSIL  | Normal                                | HSIL  | LSIL                                  |       |       |       |
| Kesin Tanı         | Normal             | 84,04 | 13,04 | 2,92                                  | 86,49 | 6,50                                  | 7,01  |       |       |
|                    | HSIL               | 26,36 | 59,14 | 14,50                                 | 56,99 | 25,54                                 | 17,47 |       |       |
|                    | LSIL               | 2,80  | 11,86 | 85,52                                 | 7,01  | 9,87                                  | 83,12 |       |       |

yanlış sınıflandırmalar, iki sınıfında HSIL olarak sınıflandırılmaları nedeniyle ortadan kalkmıştır. Bu nedenle algoritmaların performansı, MVFP ve İÖFY için sırasıyla %70,38 ve %78,20'ye yükseldi. Tablo 4.4'te gösterilen karışım matrisi incelendiğinde, MÖV-SVM<sub>poly</sub>, Düşük-SIL sınıfını Normal dokudan ayırmakta hala zorluk çekmektedir. Ayrıca, İÖFY'nin en çok karıştırdığı sınıflar %26,36'lık bir oran ile Düşük-SIL ve Normal arasındadır.

## 4.6 Sonuç

Rahim ağzı kanserinin öncü lezyonlarının doğru sınıflandırılması gözlemci içi ve gözlemciler arası değişkenliğin en aza indirilmesi hastaların doğru yönlendirilmesi için hayati öneme sahiptir. Bu alanda geliştirilen yöntemler, problemin doğasına (girintili epitel yüzey, papilla bölgeleri) ve patolojlara kabul edilebilir süre içinde sonuç verebilecek şekilde uygun olmalıdırlar. Yaygın olarak kullanılan morfometrik özellikler hücreleri tekil olarak ele almaları gereksiniminden dolayı çakışık hücre probleminin çözümü ve hücre morfolojisinin çıkartılmasında ciddi bir zaman tüketimine ihtiyaç duyarlar. Ayrıca, yüksek çözünürlük nedeniyle  $\times 20$  ve  $\times 40$  optik yakınlaştırmalı görüntülerin kullanılması hesaplama yükü daha da arttırmaktadır.

Tezin bu bölümünde, özgün olarak önerilen YHÖ ve ÇYÖ'lerin istatistiksel öznelilik olarak kullanılması ile servikal öncü lezyonları etkin bir şekilde sınıflandırabileceği gösterilmiştir. Bu öznelilikler, bazal, orta ve üst katmanların manifolduna uyumlu oldukları için lezyonların girift şekillerine karşı gürbüzdürler. Ayrıca çalışmada önerilen yöntemler hücreleri bir popülasyon olarak ele aldıkları için çakışık hücre problemi ile ilgilenmek zorunda değillerdir. Ayrıca, algoritmaya papilla sınırları dahil edilerek gerçekçi bir yöntem elde edilmiştir.

## Yüksek Çözünürlüğe Sahip Görüntülerde Çakışık Nesnelerin Bölütlenmesi

---

Geniş-ölçek veriler sadece işlenecek verinin çokluğu anlamında değil, aynı zamanda içerdiği bilginin yoğunluğu kastedilerek de kullanılmaktadır. Sayısal patoloji gibi bazı uygulamalarda sadece bir tane görüntünün GB boyutlarındaki piksellerden oluştuğu bilinmektedir. Özellikle hiperspektral gibi bant içeren görüntülerin işlenmesi geniş-ölçek veri analizi kapsamına girer. Bu tip görüntülerde bölütleme işlemi yüksek işlem yükü gereksinimiyle zorlu bir problem olarak karşımıza çıkar. Özellikle bölütlenmesi istenen nesneler üst üste çakışık bir durumda ise problem daha da karmaşık hale gelir.

Çakışık nesnelerin bölütlenmesine yönelik literatürde oldukça başarılı çalışan birçok algoritma mevcuttur. Fakat önerilen algoritmalar genellikle düşük çözünürlüğe sahip görüntüler göz önüne alınarak geliştirildiklerinden, zaman tüketimi hala önemli bir problem olarak durmaktadır. Tezin bu bölümünde, nesnelerin sadece seyreltilmiş sınır pikselleri üzerinden *Yönlü k-Ortalama Algoritması* (YKA) ile başarılı bir şekilde bölütlenebileceği gösterilmiştir. Önerilen algoritma pratik bir saha olarak biyomedikal görüntülerde çakışık hücrelerin ayrılmasında kullanılmıştır. Bölüm 4'te, serviks dokularının sınıflandırılması problemine yönelik literatürde var olan yöntemlerin büyük bölümünün hücre morfolojileri tabanlı olduğu gösterilmiştir. Bu doğrultuda, Bölüm 5 çerçevesinde öne sürülen YKA'nın hücre morfolojisi üzerinden öznitelik çıkartma işlemini hızlandıracak şekilde uyarlanmıştır.

Çakışık nesne problemi, görüntülerde üst üste binmiş nesnelerin sınırlarının ayrıştırılması ve bunların bölütlenmesi olarak tanımlanabilir. Görüntü işleme alanında kromozomların ayrıştırılmasından [146], kristal yapılarının bölütlenmesine [147], hücrelerin analizinden [148, 149], LIDAR görüntülerinin işlenmesine [150] birçok farklı sahada karşılaşılmaktadır. Literatürde varolan algoritmaların düşük çözünürlüğe sahip görüntülere yönelik geliştirildiği saptanmıştır. Bu doğrultuda YKA algoritması, problemi seyreltilmiş sınır pikselleri üzerinden tanımlanarak daha hızlı

bir çözüm önermektedir.

## 5.1 Literatür Özeti

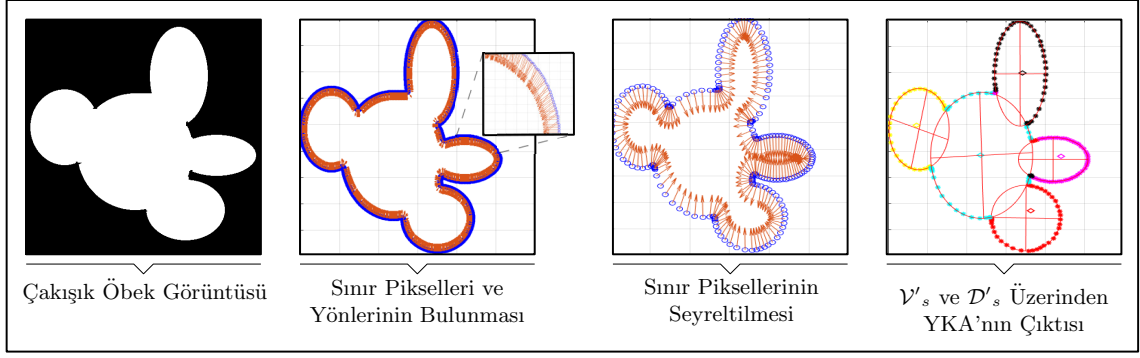
Agam ve ark. [146]'de, kromozomların sınıflandırılmasını arttıracak karyotipler üzerinden çakışık olanların ayrıştırılmasına yönelik bir çalışma yapmışlardır. İlk olarak metafaz görüntüleri iyileştirildikten sonra ilgi noktaları (iç bükey, dış bükey) belirlenmiştir. Ardından kromozom konturları üzerinden hipotezler geliştirilerek en uygun hipotezin bulunmasının ardından bölütleme yapılmıştır. Abhinav ve ark. ise [147]'teki çalışmada çakışık kristal görüntülerin ayrıştırılması için kontur tabanlı bir yöntem önermişlerdir. Bu yöntem ilk olarak kristal parçalarının sınır sınıflandırılmasını yapmakta ardından tespit edilen sınırlar üzerinden *şekil uydurma algoritması* ile kristalin tam şeklini alacak olan eğri tespit edilmektedir. Chen ve ark. [148]'teki çalışmada, histoloji görüntüleri üzerinden hücre sınırlarının bölütlenmesine yönelik çok-görevli (multi-task) bir yapıya sahip tüm-bağlı (fully-connected) ağ mimarisinde olan bir *Derin Kontur-Farkeden Ağ* önerisinde bulunmuşlardır. Ağın son katmanında kayıp fonksiyonu olarak hem sınır piksellerini hem de örüntüsel bağlamda ilişkileri içeren hibrit bir fonksiyon kullanmışlardır. Zafari ve ark. ([149]) ise gölge (silhouette) görüntülerinde çakışık nesnelerin ayrıştırılması için üç aşamalı bir yöntem önermişlerdir. İlk aşamada çakışık nesnelerin merkez noktaları *sınır erozyonu (ikili bir işlem)* ve *hızlı radyal simetri dönüşümü* ile tespit edilmekte, sonrasında bu merkezler ile ilişkili olan sınır pikselleri bulunmaktadır. Tespit edilen ve sınıflandırılan bu pikseller üzerinden elips uydurma algoritması ile öbekler ayrıştırılmıştır. Teo ve ark. [150], üç boyutlu LIDAR görüntülerinde bitişik nesnelerin (yollardaki ağaçlar trafik lambaları vb.) ayrıştırılması için renk kümelemesi tabanlı bir algoritma kullanmışlardır. Bu sayede algoritma, yol boyunca iç içe geçmiş ağaçları ve trafik işaretlerini ayrıştırabilmektedir. Çakışık nesnelerin ayrıştırılması, özellikle biyomedikal alanındaki uygulamalarda kritik bir öneme sahiptir. Bu probleme ymnelik üretilen algoritmalar, dokular içinde hücrelerin sayımının yapılabilmesini, fiziksel özelliklerinin elde edilmesini ve özneteliklerinin çıkartılması gibi hücrelerin tekil olarak ele alınmasını gerektiren bir çok alanda yaygın olarak kullanılırlar. Lu ve ark. [151]'deki çalışmada çoklu çoklu seviye kümelerini (multiple level set) kullanarak 900 tane sentetik üretilen serviks hücrelerine uygulamışlardır. Öbek içinde bulunan her hücre, farklı bir seviye kümesi fonksiyonu ile modellenmiş ve bu modeller hücre-içi (intracell) ve hücreler-arası (intercell) kısıtlar ile birleşik (joint) olarak en iyilemesi yapılmıştır. Böylece çakışık hücreleri birbirlerinden ayırırken hücre çekirdeği ve stoplazmayı da bölütleyebilmektedir. Çalışmada raporlandığına göre algoritma, 10 hücreye kadar olan öbekleri etkili bir şekilde ayırabilmektedir. Qi ve ark. ise problem için itici (repulsive) seviye kümesi algoritmasını geliştirmişlerdir [152]. Önerdikleri

yöntem iki aşamadan oluşmaktadır. İlk aşamada hücre merkezlerinin tespiti için kayan-ortalama (mean-shift) algoritması kullanılmıştır. Yazarların bu alandaki katkısı, kayan-ortalama algoritmasını Grafik İşlem Birimi (Graphical Process Unit (GPU)) yardımı ile paralelleştirerek tohum (seed) yerlerinin tespiti işlemini hızlandırmaları olmuştur. Aday hücre merkezleri tespit edildikten sonra ikinci aşamada itici seviye algoritması hücre sınırlarını interaktif bir şekilde tespit eder. Geliştirilen algoritma, 2200'den fazla kan hücresi içeren 234 görüntü üzerinden test edilmiştir. Tareef ve ark. [153] çalışmasında, çakışık hücre problemi için üç geçişli su kuyusu (watershed) algoritmasını ileri sürmüşlerdir. İlk geçişte, ön-işlemeden geçmiş bir görüntünün gradyan tabanlı kenar haritası üzerinde çekirdek yerleri tespit edilir. İkinci geçişte ise hücre şekillerine ve yerlerine adapte olmuş bir su kuyusu algoritması tekil, dokunan ve çakışık olan hücreleri bölütler. Son aşamada karşılıklı (mutual) iteratif bir su kuyusu algoritması çakışık olan hücreleri birbirinden ayırır. Probleme dönük bir başka yöntem sistematığı ise hücre öbeklerinin Gauss Karışım Modeli (Gaussian Mixture Model (GMM)) üzerinden eliptik olarak tanımlanmasıdır. Jung ve ark. [154] çalışmasında, su kuyusu algoritmasının bir aşaması olan mesafe dönüşümünü (distance transform) kullanarak hücre öbeklerinin topolojik bir haritasını çıkartmışlardır. Ardından mesafe değerleri ile orantılı olarak piksel çoğullama işlemi uygulamışlar ve bu sayede hücrelere Gauss dağılım karakteristiği kazandırmışlardır. Türetilen dağılımlar üzerinden GMM algoritması öbek sınırını hücreleri baz alarak sınıflandırmıştır. Sınır pikselleri kullanılarak doğrudan elips uydurma yöntemi [155] ile hücreler eliptik olarak bölütlenmiştir. Benzer bir yöntem [156] çalışmasında da kullanılmıştır.

