

https://dergipark.org.tr/pub/ubsud				
Araştırma Makalesi	Cilt No:	1	Sayı No:	2
Geliş Tarihi:	24.06.2025	Kabul Tarihi:	08.08.2025	Yayın Tarihi:
				31.12.2025

DOĞAL DİL İŞLEMEDE YEREL DİL ZENGİNLİĞİ: AZBERT İLE AZERBAIJAN TÜRKÇESİ METİNLERİNİN ANLAMSAL ANALİZİ

Orkhan VALIEV ¹  Gülsüm KAYABAŞI KORU ² 

Özet

Azerbaycan Türkçesinde metin sınıflandırma ve duygu analizi görevlerini yerine getirmek amacıyla tasarlanan BERT tabanlı modelin amacı, metinleri olumlu, olumsuz veya tarafsız olarak sınıflandırmaktır. Dönüştürücü mimarisine dayanan modelin eğitimi, daha önce eğitilmiş olan BERT modeli üzerinde ince ayar yapılarak gerçekleştirilir. Bu eğitimde dilin morfolojik zenginliği ve yerel ifadeleri dikkate alınmaktadır. Modelin etkinliğini değerlendirmek için doğruluk, kesinlik, geri çağırma ve F1 puanı gibi performans ölçütleri kullanılmıştır. Değerlendirme sonuçları, modelin %96 doğruluk ve %95 F1 skoru elde ettiğini ortaya koymuştur. Özellikle pozitif ve negatif sınıflarda yüksek bir doğruluk seviyesine ulaşılmıştır; ancak yerel dil ve yapısal unsurların kullanımı nedeniyle bazı sınıflandırma hataları meydana gelmiştir. Morfolojik ve semantik dil unsurlarının dikkate alınmasının önemli bir husus olduğunu gösteren bulgular, duygu analizine önemli bir katkı sağlamaktadır. Daha büyük veri kümeleri ve yerel dile özgü modellerle yapılacak daha fazla araştırmanın performansı daha da artırma olasılığı yüksektir.

Anahtar Kelimeler: Azerbaycan Türkçesi, Duygu Analizi, Metin Sınıflandırma, BERT, Doğal Dil İşleme

NATURAL LANGUAGE PROCESSING IN AZERBAIJANI: TEXT ANALYSIS AND ARTIFICIAL INTELLIGENCE APPLICATIONS

Abstract

Classifying texts as either positive, negative, or neutral is the goal of the BERT-based model that was designed for the purpose of performing text classification and sentiment analysis tasks in Azerbaijani Turkish. The training of the model, which is based on transformer architecture, is accomplished by performing fine-tuning on the BERT model that has already been trained. This training takes into consideration the morphological richness and local expressions of the language. Performance metrics such as accuracy, precision, recall, and F1 score were used in order to assess the effectiveness of the model. The results of the evaluation revealed that the model achieved a 96% accuracy and a 95% F1 score. We were able to attain a high level of accuracy, particularly in the positive and negative classes; nevertheless, there were some classification mistakes that occurred due to the use of local language and structural elements. A significant addition to sentiment analysis is shown by the findings, which indicate that taking into account morphological and semantic language aspects is an essential aspect. There is a high probability that more research with bigger datasets and models that are particular to the local language can further enhance performance.

Keywords: Azerbaijani Turkish, Sentiment Analysis, Text Classification, BERT, Natural Language Processing

¹ Yüksek lisans öğrencisi, Ostim Teknik Üniversitesi, Yazılım Mühendisliği, Orkhan878@mail.ru, <https://