## 5.2 Piksel Seyreltmeye Dayalı Yönlü k-Ortalama Algoritması

YKA kontur tabanlı bir algoritma olup, arkaplandan bölütlenmiş nesnelerin ikili görüntüleri ile çalışmaktadır. Algoritma temel olarak iki aşamadan oluşmaktadır. İlk aşama, ikili görüntüler üzerinden sınır piksellerinin bulunması ve bu piksellerin gradyan vektörlerinin tespitidir. Ayrıca bu aşamada uyarlamalı piksel seyreltme işlemi uygulanarak hesaplama yükü önemli ölçüde azaltılır. İkinci aşamada ise sınır pikselleri üzerinden seçilen ilk merkez (initialisation) noktaları ile önerilen YKA algoritması çalışmaya başlar. Sınıflandırılan sınır piksellerine elips şekilleri uydurularak nesneler tespit edilir. YKA algoritmasının işleyişi Şekil 5.1'te verilmiştir.

**(i) Sınır Piksellerinin ve Bu Piksellere Ait Gradyanların Bulunması:** Sınır piksellerine dayalı olan YKA, bölütleme sonrasında elde edilen ikili görüntü ( $I_b(x, y)$ ) üzerinden çalışmaktadır. Bu görüntü kullanılarak sınır pikselleri, herhangi bir kenar bulma operatörü (Sobel, Canny, LoG, vb.) yardımı ile kolaylıkla elde edilebilir. Bu çalışmada hesaplama hızı nedeniyle Sobel operatörü tercih edilmiştir. Sadece sınır



**Şekil 5.1** YKA'nın işleyiş sıralaması

piksel koordinatlarını içeren  $I_e(x, y)$  ikili görüntüsü, Sobel operatorlerinden gelen  $x$  ve  $y$  yönlerindeki gradyan görüntüleri olan  $G_x$  ve  $G_y$  üzerinden  $I_e(x, y) = (|G_x| + |G_y|) > 0$  ile bulunabilir.  $I_b(x, y)$  ikili resim olduğundan dolayı kenar bulma operatörü bölütlenmiş objenin iç bölgelerini sıfırlamaktadır. Bu nedenle sınır piksellerinin belirlenmesinde eşik değeri 0 olarak kullanılmıştır. Kenar bulma işleminin ardından bir piksel kalınlığında  $I_e(x, y)$  görüntüsü elde edilir. Sınır piksellerine ait yön vektörü ise,  $-\frac{\nabla I_b(x, y)}{|\nabla I_b(x, y)|}$  olarak tanımlanabilir. Negatif gradyan kullanılmasının sebebi gradyan vektörünün yönünü hücre içine doğru olmasını sağlamaktır. Gradyan vektörlerinin sadece sınır değerlerinde oluşacağı açıktır. Bu ifadeler üzerinden, sınır piksel koordinat bilgilerinin saklandığı  $\mathcal{V}$  kümesi (5.1)

$$\mathcal{V} = \left\{ \mathbf{v}_i \in \mathbb{R}^{2 \times 1} \mid \mathbf{v}_i(1) = x_i, \mathbf{v}_i(2) = y_i, I_e(x_i, y_i) = 1 \right\} \quad (5.1)$$

ve herbir koordinata ait yön vektörlerinin saklandığı  $\mathcal{D}$  kümesi (5.2)

$$\mathcal{D} = \left\{ \mathbf{d}_i \in \mathbb{R}^{2 \times 1} \mid \mathbf{d}_i = -\frac{\nabla I_b(x_i, y_i)}{|\nabla I_b(x_i, y_i)|}, I_e(x_i, y_i) = 1 \right\} \quad (5.2)$$

olarak tanımlanırlar. Birçok algoritma sadece bu kümelerdeki bilgileri kullanılarak çakışık nesnelerin bölütlemesini yapmaktadırlar. Sınır pikselleri nesnenin yüzeysel yapısını (iç-bükey ve dış-bükey kısımlarını) yeterince tanımlar. Nesne içi pikseller, öbek içinde kaç nesne olduğu ve bunların merkezlerinin kestirimi için kullanılırlar. Bunun yanında nesne içindeki piksellerin kullanılmasına dayalı bölütleme algoritmaları da mevcuttur. Fakat yüksek çözünürlüğe sahip görüntülerin ele alındığı durumlarda bu algoritmalar, işlem maliyetini bir hayli yükseltmektedirler. Diğer yandan önerilen YKA ise sınır piksellerini *Uyarlamalı Piksel Seyreltme* (UPS) algoritması (Algoritma 9) ile seyrelttikten sonra bölütleme işlemine başlamaktadır. Bu sayde işlemleri hem sınır pikselleri üzerinden yapması hem de UPS kullanması

---

**Algoritma 9** Uyarlamalı Piksel Seyreltme

---

**Girdi:**  $\mathcal{V}_s, \mathcal{D}_s, \theta_{\text{eşik}}, d_{\text{eşik}}$

**Çıktı**  $\mathcal{V}'_s, \mathcal{D}'_s$

```
1: Prosedür SEYRELTME
2:    $\mathbf{d}_t \leftarrow \mathcal{D}_s \{ \text{Uni}[1, N_s, 1] \},$        $\triangleright \mathcal{D}_s$ 'den  $[1, N_s]$  arasında rasge bir nokta atanır.
3:    $t \leftarrow 0$                                       $\triangleright t$  indisine 0 atanır.
4:    $\mathcal{V}'_s, \mathcal{D}'_s \leftarrow \emptyset$ 

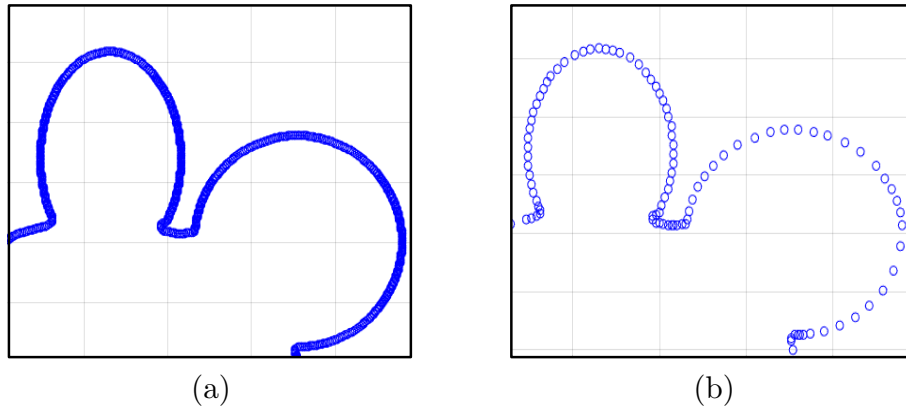
5:   Başla:
6:   while  $t \leq N_s$  do
7:      $i \leftarrow 0$ 
8:     while  $(f(\mathbf{d}_t, \mathbf{d}_i) \leq \theta_{\text{eşik}}) \ \& \ (\|\mathbf{v}_t - \mathbf{v}_i\|_2 \leq d_{\text{eşik}})$  do
9:        $i \leftarrow i + 1$ 
10:    bitti

11:     $\mathcal{V}'_s \leftarrow \mathcal{V}'_s \cup \mathbf{v}_i$  ve  $\mathcal{D}'_s \leftarrow \mathcal{D}'_s \cup \mathbf{d}_i$ 
12:     $t \leftarrow i$ 
13:  bitti
```

---

sayesinde hızlı bir bölütleme imkanı sağlar. Algoritma 9 girdi olarak sınır piksel vektörlerinin saklandığı  $\mathcal{V}_s$  ile ilgili vektörlere ait gradyanların tutulduğu  $\mathcal{D}_s$  kümelerini almaktadır. *Seyreltme* işlemi başladığında ilk olarak  $N_s$  tane sınır pikseli içinden rasgele bir noktaya ait *taban*  $\mathbf{d}_t$  seçilir. Ardından  $\mathbf{d}_t$ , kendisinden sonra gelen vektörler ile *açı* ve *mesafe* kıstaslarına göre karşılaştırılır. İki vektör arasındaki açı  $f(\cdot)$  ile hesaplanırken, sınır pikselleri arasındaki mesafe için  $\ell_2$ -norm kullanılmıştır. Bu iki parametre  $\theta_{\text{eşik}}$  ve  $d_{\text{eşik}}$  değerlerinden aşağıda ise komşu nokta atlanır, değilse o nokta ve gradyanı  $\mathcal{V}'$  ve  $\mathcal{D}'$  kümelerine dahil edilir. Yeni taban gradyanı olarak yeni bulunan nokta seçilerek işlem devam eder.

UPS algoritmasının çalışmasına bir örnek Şekil 5.2'te verilmiştir. Şekil 5.2(a), ikili görüntüye ait tüm sınırları içermektedir. Bu noktalar üzerinden UPS  $\theta_{\text{eşik}} = 10^\circ$  ve



**Şekil 5.2** Sınır piksellerinin uyarlamalı olarak seyreltilmesi.  $\theta_{\text{eşik}} = 10^\circ$  ve  $d_{\text{eşik}} = 5$  olarak seçilmiştir

$d_{\text{eşik}} = 5$  için çalıştırıldığında Şekil 5.2(b) elde edilmektedir. UPS, yassı yüzeyler (flat) üzerinde bulunan noktaları çok seyreltirken açısal değişimin fazla olduğu yerleri daha az seyreltmektedir. İndirgeme koşulunda sadece açısal değişime bakılsaydı yassılığın yüksek olduğu yerlerde çok fazla piksel seyreltmesi olacak ve gerekli olan bilginin kaybına sebep olacaktı. Bu nedenle, koşul olarak ayrıca sınır pikselleri arasındaki mesafeye bakılması da eklenmiştir.

Sınır pikselleri seyreltilip  $\mathcal{V}'$  ve  $\mathcal{D}'$  kümeleri bulunduktan sonra YKA çalışmaya başlar. Bu sayede YKA, seçilen eşik değerlerine göre  $[2, 10]$  kata kadar daha az nokta ile çalışır.

**(ii) Yönü  $k$ -Ortalama Algoritması ile Bölütleme:**  $k$ -Ortalama algoritması makine öğrenmesinin en köklü ve en çok kullanılan eğitici-siz bölütleme algoritmalarından birisidir. Ancak bölütleme alanında birçok başarılı çalışmaları olan  $k$ -Ortalama algoritmasının yetersiz olduğu problemler de mevcuttur. Bu eksikliklerinin giderilmesi için  $k$ -Ortalama algoritmasının birçok türevleri geliştirilmiştir [27–30, 157]. Örnek olarak  $k$ -Ortalama algoritması çözüm iterasyonu esnasında örneklerin sınıflara atamasında keskin eşikleme (hard thresholding) uygulamaktadır. Bu nedenle örnekler sadece bir kümeye dahil olabilirler. Buna karşın *Bulanık  $c$ -Ortalama* (Fuzzy  $c$ -Means) algoritmasında örneklerin ataması üyelik fonksiyonu kullanılarak gerçekleştirilir. Örnek en yüksek değeri hangi fonksiyon altında veriyorsa o sınıfa dahil edilir [157].  $k$ -Ortalama algoritması ile ilgili başka bir problem ise veriyi kümelere kaç sınıfa bölüneceğinin belirli olmasından kaynaklanmaktadır. Bu doğrultuda verinin kümeleme aşamalarında Akaike Bilgi Kriteri (Akaike Information Criterion) ve Bayes Bilgi Kriteri (Bayesian Information Criterion) ölçümlerini göz önüne alarak bölütleme yapan X-Ortalama algoritması geliştirilmiştir [29].  $k$ -Ortalama algoritması ilk değer atamasına karşı farklı sonuçlar üretebilmektedir. Buna yönelik  $k$ -Ortalama++ ve bu algoritmayı hızlandırmak için Hiyerarşik  $k$ -Ortalama algoritmaları önerilmiştir [28]. Diğer taraftan  $k$ -Ortalama algoritmasının bu kadar yaygın kullanılmasının nedenlerinden birisi de çözüm sistematığının her zaman yakınsak olmasıdır. [64]’da gösterildiği üzere  $k$ -Ortalama tabanlı geliştirilen algoritmaların çözümleri de yakınsaktır.

Tezin bu bölümü kapsamında, çakışık nesnelerin bölütlenmesi için  $k$ -Ortalama algoritması, problemin doğasına uygun olarak algoritmaya yön bilgisinin eklenmesi ile güncellenmiş ve yeni bir yöntem önerilmiştir. Önerilen algoritma seyreltilmiş sınır piksellerinin tutulduğu  $\mathcal{V}'_s$  ile bu piksellere ait yön vektörlerinin bulunduğu  $\mathcal{D}'_s$  kümeleri ile öbek içinde kaç tane nesnenin bulunduğu sayısına ihtiyaç duyulur. Önerilen YKA’nın temel amacı,

$$\mathbf{J} = \mathbb{E} \left[ \sum_{\mathbf{v}_i \in \mathcal{V}'_s} \frac{1}{2} \min_{\mathbf{v}_k} \{g(\theta_{ik}) \|\mathbf{v}_i - \mathbf{c}_k\|^2\} \right] \quad (5.3)$$

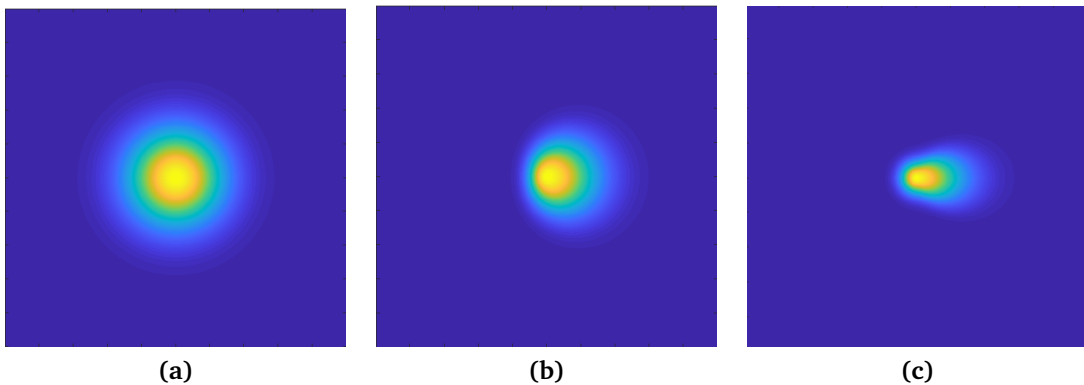
ile ifade edilen kayıp fonksiyonunu en aza indirmektedir. (5.3) ifadesinde bulunan  $g(\theta_{ik})$  fonksiyonu, sınır pikselleri ( $\mathbf{v}_i$ ) ile merkez koordinatları ( $\mathbf{c}_k$ ) arasındaki mesafeyi ağırlıklandırılmaktadır.  $g(\theta_{ik})$  fonksiyonu,

$$g(\theta_{ik}) = \frac{1}{\epsilon + e^{-0.5 \left( \frac{\theta_{ik}}{\theta_t} \right)^2}}, \quad \theta_{ik} = \cos^{-1} \left( \frac{\mathbf{d}_i \cdot (\mathbf{c}_k - \mathbf{v}_i)}{|\mathbf{d}_i| \cdot |\mathbf{c}_k - \mathbf{v}_i|} \right) \quad (5.4)$$

olarak tanımlanmaktadır. (5.4)'te  $g(\theta_{ik})$  fonksiyonu, ( $\mathbf{v}_i$ ) noktasına ( $\mathbf{d}_i$ ) gradyan doğrultusunda  $[-\theta_t, \theta_t]$  aralığında bir görüş açısı getirir. Eğer ( $\mathbf{c}_k$ ) merkez noktası bu görüş açısı içindeyse  $g(\theta_{ik})$  ağırlık katsayısı düşük olacak ve algoritma merkez nokta ile sınır piksel arasını daha yakın olarak işleyecektir. Dışındaysa  $e^{-0.5 \left( \frac{\theta_{ik}}{\theta_t} \right)^2}$  terimi 0 değerine yaklaşacağı için  $g(\theta_{ik}) = \frac{1}{\epsilon}$  olacaktır.  $\epsilon \cong 0,1$  seçilmesi durumunda görüş açısı dışında kalan noktalar arasındaki mesafe  $\times 10$  çarpanı ile ağırlıklandırılacak ve algoritma bu noktaları daha uzaktaymış gibi hesaba katacaktır.

$\mathbf{v} = [0,0]^T$  ve  $\mathbf{d} = [1,0]^T$  olduğu bir senaryo Şekil 5.3 bu duruma örnek olarak verilmiştir. Şekil 5.3(a),  $k$ -Ortalama algoritmasının uzaklık alanıdır. Sarı ve mavi tonlar sırasıyla  $\mathbf{v}$ 'ye en yakın ve uzak noktaları göstermektedir.  $k$ -Ortalama'nın *kapsam açıklığı*  $360^\circ$ 'dir. Diğer taraftan YKA için  $\theta_t = 60^\circ$  iken Şekil 5.3(b)'deki gibi  $120^\circ$ 'lik bir kapsam açıklığı oluşurken,  $\theta_t = 30^\circ$  seçildiğinde Şekil 5.3(c)'de görülen kapsam açıklığı oluşur.

YKA algoritmasının çözüm sistematığı, Bölüm 2'de verilen  $k$ -Ortalama algoritması ile



**Şekil 5.3**  $\theta_t$  eşik değerine göre YKA'nın uzaklık değerlendirmesi (Sarı: En yakın, Mavi: En uzak, Yön vektörü:  $\mathbf{d} = [1,0]^T$ ): (a)  $k$ -Ortalama algoritması isotropik olarak mesafe değerlendirmesi yaparken YKA, (b)  $\theta_t = 60^\circ$  eşik değerinde  $120^\circ$ , (c)  $\theta_t = 30^\circ$ 'ken  $60^\circ$ 'lik bir kapsam açıklığına sahiptir

aynıdır (Algoritma 2.2). YKA, Algoritma 2.2'ye gradyan doğrultusuna bağlı olarak ağırlıklandırma terimi eklemiştir. Bu sayede, Algoritma 2.2'nin uzaydaki noktaları isotropik bir değerlendirme yerine yönelim doğrultusu ile olan açısını da hesaba katarak problemin çözülmesi sağlanmıştır. YKA algoritması ile sınır pikselleri sınıflandırıldıktan sonra her bir sınıfa ait noktalara en iyi uyan elipsin parametrelerinin bulunması aşamasına gelinmiştir.

Elips parametrelerinin kestirimi için *Doğrudan Elips Uydurma* algoritması [155] kullanılmıştır. Bu algoritma elipse ait parametre kestirimini, elipsin geçeceği noktalar üzerinden problemi kuadratik en iyileme olarak ifade ettikten sonra yapmaktadır. İki boyutlu elips şeklini de içeren en genel konik ailesi,

$$F(\mathbf{u}, \mathbf{x}) = \mathbf{u} \cdot \mathbf{x} = ax^2 + bxy + xy^2 + dx + ey + f = 0 \quad (5.5)$$

olarak kapalı form biçiminde tanımlanabilir. (5.5)'de  $\mathbf{u} = [a, b, c, d, e, f]^T \in \Re^{6 \times 1}$  katsayı vektörünü ve  $\mathbf{x} = [x^2, xy, y^2, x, y, 1]^T \in \Re^{6 \times 1}$  ise her hangi bir  $\mathbf{v} = [x, y]^T$  koordinatlarına ait veri vektörünü simgeler. Sınır pikselleri üzerinden geçecek olan konik  $F(\mathbf{u}, \mathbf{x}) = 0$  olarak ifade edilirse, tüm veri seti için  $\mathbf{u}$  parametre uzayı ile oluşan denkleme ait kayıp fonksiyonu,

$$L(\mathbf{u}) = \sum_{i=1}^{N_s} F(\mathbf{u}, \mathbf{x}_i)^2 \quad (5.6)$$

ile hesaplanır. (5.6) içinde bulunan  $L(\mathbf{u})$  ifadesinin matris formu olarak  $L(\mathbf{u}) = \|\mathbf{D}\mathbf{u}\|_2^2$  olarak gösterilebilir. Burada  $\mathbf{D} = [\mathbf{x}_1, \mathbf{x}_2, \mathbf{x}_3, \dots, \mathbf{x}_{N_s}]^T \in \Re^{N_s \times 6}$  veri matrisidir. Dolayısıyla kayıp fonksiyonu,  $L(\mathbf{u}) = \mathbf{u}^T \mathbf{S} \mathbf{u}$  olarak  $\mathbf{S} = \mathbf{D}^T \mathbf{D}$  yayılım matrisi kullanılarak yeniden düzenlenebilir. Tanımlanan  $L(\mathbf{u})$  konik kayıp fonksiyonunun elips formuna sahip olabilmesi için kapalı forma ait diskriminant değerinin  $4ac - b^2 = 1$  olması gerekmektedir. Bu kısıt (constrain) altında veriler üzerinden geçecek olan elips problemi,

$$\underset{\mathbf{u}}{\operatorname{argmin}}\{L(\mathbf{u})\} \quad s.t \quad \mathbf{u}^T \mathbf{C} \mathbf{u} = 1, \quad \mathbf{C} = \begin{bmatrix} 0 & 0 & 2 & 0 & 0 & 0 \\ 0 & -1 & 0 & 0 & 0 & 0 \\ 2 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 \end{bmatrix} \quad (5.7)$$

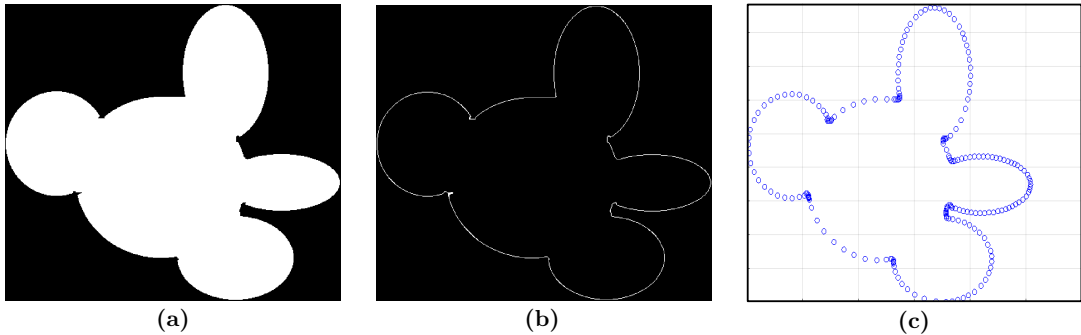
olarak kuadratik en iyileme problemi olarak tanımlanabilir [155]. (5.7) içinde bulunan  $\mathbf{C}$  matrisi ile sağlanan kısıt sayesinde  $L(\mathbf{u})$  minimize edilirken,  $\mathbf{u}$  parametreleri ile tanımlanan şeklin elips olmasına zorlanmıştır. Ayrıca (5.7)'nin iç-bükey olması onu en iyileme metotları ile çözülebilir kılmaktadır. YKA'nın işleyişi aşamalar ile verildikten sonra algoritmanın çalışması sentetik bir problem üzerinden görselleştirilmiştir.

### 5.2.1 Sentetik Problem Üzerinde YKA'nın Testi

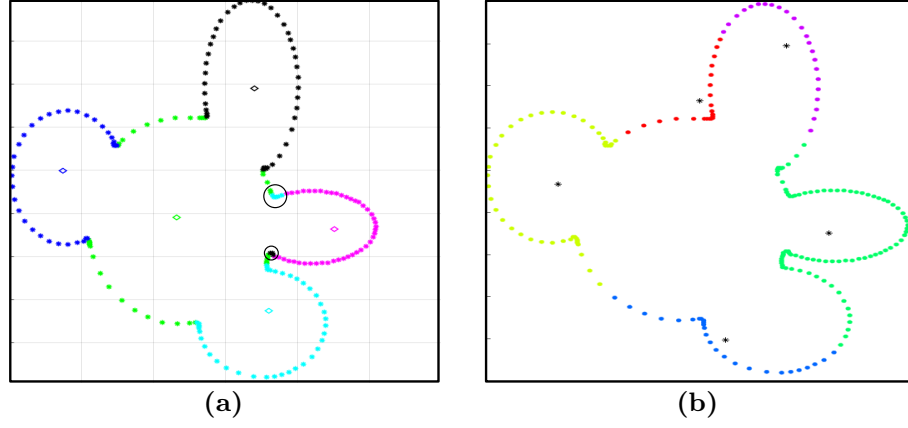
Sentetik olarak üretilen Şekil 5.4(a)'daki ikili görüntü, yüksek çözünürlüklü çakışık bir öbeği temsil etmektedir. Sobel operatörü ile öbeğe ait sınır piksellerinin çıkartılması Şekil 5.4(b)'de gösterilmiştir. Toplam 1690 tane sınır pikseline ait gradyan yönelimlerinin hesaplanmasının ardından UPS algoritması ile bu sayı 260 örneğe kadar indirgenmiştir. Seyreltilmiş sınır piksellerine dair görsel Şekil 5.4(c)'de sunulmuştur. UPS sonrasında elde edilen  $\mathcal{V}'$  ve  $\mathcal{D}'$  kümeleri ile  $\theta_t = 50^\circ$  ve  $N_c = 5$  parametreleri ile YKA'ya verilmiştir.

YKA döngü içinde bulunan sınıf merkezleri üzerinden tanımlanan  $\|\mathbf{c}_{k+1} - \mathbf{c}_k\| \leq 0.01$  koşulundan dolayı 10 iterasyon sonra durmuştur. Algoritma çıktısı Şekil 5.5(a)'da gösterilmektedir. Sınır piksellerine eklenen yönsel kapsam açıklığı sayesinde sınıf merkezleri başka sınıflara ait noktalara yakın olsalar da bu noktalar, yönelimlerinden dolayı yanlış sınıflandırılmaktan korunmuşlardır. Diğer taraftan  $k$ -Ortalama algoritması uzaydaki noktaları eş-uzaklık üzerinden değerlendirdiği için öbeklerin ayrıştırılması için yetkin değildir. Algoritma çalıştırıldığında Şekil 5.5(b) elde edilmektedir. Tüm noktaların aynı ağırlıkla değerlendirilmesi problemin çözümüne katkı sunmamaktadır.

Şekil 5.5(a) üzerinde özellikle dönüm-büküm noktalarında (çemberler içinde) aykırı değerler (outliers) oluşabilmektedir. Bu değerler özellikle elips uydurma aşamasında



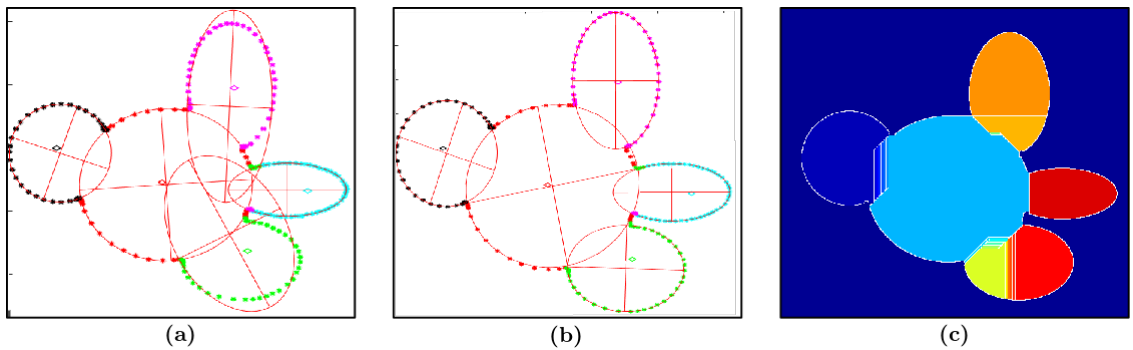
**Şekil 5.4** Sentetik olarak oluşturulan bir nesne: Yapısında 5 tane çakışık nesnenin bulunduğu şekil (a) üzerinden ilk olarak sınır pikselleri çıkartılmıştır (b). Sonrasında UPS algoritması kullanılarak sınır pikselleri indirgenmiştir (c). Bu sayede algoritma 1690 tane sınır pikseli yerine 260 örnek üzerinden problemi çözmektedir



**Şekil 5.5** Öbeğin sınır piksellerinin sınıflandırılması: YKA, girdi olarak seyreltilmiş sınır pikselleri ile öbek içinde kaç nesnenin bulunduğu bilgisin aldıktan sonra **(a)**'yı üretmektedir. **(b)** ise  $k$ -Ortalama algoritmasının çıktısıdır

ciddi bozulmalara yol açmaktadır. Buna örnek olarak Şekil 5.6(a) verilebilir. Elips uydurma algoritması tüm noktalar için minimize etmeye çalıştığı için elipsin yönelimi aykırı değerlere doğru olabilmekte bu nedenle genel nokta dağılımından uzaklaşmaktadır. Şekil 5.6(a)'da yeşil ve magenta renkleri ile gösterilmiş sınıflarda bu durum ortaya çıkmıştır. Aykırı değerlerin piksel seyreltmeden dolayı sayıları yüksek olmamaktadır. Bunların ayrıştırılması için Bölüm 2'de ele alınan DBSCAN algoritmasından faydalanılmıştır. DBSCAN, her sınıf içinde az olan yerel piksel kümelerini tespit etmiştir. Ardından  $N_p = 5$  parametresi ile 5'ten daha az sayıda örnek içeren yerel kümeler elips uydurma işleminin dışında tutulmuştur. DBSCAN kullanılması ile Şekil 5.6(b) elde edilmektedir. Buradan öbek içindeki tüm nesnelerin elips olarak bölütlenebildiği görülmektedir.

YKA aynı zamanda literatürde çakışık nesnelerin bölütlenmesi için sıkça kullanılan



**Şekil 5.6** Sınır piksellerinin sınıflandırılmasında aykırı sınıflandırılmış örnekler çıkmaktadır. Bu örnekler elips uydurma algoritmasına bozucu etkileri olur **(a)**. Bunun giderilmesi için DBSCAN ile aykırı değerler elenmiştir. Bu aşamadan sonra YKA'nın nihayi çıktısı **(b)** elde edilir. **(c)** ise aynı problem için su kuyusu algoritmasının çıktısıdır

su kuyusu algoritması [158] ile karşılaştırılmıştır. Şekil 5.6(c) su kuyusu algoritması tarafından üretilmiştir. Bu algoritma ikili görüntüleri kullanarak mesafe dönüşümü üzerinden çalışmaktadır. Geniş-ölçek öbeklerde yerel en yüksek noktalar fazla çıktığından dolayı algoritma öbeği aşırı bölütler (over-segmentation). Diğer taraftan YKA algoritması gürbüz çalışması seçilen parametrelere duyarlıdır. Bu durumun giderilmesi için YKA, veri-uyarlamalı parametre en iyilemesi sürecinden geçmelidir.

Bölüm 5 kapsamında ortaya sunulan UPS algoritması ve YKA, uygulama olarak, biyomedikal alanda dokuların sınıflandırılması ve hastalıkların derecelendirilmesi için kilit bir role sahip olan çakışık hücre problemi üzerinde kullanılmıştır.

### 5.3 Uygulama Alanı: Çakışık Hücrelerin Bölütlenmesi

Yardımcı histolojik görüntü analiz yöntemlerinin geliştirilmesi için hücre bölütleme teknikleri hayati bir öneme sahiptir. Çakışık hücre problemi düşük kontrast, zayıf sınırlar ve örtüşen hücreler nedeniyle çok zor bir problemdir. Bu sebeplerden dolayı bölütlemenin otomatik olarak yapılması için kararlı ve makul süre içinde cevap verebilen algoritmalara ihtiyaç vardır. Bölütlenen hücrelerin morfolojik özellikleri üzerinden elde edilen istatistikler, patolojlara doğru tanıyı koymalarında yardımcı olmaktadır. Sayısal verilere dayalı tanımlar gözlemci-içi ve gözlemciler arası değişimliliği minimize etmektedir. Doğru tanı sayesinde hastalıklar için gerekli olan tedavi süreci erken evrede başlatılabilmektedir.

#### 5.3.1 Literatür Özeti

Çakışık hücre problemi için literatürde birçok çalışma bulunmaktadır. Akram ve ark. [159]'de yaptıkları çalışmada çakışık hücrelerin tespiti (detection) ve tek tek hücrelerinin ayrıştırılması için ESA tabanlı bir çözüm önermişlerdir. Geliştirilen ESA modeli iki alt ağıdan oluşmaktadır. İlk ağın amacı hücre adaylarının etrafına sınırlayıcı kutu (bounding box) koymaktır. Ağ çıkışında oluşan birçok adaya ait kutular, skor değerleri üzerinden maksimum olmayan baskılama (non-maxima suppression) ile en uygun olanları seçilmiştir. Belirlenen uygun adaylar ikinci ağa aktarılarak sınırları tespit edilmiştir. Bu nedenle önerilen yöntem hem merkezlerinin bulunduğu hem de sınır piksellerinin belirlendiği veri setlerine ihtiyaç duymaktadır. Algoritma H&E görüntülerinin yanında floresan ışığı ile elde edilen görüntüler kullanılarak test edilmiştir. Song ve ark. ise problemin çözümü için kademeli seyrek regresyon zinciri modelini öne sürmüşlerdir [160]. Çevrim-içi (online) öğrenme yetisine sahip ayrılabilir (separable) ESA yapısına sahip modelin özgün yanı çekirdek ve hücre sınırlarını eş zamanlı olarak aynı çerçevede regresyon problemi olarak tespit

edebilmesidir. Sınır ve çekirdekler belirlendikten sonra tekil hücreler, en kısa mesafe aidiyeti ve bağ interpolasyonu (B-Spline) ile belirlenmiştir. Gradyan arttırım regresyon (gradient boosting regression) algoritması ile daha gürbüz sınır piksel kestirimi yapılabilmektedir. Önerilen çalışma farklı hastalıklara ait farklı veri setlerinde test edilmiştir. [161]'deki çalışmada ise serviks pap smear görüntüleri için eliptik modellemeye karşın yıldız-şekli öncülünü kullanarak çakışık hücreler, Voroni enerji denklemi ile modellendikten sonra ayrıştırılmaktadır. Yöntemde ilk olarak HoG ile çekirdek adayları belirlenmiştir. Ardından karar ağacı sınıflandırıcısı ile adaylar arasından uygun olanlar sınıflandırılarak sınır eşleşmesi için enerji fonksiyonunda kullanılmıştır. Polar koordinatlarda geliştirilen enerji fonksiyonu üç temel enerji bileşeninin ağırlıklı toplamından oluşmaktadır. Bunlar, yıldız-şekli öncülünü sağlayan terim ( $E_s$ ), ne kadar hücrenin üstü üstte çakışacağını kontrol eden Voroni enerji terimi ( $E_v$ ) ve bölgesel veri terimidir ( $E_r$ ). Bu terimlerin yanında ek olarak düzenleyici (regularization) terimi ( $R$ ) eklenmiştir. Yazarların bildirdiğine göre çalışma herhangi bir eğitim olmadan problemi analitik olarak hızlı bir şekilde çözebilmektedir. Pap smear görüntüsü üzerine yapılan bir diğer çalışma ise [162]'dir. Song ve ark.'larının yaptığı çalışmada, hücre bileşenlerinin tahmini için çok-ölçekli (multi-scale) ESA yapısı kullanılmıştır. Görüntü içindeki her bir piksel için hücreye ait bir parça olma olasılığı hesaplandıktan sonra, dinamik çoklu şablon deformasyonu adı verilen seviye-kümesi (level-set) temelinde oluşturulan bir enerji fonksiyonu kullanılarak çakışık hücreler sınıflandırılmıştır. Plissiti ve ark. ise çakışık hücre problemi için uyarlamalı aktif fiziksel model geliştirmişlerdir [163]. Önerdikleri model görüntü içindeki birinci ve ikinci gradyanlardan faydalanarak örtüşen bölgelerin ve çekirdek öbeklerinin fiziksel modelini çıkartmaktadır. Model bir eğitim fazına ihtiyaç duymaktadır. Uzman bir patolog tarafından çizilen hücre sınırları üzerinden fiziksel model, şekle denk gelecek parametreleri uyarlamalı olarak kestirmektedir. Test aşamasında ise mesafe dönüşümü temelinde yapılan ilkleme sonrasında yerel modeller çakışık hücreleri bölütlemeye başlarlar. [164]'deki çalışmada en kısa açıklama uzunluğu temelinde analitik bir yöntem geliştirilmiştir. İki aşamadan oluşan yöntemin ilk aşamasında uzaklık-haritalaması kısıtı altında öne sürülen *Laplacian of Gaussian (LoG)* filtreler ile hücre merkezlerinin adayları tespit edilmiştir. Bu adaylar arasında çekirdek belirteçleri (marker) ile belirlenenler yardımıyla, çekirdek sınırı üzerinden ayırım-nokta (split-point) adayları hesaplanmaktadır. İteratif morfolojik erozyon ile eşleştirilen ayırım-nokta çiftleri üzerinden kabaca hücreler ayrılmışlardır. Bu aşamadan sonra merkez noktaları tespit edilerek hücre sınırlarının sınıflandırılması olan ikinci aşamaya geçilmiştir. Belirlenen hücre merkezleri ile sınır pikselleri arasında en kısa mesafeye dayalı bir eşleştirme yapılmaktadır. Bu eşleştirmede aday hücre merkezi ile sınır pikselin arasındaki açı da hesaba katılmaktadır. Sınır piksellerin sınıflandırılmasından sonra b-spline ile iki boyutlu hücre şekli elde edilir.

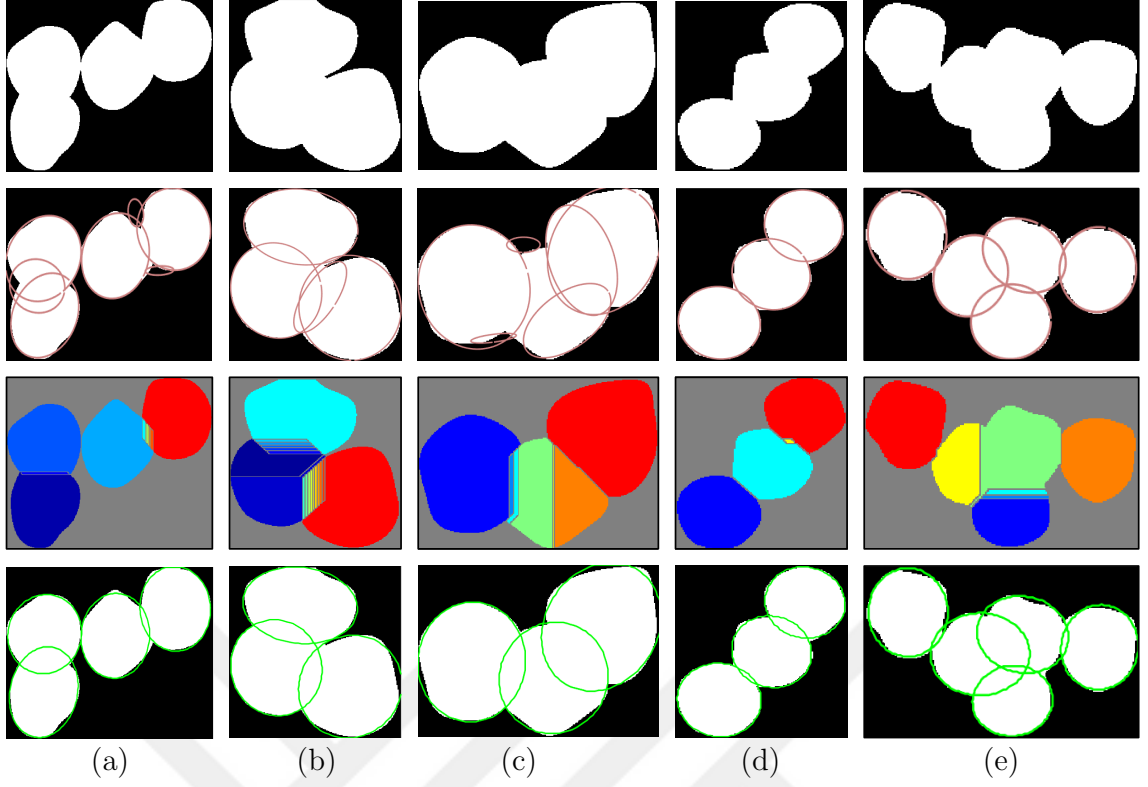
Literatürde incelenen yöntemler göz önüne alındığında, yüksek çözünürlüğe sahip yoğun sınır pikseli içeren öbeklerin sınıflandırılmasında algoritmaların işlem karmaşıklığının indirgenmesi ön planda tutulmamaktadır. Buradan yola çıkarak tezin bu bölümünde, UPS algoritması ile desteklenmiş YKA'nın çakışık hücrelerin ayrıştırılmasında kullanılmasına karar verilmiştir. Test aşamasında önerilen algoritma su-kuyusu algoritması ve elips modelleme temelinde çalışan [123]'deki algoritma ile test edilmiştir.

### 5.3.2 YKA ile Problemin Çözümü

YKA algoritmasının kullanım alanı olarak, literatürde böylesi algoritmaların en çok kullanıldığı histopatolojideki çakışık hücrelerin bölütlenmesi alanı tercih edilmiştir. Hücrelerin ayrı ayrı ele alınarak geometrik ve fizyolojik özniteliklerin çıkartılması, patoloğların nesnel tanı koymaları için önemli bir bilgi sağlamaktadır. Doğru tanının tespiti ile hem hastanın hayati durumu için uygun adımın atılması sağlanırken, hem de gereksiz bir tedavi sürecinin önüne geçilerek fazladan oluşabilecek maddi harcamaların önüne geçilebilir. Bu aşamada, patoloğlara hızlı ve faydalı bilgiyi verebilecek algoritmalar önem kazanır.

Tez çerçevesinde geliştirilen YKA, bu problem için önemli bir çözüm aracı olarak kullanılabilir. Yüksek çözünürlüğe sahip görüntüler üzerinden elde edilen hücre öbeklerinin sınır pikselleri UPS ile seyreltilerek makul süreler içerisinde yanıt alınabilir. Algoritmanın testi için SimCep [165] programı ile üretilen sentetik florosan hücre görüntüleri kullanılmıştır. Karşılaştırmalar için su-kuyusu algoritması ile birlikte [149] tercih edilmiştir. İlk olarak Şekil 5.7(a) (b) (c)'de bulunan yüksek çözünürlüğe sahip görüntüler üzerinden algoritmalar test edilmiştir. Ardından Şekil 5.7(d) (e) düşük çözünürlüklü hücre öbekleri ile işlemler tekrarlanmıştır.

Şekil 5.7'de verilen görüntüde ilk sıra, öbek halinde bulunan florasan hücre yığınlarını göstermektedir. İkinci sıra ise [149] algoritmasının çıkışını gösterirken üçüncü satır ise su kuyusu algoritmasının çıktısını vermektedir. Son satır olan dördüncü satırda da YKA'nın bulguları sunulmuştur. Şekil 5.7(a) (b) (c)'ye verilen görüntülerde sırası ile 1140, 746, 709 sınır pikseli bulunmaktadır. UPS algoritması ile bu değerler YKA için 167 (6.80), 90 (8.27) ve 91 (7.78) değerlerine kadar indirgenmektedir. Parantez içindeki değerler indirgenme oranını gösterir. Öbek içeriğine ait pikseller yoğun olduğu için su-kuyusu ve [149] algoritması doğru merkezleri tespit edememektedir. YKA ise kendisine verilen hücre sayısı kadar merkezi iteratif olarak doğru bir şekilde tespit edebilmektedir. Şekil 5.7(d) (e) örneklerinde ise çözünürlük düşüktür. Öbeklere ait sınır pikselleri sırası ile 416 ve 627'dir. YKA için bu değerler 92 (4.52) ve 146 (4.29) olarak hesaplanmıştır. Çözünürlük düştüğü içi ardışıl sınır pikselleri



**Şekil 5.7** YKA'nın [149] ve su-kuyusu algoritmaları ile karşılaştırılması: İlk sırada florasın hücrelerden elde edilmiş ikili görüntüler bulunmaktadır. İkinci ve üçüncü sıralarda ise [149] ve su-kuyusu algoritmalarının çıktıları ve dördüncü satırda da YKA'nın sonuçları gösterilmiştir. (a-c) olan hücreler yüksek çözünürlüğe sahipken (d) ve (e) düşük çözünürlük ile elde edilmiştir

arasındaki açısal değişim yüksek olmaktadır. Bu nedenle indirgeme oranları da uyarlamalı olarak düşük çıkmaktadır. UPS algoritması sayesinde YKA için elverişli olan pikseller, çözünürlüğe bağlı olarak seçilebilmektedir. Düşük çözünürlük altında [149] algoritmasının da YKA gibi doğru çözümler sunduğu görülmektedir. Piksel yoğunluğunun azalması ile birlikte merkezi ilkendirme başarılı bir şekilde yapılabilmekte ve bu sayede sınır pikselleri etkili bir şekilde sınıflandırılabilir. Su-kuyusu algoritması ise Şekil 5.7(d)'de doğru sonuç verirken, Şekil 5.7(e)'de ise başarısız bir bölütlemeye çalışmıştır. Bunun nedeni ise aday-merkez tespitinde kullanılan mesafe dönüşümünün yetersiz kalmasından kaynaklanmaktadır.

## 5.4 Sonuç

YKA, UPS algoritmasının desteği ile seyrek örnekler üzerinden çakışık nesne problemi için etkin bir çözüm olmaktadır.  $k$ -Ortalama algoritmasına eklenen açısal görüş ağırlıklandırması ile merkezler iteratif olarak bulunabilmektedir. Böylece, merkez ile sınır pikselleri arasındaki mesafenin yanı sıra, sınır üzerinde oluşan gradyan yöneliminin sağladığı önem faktörü ile  $k$ -Ortalama bakış-açısı-etkisi doğrultusunda

kayıp fonksiyonunu minimize etmektedir. Algoritma [149] ve su-kuyusu algoritması ile kıyaslandığında yüksek çözünürlükte üstün bir başarı sergilemektedir. Düşük çözünürlük altında ise [149] ile benzer sonuçlar gösterir. YKA'nın tam otomatik olarak çalışması için ön-işleme adımı olarak doğru yapılmış bölütleme ile öbek içinde bulunan hücre sayısının bulunması gerekir. Bu adımların geliştirilmesi YKA merkezinde yapılması planlanan çalışmalar arasındadır. Özellikle bölütleme ve hücre merkezlerinin veya sayılarının tespiti için ESA tabanlı bir model geliştirilmesi hedeflenmektedir.



## 6 Sonuç ve Öneriler

---

Ortaya koyulan tez kapsamında büyük veri dahilinde ele alınan geniş ölçekli veri setleri için sınıflandırma ve bölütleme problemlerine yönelik ESA ve istatistiksel model tabanlı çözümler sunulmuştur. Bölüm-3 ve Bölüm-4'te sınıflandırma için geliştirilen model ve algoritma işlenmiş, Bölüm-5'te ise bölütleme probleminin çözümü için öne sürülen özgün algoritma açıklanmıştır.

Ele alınan sınıflandırma problemlerinden ilki, Bölüm-3 dahilinde işlenen çok sınıflı ve sınıf başına az örneğin bulunduğu veri setlerinin sınıflandırılmasıdır. Çözüm olarak ESA tabanlı *G-ESA* modeli bir ön-işleme bloğunun desteği ile sunulmuştur. Verilerin ham olarak işlenmesi yerine, onların model parametrelerine uygun bir yapıya dönüştürülmesinin avantajları gösterilmiştir. Ayrıca gömülü uzayda sınıf-içi veri miktarının az olmasından yola çıkarak, test örneklerinin 1-NN tabanlı sınıfların merkez vektörleri üzerinden sınıflandırılması önerilmiştir. Geliştirilen model imza tanıma alanında uygulanmıştır. Bu doğrultuda gerçek kişilerden toplanan MCYT, CEDAR veri setleri ile birlikte sentetik olarak 4k kişi üzerinden üretilen GPDS-Syntetic veri seti kullanılmıştır. Öne sürülen modelin, VGG tipi ağlar ile kıyaslandığında üstün bir başarı sergilediği gösterilmiştir.

İkinci problem olarak az sınıflı ve sınıflara ait özniteliklerin istatistiksel olarak çıkartılabileceği durumlara yönelik yerel histogram tabanlı algoritma, Bölüm-4'te öne sürülmüştür. Elde edilen istatistiksel öznitelikler arasında simetrik KL diverjansı kullanılarak benzerlik metriği ölçülmüştür. Bu değerler üzerinden ağırlıklı *k*-NN algoritması kullanılarak örnekler sınıflandırılmıştır. Önerilen yöntemin testi için serviks kanserinde dokuların sınıflandırılması seçilmiştir. Veri seti olarak İstanbul Medipol Üniversitesi Patoloji Laboratuvarının sağladığı veri seti kullanılmıştır. Patologların doğruluk başarımının %85 civarında olduğu bu zorlu veri setinde, önerilen algoritma %75 civarında oldukça önemli bir doğruluk başarımı kaydetmiştir. Böylesi problemlerde literatürde sıklıkla tercih edilen morfometrik öznitelikler ile kıyaslanmasında ise öne sürülen algoritmanın F-skor ve doğruluk kıstasları bakımından daha etkin ve kullanışlı olduğu gösterilmiştir.

Üçüncü problem olarak bölütleme, Bölüm-5'te dahilinde irdelenmiştir. Eğiticişiz öğrenme altında yüksek çözünürlüğe sahip görüntülerde çakışık nesnelerin ayrıştırılmasına yönelik UPS tabanlı YKA ortaya koyulmuştur. Algoritmanın temelini,  $k$ -Ortalama algoritmasının çakışık nesneye ait sınır pikselleri ve sınıf merkezleri arasındaki uzaklık ile açısai bilginin de hesaba katılması oluşturmaktadır. Sentetik hücre görüntüleri üzerinden test edilen algoritma seyreltilmiş sınır pikselleri üzerinden hücre sınırlarını modelleyerek önemli bir katkı sunmuştur.



- [1] M. Chen, S. Mao, and Y. Liu, “Big data: A survey,” *Mobile Networks and Applications*, vol. 19, no. 2, pp. 171–209, 2014.
- [2] J. S. Ward and A. Barker, “Undefined by data: A survey of big data definitions,” *arXiv preprint arXiv:1309.5821*, 2013.
- [3] X. Yao and G. Li, “Big spatial vector data management: A review,” *Big Earth Data*, vol. 2, no. 1, pp. 108–129, 2018.
- [4] W. Zou, W. Jing, G. Chen, Y. Lu, and H. Song, “A survey of big data analytics for smart forestry,” *IEEE Access*, vol. 7, pp. 46 621–46 636, 2019.
- [5] R. Lu, H. Zhu, X. Liu, J. K. Liu, and J. Shao, “Toward efficient and privacy-preserving computing in big data era,” *IEEE Network*, vol. 28, no. 4, pp. 46–50, 2014.
- [6] M. Strohbach, H. Ziekow, V. Gazis, and N. Akiva, “Towards a big data analytics framework for iot and smart city applications,” in *Modeling and Processing for Next-Generation Big-Data Technologies*, Springer, 2015, pp. 257–282.
- [7] Z. Nyikes and Z. Rajnai, “Big data, as part of the critical infrastructure,” in *2015 IEEE 13th International Symposium on Intelligent Systems and Informatics (SISY)*, IEEE, 2015, pp. 217–222.
- [8] H. Guo, L. Wang, F. Chen, and D. Liang, “Scientific big data and digital earth,” *Chinese Science Bulletin*, vol. 59, no. 35, pp. 5066–5073, 2014.
- [9] Y. Demchenko, C. De Laat, and P. Membrey, “Defining architecture components of the big data ecosystem,” in *2014 International Conference on Collaboration Technologies and Systems (CTS)*, IEEE, 2014, pp. 104–112.
- [10] R. Kitchin, “The real-time city? big data and smart urbanism,” *GeoJournal*, vol. 79, no. 1, pp. 1–14, 2014.
- [11] M. Batty, “Big data, smart cities and city planning,” *Dialogues in Human Geography*, vol. 3, no. 3, pp. 274–279, 2013.
- [12] H. Song, R. Srinivasan, T. Sookoor, and S. Jeschke, *Smart cities: foundations, principles, and applications*. John Wiley and Sons, 2017.
- [13] Z. Tufekci, “Big questions for social media big data: Representativeness, validity and other methodological pitfalls,” in *Eighth International AAAI Conference on Weblogs and Social Media*, 2014.
- [14] J. Wang, Y. Wu, N. Yen, S. Guo, and Z. Cheng, “Big data analytics for emergency communication networks: A survey,” *IEEE Communications Surveys and Tutorials*, vol. 18, no. 3, pp. 1758–1778, 2016.

- [15] G. V. Attigeri, M. P. MM, R. M. Pai, and A. Nayak, "Stock market prediction: A big data approach," in *TENCON 2015-2015 IEEE Region 10 Conference*, IEEE, 2015, pp. 1–5.
- [16] S. Sohangir, D. Wang, A. Pomeranets, and T. M. Khoshgoftaar, "Big data: Deep learning for financial sentiment analysis," *Journal of Big Data*, vol. 5, no. 1, p. 3, 2018.
- [17] R. Lin, Z. Ye, H. Wang, and B. Wu, "Chronic diseases and health monitoring big data: A survey," *IEEE reviews in Biomedical Engineering*, vol. 11, pp. 275–288, 2018.
- [18] K. Shvachko, H. Kuang, S. Radia, R. Chansler, *et al.*, "The hadoop distributed file system," in *MSST*, vol. 10, 2010, pp. 1–10.
- [19] N. K. Alham, M. Li, Y. Liu, and S. Hammoud, "A mapreduce-based distributed svm algorithm for automatic image annotation," *Computers and Mathematics with Applications*, vol. 62, no. 7, pp. 2801–2811, 2011.
- [20] G. Caruana, M. Li, and M. Qi, "A mapreduce based parallel svm for large scale spam filtering," in *2011 8. International Conference on Fuzzy Systems and Knowledge Discovery (FSDK)*, IEEE, vol. 4, 2011, pp. 2659–2662.
- [21] C.-T. Chu, S. K. Kim, Y.-A. Lin, Y. Yu, G. Bradski, K. Olukotun, and A. Y. Ng, "Map-reduce for machine learning on multicore," in *Advances in Neural Information Processing Systems*, 2007, pp. 281–288.
- [22] S. Shalev-Shwartz, Y. Singer, N. Srebro, and A. Cotter, "Pegasos: Primal estimated sub-gradient solver for svm," *Mathematical Programming*, vol. 127, no. 1, pp. 3–30, 2011.
- [23] P. Rebentrost, M. Mohseni, and S. Lloyd, "Quantum support vector machine for big data classification," *Physical Review Letters*, vol. 113, no. 13, p. 130 503, 2014.
- [24] I. W. Tsang, J. T. Kwok, and P.-M. Cheung, "Core vector machines: Fast svm training on very large data sets," *Journal of Machine Learning Research*, vol. 6, no. Apr, pp. 363–392, 2005.
- [25] Z. Deng, X. Zhu, D. Cheng, M. Zong, and S. Zhang, "Efficient knn classification algorithm for big data," *Neurocomputing*, vol. 195, pp. 143–148, 2016.
- [26] J. Maillo, S. Ramirez, I. Triguero, and F. Herrera, "Knn-1s: An iterative spark-based design of the k-nearest neighbors classifier for big data," *Knowledge-Based Systems*, vol. 117, pp. 3–15, 2017.
- [27] A. K. Jain, "Data clustering: 50 years beyond k-means," in *Joint European Conference on Machine Learning and Knowledge Discovery in Databases*, Springer, 2008, pp. 3–4.
- [28] D. Arthur and S. Vassilvitskii, "K-means++: The advantages of careful seeding," in *Proceedings of the Eighteenth Annual ACM-SIAM Symposium on Discrete Algorithms*, Society for Industrial and Applied Mathematics, 2007, pp. 1027–1035.
- [29] D. Pelleg, A. W. Moore, *et al.*, "X-means: Extending k-means with efficient estimation of the number of clusters.," in *In Proceedings of the 17th International Conf. on Machine Learning (ICML)*, vol. 1, 2000, pp. 727–734.

- [30] E. Schubert and P. J. Rousseeuw, “Faster k-medoids clustering: Improving the pam, clara, and clarans algorithms,” *arXiv preprint arXiv:1810.05691*, 2018.
- [31] L. Bottou and Y. Bengio, “Convergence properties of the k-means algorithms,” in *Advances in Neural Information Processing Systems*, 1995, pp. 585–592.
- [32] X. Cai, F. Nie, and H. Huang, “Multi-view k-means clustering on big data,” in *Twenty-Third International Joint Conference on Artificial Intelligence*, 2013.
- [33] X. Cui, P. Zhu, X. Yang, K. Li, and C. Ji, “Optimized big data k-means clustering using mapreduce,” *The Journal of Supercomputing*, vol. 70, no. 3, pp. 1249–1259, 2014.
- [34] Y. LeCun, L. Bottou, Y. Bengio, P. Haffner, *et al.*, “Gradient-based learning applied to document recognition,” *Proceedings of the IEEE*, vol. 86, no. 11, pp. 2278–2324, 1998.
- [35] A. Krizhevsky, I. Sutskever, and G. E. Hinton, “Imagenet classification with deep convolutional neural networks,” in *Advances in Neural Information Processing Systems*, 2012, pp. 1097–1105.
- [36] J. Deng, W. Dong, R. Socher, L.-J. Li, K. Li, and L. Fei-Fei, “Imagenet: A large-scale hierarchical image database,” in *Computer Vision and Pattern Recognition*, 2009.
- [37] C. Szegedy, W. Liu, Y. Jia, P. Sermanet, S. E. Reed, D. Anguelov, D. Erhan, V. Vanhoucke, and A. Rabinovich, “Going deeper with convolutions,” *2015 IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, vol. abs/1409.4842, 2014.
- [38] K. Simonyan and A. Zisserman, “Very deep convolutional networks for large-scale image recognition,” *arXiv preprint arXiv:1409.1556*, 2014.
- [39] K. He, X. Zhang, S. Ren, and J. Sun, “Deep residual learning for image recognition,” in *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 2016, pp. 770–778.
- [40] A. Garcia-Garcia, S. Orts-Escolano, S. Oprea, V. Villena-Martinez, and J. Garcia-Rodriguez, “A review on deep learning techniques applied to semantic segmentation,” *arXiv preprint arXiv:1704.06857*, 2017.
- [41] J. S. Chung and A. Zisserman, “Lip reading in the wild,” in *Asian Conference on Computer Vision*, Springer, 2016, pp. 87–103.
- [42] J. Donahue, L. Anne Hendricks, S. Guadarrama, M. Rohrbach, S. Venugopalan, K. Saenko, and T. Darrell, “Long-term recurrent convolutional networks for visual recognition and description,” in *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 2015, pp. 2625–2634.
- [43] A. Tuan Tran, T. Hassner, I. Masi, and G. Medioni, “Regressing robust and discriminative 3d morphable models with a very deep neural network,” in *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 2017, pp. 5163–5172.
- [44] D. H. Hubel and T. N. Wiesel, “Receptive fields and functional architecture of monkey striate cortex,” *The Journal of physiology*, vol. 195, no. 1, pp. 215–243, 1968.

- [45] C. Nwankpa, W. Ijomah, A. Gachagan, and S. Marshall, "Activation functions: Comparison of trends in practice and research for deep learning," *arXiv preprint arXiv:1811.03378*, 2018.
- [46] L. B. Godfrey and M. S. Gashler, "A continuum among logarithmic, linear, and exponential functions, and its potential to improve generalization in neural networks," in *2015 7th International Joint Conference on Knowledge Discovery, Knowledge Engineering and Knowledge Management (IC3K)*, IEEE, vol. 1, 2015, pp. 481–486.
- [47] W. Duch and N. Jankowski, "Survey of neural transfer functions," *Neural Computing Surveys*, vol. 2, no. 1, pp. 163–212, 1999.
- [48] V. Nair and G. E. Hinton, "Rectified linear units improve restricted boltzmann machines," in *Proceedings of the 27th International Conference on Machine Learning (ICML-10)*, 2010, pp. 807–814.
- [49] B. Xu, N. Wang, T. Chen, and M. Li, "Empirical evaluation of rectified activations in convolutional network," *arXiv preprint arXiv:1505.00853*, 2015.
- [50] D. Scherer, A. Müller, and S. Behnke, "Evaluation of pooling operations in convolutional architectures for object recognition," in *International Conference on Artificial Neural Networks*, Springer, 2010, pp. 92–101.
- [51] S. Xie, R. Girshick, P. Dollár, Z. Tu, and K. He, "Aggregated residual transformations for deep neural networks," in *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 2017, pp. 1492–1500.
- [52] S. Ioffe and C. Szegedy, "Batch normalization: Accelerating deep network training by reducing internal covariate shift," in *International Conference on Machine Learning*, 2015, pp. 448–456.
- [53] G. Huang, Y. Sun, Z. Liu, D. Sedra, and K. Q. Weinberger, "Deep networks with stochastic depth," in *European Conference on Computer Vision*, Springer, 2016, pp. 646–661.
- [54] I. J. Goodfellow, D. Warde-Farley, M. Mirza, A. Courville, and Y. Bengio, "Maxout networks," *arXiv preprint arXiv:1302.4389*, 2013.
- [55] Y. Wu and K. He, "Group normalization," in *Proceedings of the European Conference on Computer Vision (ECCV)*, 2018, pp. 3–19.
- [56] V. Vapnik, *The nature of statistical learning theory*. Springer science and business media, 2013.
- [57] C. Cortes and V. Vapnik, "Support-vector networks," *Machine Learning*, vol. 20, no. 3, pp. 273–297, 1995.
- [58] B. E. Boser, I. M. Guyon, and V. N. Vapnik, "A training algorithm for optimal margin classifiers," in *Proceedings of the Fifth Annual Workshop on Computational Learning Theory*, ser. COLT '92, Pittsburgh, Pennsylvania, USA: ACM, 1992, pp. 144–152, ISBN: 0-89791-497-X. DOI: 10.1145/130385.130401. [Online]. Available: <http://doi.acm.org/10.1145/130385.130401>.
- [59] J. Duchi, E. Hazan, and Y. Singer, "Adaptive subgradient methods for online learning and stochastic optimization," *Journal of Machine Learning Research*, vol. 12, no. Jul, pp. 2121–2159, 2011.

- [60] S. Ruder, “An overview of gradient descent optimization algorithms,” *arXiv preprint arXiv:1609.04747*, 2016.
- [61] D. Kingma and J. Ba, “Adam: A method for stochastic optimization,” *arXiv preprint arXiv:1412.6980*, 2014.
- [62] C. E. Shannon, “A mathematical theory of communication,” *Bell System Technical Journal*, vol. 27, no. 3, pp. 379–423, 1948.
- [63] M. Ester, H.-P. Kriegel, J. Sander, X. Xu, *et al.*, “A density-based algorithm for discovering clusters in large spatial databases with noise,” in *Proceedings of the Second International Conference on Knowledge Discovery and Data Mining*, vol. 96, 1996, pp. 226–231.
- [64] S. Z. Selim and M. A. Ismail, “K-means-type algorithms: A generalized convergence theorem and characterization of local optimality,” *IEEE Transactions on Pattern Analysis and Machine Intelligence*, no. 1, pp. 81–87, 1984.
- [65] M. Ankerst, M. M. Breunig, H.-P. Kriegel, and J. Sander, “Optics: Ordering points to identify the clustering structure,” in *ACM Sigmod Record*, ACM, vol. 28, 1999, pp. 49–60.
- [66] A. Y. Ng, M. I. Jordan, and Y. Weiss, “On spectral clustering: Analysis and an algorithm,” in *Advances in Neural Information Processing Systems*, 2002, pp. 849–856.
- [67] K. He, X. Zhang, S. Ren, and J. Sun, “Deep residual learning for image recognition,” *2016 IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, vol. abs/1512.03385, 2015.
- [68] N. Tajbakhsh and K. Suzuki, “Comparing two classes of end-to-end machine-learning models in lung nodule detection and classification: MTANNs vs. CNNs,” *Pattern Recognition*, vol. 63, pp. 476–486, 2017.
- [69] L. G. Hafemann, R. Sabourin, and L. S. Oliveira, “Offline handwritten signature verification — literature review,” pp. 1–8, 2017.
- [70] M. Taşkıran and Z. G. Çam, “Offline signature identification via hog features and artificial neural networks,” in *2017 IEEE 15th International Symposium on Applied Machine Intelligence and Informatics (SAMI)*, Jan. 2017, pp. 000 083–000 086. DOI: 10.1109/SAMI.2017.7880280.
- [71] G. Rajput and P. Patil, “Writer-independent offline signature recognition based upon fourier descriptors,” *International Journal of Computer Applications*, vol. 162, no. 5, 2017.
- [72] A. B. Adeyemo and A. O. Abiodun, “Adaptive sift/surf algorithm for off-line signature recognition,” *Egyptian Computer Science Journal*, vol. 39, no. 1, 2015.
- [73] L. S. Oliveira, E. Justino, C. Freitas, and R. Sabourin, “The graphology applied to signature verification,” in *12th Conference of the International Graphonomics Society*, 2005, pp. 286–290.
- [74] H. Baltzakis and N. Papamarkos, “A new signature verification technique based on a two-stage neural network classifier,” *Engineering Applications of Artificial Intelligence*, vol. 14, no. 1, pp. 95–103, 2001.

- [75] R. Bajaj and S. Chaudhury, "Signature verification using multiple neural classifiers," *Pattern Recognition*, vol. 30, no. 1, pp. 1–7, 1997.
- [76] J. G. A. Dolfig, E. H. L. Aarts, and J. J. G. M. van Oosterhout, "On-line signature verification with hidden markov models," in *14th International Conference on Pattern Recognition (Cat. No.98EX170)*, vol. 2, Aug. 1998, 1309–1312 vol.2. DOI: 10.1109/ICPR.1998.711942.
- [77] J. Fierrez, J. Ortega-Garcia, D. Ramos, and J. Gonzalez-Rodriguez, "HMM-based on-line signature verification: Feature extraction and signature modeling," *Pattern Recognition Letters*, vol. 28, no. 16, pp. 2325–2334, 2007.
- [78] G. Agam and S. Suresh, "Warping-based offline signature recognition," *IEEE Transactions on Information Forensics and Security*, vol. 2, no. 3, pp. 430–437, 2007.
- [79] R. Zhou and C. Quek, "An automatic fuzzy neural network driven signature verification system," in *IEEE International Conference on Neural Networks*, vol. 2, 1996, pp. 1034–1039.
- [80] L. Nanni and A. Lumini, "Advanced methods for two-class problem formulation for on-line signature verification," 7-9, vol. 69, *Neurocomputing*, 2006, pp. 854–857.
- [81] S. Shalev-Shwartz, Y. Singer, N. Srebro, and A. Cotter, "Pegasos: Primal estimated sub-gradient solver for svm," *Mathematical Programming*, vol. 127, no. 1, pp. 3–30, 2011.
- [82] S. Y. Ooi, A. B. J. Teoh, Y. H. Pang, and B. Y. Hiew, "Image-based handwritten signature verification using hybrid methods of discrete radon transform, principal component analysis and probabilistic neural network," *Applied Soft Computing*, vol. 40, pp. 274–282, 2016.
- [83] L. G. Hafemann, L. S. Oliveira, and R. Sabourin, "Fixed-sized representation learning from offline handwritten signatures of different sizes," *International Journal on Document Analysis and Recognition (IJ DAR)*, pp. 1–14, 2018.
- [84] J. Ortega-Garcia, J. Fierrez-Aguilar, D. Simon, J. Gonzalez, M. Faundez-Zanuy, V. Espinosa, A. Satue, I. Hernaez, J.-J. Igarza, C. Vivaracho, *et al.*, "Mcyt baseline corpus: A bimodal biometric database," *IEE Proceedings-Vision, Image and Signal Processing*, vol. 150, no. 6, pp. 395–401, 2003.
- [85] S. N. Srihari, A. Xu, and M. K. Kalera, "Learning strategies and classification methods for off-line signature verification," in *9. International Workshop on Frontiers in Handwriting Recognition (IWFHR)*, IEEE, 2004, pp. 161–166.
- [86] M. A. Ferrer, M. D. Cabrera, C. Carmona-Duarte, and A. Morales, "A behavioral handwriting model for static and dynamic signature synthesis.," *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 39, no. 6, pp. 1041–1053, 2017.
- [87] M. A. Ferrer, J. F. Vargas, A. Morales, and A. Ordonez, "Robustness of offline signature verification based on gray level features," *IEEE Transactions on Information Forensics and Security*, vol. 7, no. 3, pp. 966–977, Jun. 2012, ISSN: 1556-6013. DOI: 10.1109/TIFS.2012.2190281.

- [88] A. B. Jagtap and R. S. Hegadi, "Eigen value based features for offline handwritten signature verification using neural network approach," in *International Conference on Recent Trends in Image Processing and Pattern Recognition*, Springer, 2016, pp. 39–48.
- [89] S. Choudhary, S. Mishra, S. Kaul, and J. Arun, "Design of handwritten signature verification using java–python platform," in *Emerging Research in Computing, Information, Communication and Applications*, Springer, 2016, pp. 75–86.
- [90] H. Khalajzadeh, M. Mansouri, and M. Teshnehlab, "Persian signature verification using convolutional neural networks," *International Journal of Engineering Research and Technology*, vol. 1, 2012.
- [91] L. G. Hafemann, R. Sabourin, and L. S. Oliveira, "Learning features for offline handwritten signature verification using deep convolutional neural networks," *Pattern Recognition*, vol. 70, pp. 163–176, 2017.
- [92] S. Dey, A. Dutta, J. I. Toledo, S. K. Ghosh, J. Lladós, and U. Pal, "Signet: Convolutional siamese network for writer independent offline signature verification," *arXiv preprint arXiv:1707.02131*, 2017.
- [93] M. B. Yilmaz, B. Yanikoglu, C. Tirkaz, and A. Kholmatov, "Offline signature verification using classifier combination of HOG and LBP features," in *2011 IEEE International Joint Conference on Biometrics (IJCB)*, 2011, pp. 1–7.
- [94] B. Erkmen, N. Kahraman, R. A. Vural, and T. Yildirim, "Conic section function neural network circuitry for offline signature recognition," *IEEE Transactions on Neural Networks*, vol. 21, no. 4, pp. 667–672, Apr. 2010, ISSN: 1045-9227. DOI: 10.1109/TNN.2010.2040751.
- [95] S. Pal, A. Alireza, U. Pal, and M. Blumenstein, "Off-line signature identification using background and foreground information," in *2011 IEEE International Conference on Digital Image Computing Techniques and Applications (DICTA)*, 2011, pp. 672–677.
- [96] N. R. Council, W. B. Committee, *et al.*, *Biometric recognition: challenges and opportunities*. National Academies Press, 2010.
- [97] A. K. Jain, K. Nandakumar, and A. Ross, "50 years of biometric research: Accomplishments, challenges, and opportunities," *Pattern Recognition Letters*, vol. 79, pp. 80–105, 2016.
- [98] J. Unar, W. C. Seng, and A. Abbasi, "A review of biometric technology along with trends and prospects," *Pattern Recognition*, vol. 47, no. 8, pp. 2673–2688, 2014.
- [99] M. A. Ferrer, M. Diaz-Cabrera, and A. Morales, "Static signature synthesis: A neuromotor inspired approach for biometrics," *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. 37, no. 3, pp. 667–680, 2015.
- [100] N. Otsu, "A threshold selection method from gray-level histograms," *IEEE Transactions on Systems, Man, and Cybernetics*, vol. 9, no. 1, pp. 62–66, Jan. 1979, ISSN: 0018-9472. DOI: 10.1109/TSMC.1979.4310076.
- [101] A. Vedaldi and K. Lenc, "Matconvnet – convolutional neural networks for matlab," in *International Conference on Multimedia (ACM)*, 2015.

- [102] K. Chatfield, K. Simonyan, A. Vedaldi, and A. Zisserman, "Return of the devil in the details: Delving deep into convolutional nets," in *British Machine Vision Conference*, 2014.
- [103] X. Glorot and Y. Bengio, "Understanding the difficulty of training deep feedforward neural networks," in *International Conference on Artificial Intelligence and Statistics (AISTATS'10)*, 2010.
- [104] N. Srivastava, G. Hinton, A. Krizhevsky, I. Sutskever, and R. Salakhutdinov, "Dropout: A simple way to prevent neural networks from overfitting," *Journal of Machine Learning Research*, vol. 15, pp. 1929–1958, 2014. [Online]. Available: <http://jmlr.org/papers/v15/srivastava14a.html>.
- [105] M. D. Prieto, G. Cirrincione, A. G. Espinosa, J. A. Ortega, and H. Henao, "Bearing fault detection by a novel condition-monitoring scheme based on statistical-time features and neural networks," *IEEE Transactions on Industrial Electronics*, vol. 60, no. 8, pp. 3398–3407, 2012.
- [106] G. Latif, D. A. Iskandar, J. M. Alghazo, and N. Mohammad, "Enhanced mr image classification using hybrid statistical and wavelets features," *IEEE Access*, vol. 7, pp. 9634–9644, 2018.
- [107] S. Ahmadi and A. S. Spanias, "Cepstrum-based pitch detection using a new statistical v/uv classification algorithm," *IEEE Transactions on Speech and Audio Processing*, vol. 7, no. 3, pp. 333–338, 1999.
- [108] A. Bansal, R. Agarwal, and R. Sharma, "Statistical feature extraction based iris recognition system," *Sādhana*, vol. 41, no. 5, pp. 507–518, 2016.
- [109] L. Debiasi and A. Uhl, "Techniques for a forensic analysis of the casia-iris v4 database," in *3rd International Workshop on Biometrics and Forensics (IWBF 2015)*, IEEE, 2015, pp. 1–6.
- [110] S. Chuai-Aree, C. Lursinsap, P. Sophasathit, and S. Siripant, "Fuzzy c-mean: A statistical feature classification of text and image segmentation method," *International Journal of Uncertainty, Fuzziness and Knowledge-Based Systems*, vol. 9, no. 06, pp. 661–671, 2001.
- [111] K. Hechenbichler and K. Schliep, "Weighted k-nearest-neighbor techniques and ordinal classification," 2004.
- [112] F. Bray, J. Ferlay, I. Soerjomataram, R. L. Siegel, L. A. Torre, and A. Jemal, "Global cancer statistics 2018: Globocan estimates of incidence and mortality worldwide for 36 cancers in 185 countries," *CA: A Cancer Journal for Clinicians*, 2018.
- [113] H. zur Hausen, "Papillomaviruses in the causation of human cancers — a brief historical account," *Virology*, vol. 384, no. 2, pp. 260–265, 2009, Small Viruses, Big Discoveries: The Interwoven Story of the Small DNA Tumor Viruses, ISSN: 0042-6822. DOI: <https://doi.org/10.1016/j.virol.2008.11.046>. [Online]. Available: <http://www.sciencedirect.com/science/article/pii/S0042682208007721>.
- [114] Y. Song, E. Tan, X. Jiang, J. Cheng, D. Ni, S. Chen, B. Lei, and T. Wang, "Accurate cervical cell segmentation from overlapping clumps in pap smear images," *IEEE Transactions on Medical Imaging*, vol. 36, no. 1, pp. 288–300, Jan. 2017, ISSN: 0278-0062. DOI: 10.1109/TMI.2016.2606380.

- [115] A. Alush, H. Greenspan, and J. Goldberger, "Automated and interactive lesion detection and segmentation in uterine cervix images," *IEEE Transactions on Medical Imaging*, vol. 29, no. 2, pp. 488–501, Feb. 2010, ISSN: 0278-0062. DOI: 10.1109/TMI.2009.2037201.
- [116] A. M. Jastreboff and T. Cymet, "Role of the human papilloma virus in the development of cervical intraepithelial neoplasia and malignancy," *Postgraduate Medical Journal*, vol. 78, no. 918, pp. 225–228, 2002, ISSN: 0032-5473. DOI: 10.1136/pmj.78.918.225. eprint: <https://pmj.bmj.com/content/78/918/225.full.pdf>. [Online]. Available: <https://pmj.bmj.com/content/78/918/225>.
- [117] A. Kobayashi, C. Miaskowski, M. Wallhagen, and K. Smith-McCune, "Recent developments in understanding the immune response to human papilloma virus infection and cervical neoplasia.," in *Oncology Nursing Forum*, vol. 27, 2000, pp. 643–51.
- [118] K. P. Braaten and M. R. Laufer, "Human papillomavirus (hpv), hpv-related disease, and the hpv vaccine," *Reviews in Obstetrics and Gynecology*, vol. 1, no. 1, p. 2, 2008.
- [119] M. H. Stoler, "Human papillomaviruses and cervical neoplasia: A model for carcinogenesis," *International Journal of Gynecological Pathology*, vol. 19, no. 1, pp. 16–28, 2000.
- [120] M. Schiffman and N. Wentzensen, "Human papillomavirus infection and the multistage carcinogenesis of cervical cancer," *Cancer Epidemiology, Biomarkers Prevention*, vol. 22, no. 4, pp. 553–560, 2013.
- [121] W. McCluggage, M. Walsh, C. Thornton, P. Hamilton, A. Date, L. Caughley, and H. Bharucha, "Inter-and intra-observer variation in the histopathological reporting of cervical squamous intraepithelial lesion using a modified Bethesda grading system," *BJOG: An International Journal of Obstetrics and Gynaecology*, vol. 105, no. 2, pp. 206–210, 1998.
- [122] L. Valentin and I. Bergelin, "Intra-and interobserver reproducibility of ultrasound measurements of cervical length and width in the second and third trimesters of pregnancy," *Ultrasound in Obstetrics and Gynecology*, vol. 20, no. 3, pp. 256–262, 2002.
- [123] A. Tareef, Y. Song, H. Huang, D. Feng, M. Chen, Y. Wang, and W. Cai, "Multi-pass fast watershed for accurate segmentation of overlapping cervical cells," *IEEE Transactions on Medical Imaging*, vol. 37, no. 9, pp. 2044–2059, Sep. 2018, ISSN: 0278-0062. DOI: 10.1109/TMI.2018.2815013.
- [124] H. Irshad, A. Veillard, L. Roux, and D. Racoceanu, "Methods for nuclei detection, segmentation, and classification in digital histopathology: A review—current status and future potential," *IEEE Reviews in Biomedical Engineering*, vol. 7, pp. 97–114, 2014, ISSN: 1937-3333. DOI: 10.1109/RBME.2013.2295804.

- [125] T. Guan, D. Zhou, and Y. Liu, "Accurate segmentation of partially overlapping cervical cells based on dynamic sparse contour searching and gvf snake model," *IEEE Journal of Biomedical and Health Informatics*, vol. 19, no. 4, pp. 1494–1504, Jul. 2015, ISSN: 2168-2194. DOI: 10.1109/JBHI.2014.2346239.
- [126] J. W. Bacus, C. W. Boone, J. V. Bacus, M. Follen, G. J. Kelloff, V. Kagan, and S. M. Lippman, "Image morphometric nuclear grading of intraepithelial neoplastic lesions with applications to cancer chemoprevention trials," *Cancer Epidemiology and Prevention Biomarkers*, vol. 8, no. 12, pp. 1087–1094, 1999.
- [127] L. Deligdisch, C. R. de Resende Miranda, H.-S. Wu, and J. Gil, "Human papillomavirus-related cervical lesions in adolescents: A histologic and morphometric study," *Gynecologic Oncology*, vol. 89, no. 1, pp. 52–59, 2003.
- [128] S. J. Keenan, J. Diamond, W. Glenn McCluggage, H. Bharucha, D. Thompson, P. H. Bartels, and P. W. Hamilton, "An automated machine vision system for the histological grading of cervical intraepithelial neoplasia (cin)," *The Journal of Pathology*, vol. 192, no. 3, pp. 351–362, 2000.
- [129] M. R. Castellanos, A. Szerszen, S. Gundry, E. C. Pirog, M. Maiman, S. Rajupet, J. P. Gomez, A. Davidov, P. R. Debata, P. Banerjee, *et al.*, "Diagnostic imaging of cervical intraepithelial neoplasia based on hematoxylin and eosin fluorescence," *Diagnostic Pathology*, vol. 10, no. 1, p. 119, 2015.
- [130] Q. Ji, J. Engel, and E. Craine, "Texture analysis for classification of cervix lesions," *IEEE Transactions on Medical Imaging*, vol. 19, no. 11, pp. 1144–1149, Nov. 2000, ISSN: 0278-0062. DOI: 10.1109/42.896790.
- [131] J. Gu, C. Y. Fu, B. K. Ng, L. B. Liu, S. K. Lim-Tan, and C. G. L. Lee, "Enhancement of early cervical cancer diagnosis with epithelial layer analysis of fluorescence lifetime images," *PloS One*, vol. 10, no. 5, e0125706, 2015.
- [132] L. Zhang, L. Lu, I. Nogues, R. M. Summers, S. Liu, and J. Yao, "Deeppap: Deep convolutional networks for cervical cell classification," *IEEE Journal of Biomedical and Health Informatics*, vol. 21, no. 6, pp. 1633–1643, Nov. 2017, ISSN: 2168-2194. DOI: 10.1109/JBHI.2017.2705583.
- [133] C. Konstandinou, D. Glotsos, S. Kostopoulos, I. Kalatzis, P. Ravazoula, G. Michail, E. Lavdas, D. Cavouras, and G. Sakellaropoulos, "Multifeature quantification of nuclear properties from images of h&e-stained biopsy material for investigating changes in nuclear structure with advancing cin grade," *Journal of Healthcare Engineering*, vol. 2018, 2018.
- [134] P. Guo, K. Banerjee, R. J. Stanley, R. Long, S. Antani, G. Thoma, R. Zuna, S. R. Frazier, R. H. Moss, and W. V. Stoecker, "Nuclei-based features for uterine cervical cancer histology image analysis with fusion-based classification," *IEEE Journal of Biomedical and Health Informatics*, vol. 20, no. 6, pp. 1595–1607, Nov. 2016, ISSN: 2168-2194. DOI: 10.1109/JBHI.2015.2483318.
- [135] T. Denkçeken, T. Şimşek, G. Erdoğan, E. Peştereli, Ş. Karaveli, D. Özel, U. Bilge, and M. Canpolat, "Elastic light single-scattering spectroscopy for the detection of cervical precancerous ex vivo," *IEEE Transactions on Biomedical Engineering*, vol. 60, no. 1, pp. 123–127, Jan. 2013, ISSN: 0018-9294. DOI: 10.1109/TBME.2012.2225429.

- [136] A. Albayrak, A. Unlu, N. Calik, A. Capar, G. Bilgin, B. U. Toreyin, B. Muezzinoglu, I. Turkmen, and L. Durak-Ata, "A whole slide image grading benchmark and tissue classification for cervical cancer precursor lesions with inter-observer variability," *arXiv preprint arXiv:1812.2521451*, 2018.
- [137] A. M. Khan, N. Rajpoot, D. Treanor, and D. Magee, "A nonlinear mapping approach to stain normalization in digital histopathology images using image-specific color deconvolution," *IEEE Transactions on Biomedical Engineering*, vol. 61, no. 6, pp. 1729–1738, 2014.
- [138] M. Macenko, M. Niethammer, J. S. Marron, D. Borland, J. T. Woosley, X. Guan, C. Schmitt, and N. E. Thomas, "A method for normalizing histology slides for quantitative analysis," in *IEEE International Symposium on Biomedical Imaging (ISBI'09)*, IEEE, 2009, pp. 1107–1110.
- [139] A. C. Ruifrok, D. A. Johnston, *et al.*, "Quantification of histochemical staining by color deconvolution," *Analytical and Quantitative Cytology and Histology*, vol. 23, no. 4, pp. 291–299, 2001.
- [140] V. Badrinarayanan, A. Kendall, and R. Cipolla, "Segnet: A deep convolutional encoder-decoder architecture for image segmentation," *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. 39, pp. 2481–2495, 2017.
- [141] L. v. d. Maaten and G. Hinton, "Visualizing data using t-sne," *Journal of Machine Learning Research*, vol. 9, no. Nov, pp. 2579–2605, 2008.
- [142] C. Molloy, C. Dunton, P. Edmonds, M. F. Cunnane, and T. Jenkins, "Evaluation of colposcopically directed cervical biopsies yielding a histologic diagnosis of cin 1, 2," *Journal of Lower Genital Tract Disease*, vol. 6, no. 2, pp. 80–83, 2002.
- [143] S. De, R. J. Stanley, C. Lu, R. Long, S. Antani, G. Thoma, and R. Zuna, "A fusion-based approach for uterine cervical cancer histology image classification," *Computerized Medical Imaging and Graphics*, vol. 37, no. 7-8, pp. 475–487, 2013.
- [144] J. W. Bacus, C. W. Boone, J. V. Bacus, M. Follen, G. J. Kelloff, V. Kagan, and S. M. Lippman, "Image morphometric nuclear grading of intraepithelial neoplastic lesions with applications to cancer chemoprevention trials," *Cancer Epidemiology and Prevention Biomarkers*, vol. 8, no. 12, pp. 1087–1094, 1999.
- [145] L. Deligdisch, C. R. de Resende Miranda, H.-S. Wu, and J. Gil, "Human papillomavirus-related cervical lesions in adolescents:: A histologic and morphometric study," *Gynecologic Oncology*, vol. 89, no. 1, pp. 52–59, 2003.
- [146] G. Agam and I. Dinstein, "Geometric separation of partially overlapping nonrigid objects applied to automatic chromosome classification," *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. 19, no. 11, pp. 1212–1222, Nov. 1997, ISSN: 0162-8828. DOI: 10.1109/34.632981.
- [147] K. Abhinav, J. S. Chauhan, and D. Sarkar, "Image segmentation of multi-shaped overlapping objects," *arXiv preprint arXiv:1711.02217*, 2017.
- [148] H. Chen, X. Qi, L. Yu, Q. Dou, J. Qin, and P.-A. Heng, "Dcan: Deep contour-aware networks for object instance segmentation from histology images," *Medical Image Analysis*, vol. 36, pp. 135–146, 2017.

- [149] S. Zafari, T. Eerola, J. Sampo, H. Kälviäinen, and H. Haario, "Segmentation of overlapping elliptical objects in silhouette images," *IEEE Transactions on Image Processing*, vol. 24, no. 12, pp. 5942–5952, 2015.
- [150] T. Teo and C. Chiu, "Pole-like road object detection from mobile lidar system using a coarse-to-fine approach," *IEEE Journal of Selected Topics in Applied Earth Observations and Remote Sensing*, vol. 8, no. 10, pp. 4805–4818, Oct. 2015, ISSN: 1939-1404.
- [151] Z. Lu, G. Carneiro, and A. P. Bradley, "An improved joint optimization of multiple level set functions for the segmentation of overlapping cervical cells," *IEEE Transactions on Image Processing*, vol. 24, no. 4, pp. 1261–1272, 2015.
- [152] X. Qi, F. Xing, D. J. Foran, and L. Yang, "Robust segmentation of overlapping cells in histopathology specimens using parallel seed detection and repulsive level set," *IEEE Transactions on Biomedical Engineering*, vol. 59, no. 3, pp. 754–765, 2011.
- [153] A. Tareef, Y. Song, H. Huang, D. Feng, M. Chen, Y. Wang, and W. Cai, "Multi-pass fast watershed for accurate segmentation of overlapping cervical cells," *IEEE Transactions on Medical Imaging*, vol. 37, no. 9, pp. 2044–2059, 2018.
- [154] C. Jung, C. Kim, S. W. Chae, and S. Oh, "Unsupervised segmentation of overlapped nuclei using bayesian classification," *IEEE Transactions on Biomedical Engineering*, vol. 57, no. 12, pp. 2825–2832, 2010.
- [155] A. Fitzgibbon, M. Pilu, and R. B. Fisher, "Direct least square fitting of ellipses," *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. 21, no. 5, pp. 476–480, 1999.
- [156] M. Winter, W. Mankowski, E. Wait, E. C. De La Hoz, A. Aguinaldo, and A. R. Cohen, "Separating touching cells using pixel replicated elliptical shape models," *IEEE Transactions on Medical Imaging*, vol. 38, no. 4, pp. 883–893, 2019.
- [157] J. C. Bezdek, R. Ehrlich, and W. Full, "Fcm: The fuzzy c-means clustering algorithm," *Computers and Geosciences*, vol. 10, no. 2-3, pp. 191–203, 1984.
- [158] S. Xu, H. Liu, and E. Song, "Marker-controlled watershed for lesion segmentation in mammograms," *Journal of Digital Imaging*, vol. 24, no. 5, pp. 754–763, 2011.
- [159] S. U. Akram, J. Kannala, L. Eklund, and J. Heikkilä, "Cell segmentation proposal network for microscopy image analysis," in *Deep Learning and Data Labeling for Medical Applications*, Springer, 2016, pp. 21–29.
- [160] J. Song, L. Xiao, and Z. Lian, "Contour-seed pairs learning-based framework for simultaneously detecting and segmenting various overlapping cells/nuclei in microscopy images," *IEEE Transactions on Image Processing*, vol. 27, no. 12, pp. 5759–5774, 2018.
- [161] M. S. Nosrati and G. Hamarneh, "Segmentation of overlapping cervical cells: A variational method with star-shape prior," in *2015 IEEE 12th International Symposium on Biomedical Imaging (ISBI)*, IEEE, 2015, pp. 186–189.

- [162] Y. Song, J.-Z. Cheng, D. Ni, S. Chen, B. Lei, and T. Wang, "Segmenting overlapping cervical cell in pap smear images," in *2016 IEEE 13th International Symposium on Biomedical Imaging (ISBI)*, IEEE, 2016, pp. 1159–1162.
- [163] M. E. Plissiti, C. Nikou, and A. Charchanti, "Automated detection of cell nuclei in pap smear images using morphological reconstruction and clustering," *IEEE Transactions on Information Technology in Biomedicine*, vol. 15, no. 2, pp. 233–241, Mar. 2011, ISSN: 1089-7771. DOI: 10.1109/TITB.2010.2087030.
- [164] J. Song, L. Xiao, and Z. Lian, "Boundary-to-marker evidence-controlled segmentation and mdl-based contour inference for overlapping nuclei," *Journal of Biomedical and Health Informatics*, vol. 21, no. 2, pp. 451–464, 2015.
- [165] D. of Signal Processing Tampere University of Technology, *Framework for simulating fluorescent microscope images with cell populations*, <http://www.cs.tut.fi/sgn/csb/simcep/index.html>, 2007.



## Tezden Üretilmiş Yayınlar

---

İletişim Bilgileri: nurullah.calik@gmail.com

### Makale

1. N. ÇALIK, O. C. KURBAN, A. R. YILMAZ, T. YILDIRIM ve L. DURAK ATA, "Large-scale offline signature recognition via deep neural networks and feature embedding", *Neurocomputing*, 359.1 (2019), 1-14.

### Konferans Bildirisi

1. N. ÇALIK, O. C. KURBAN, A. R. YILMAZ, L. DURAK ATA ve T. YILDIRIM, (2017). "Signature recognition application based on deep learning", in 25th Signal Processing and Communications Applications Conference (SIU) (pp. 1-4). IEEE.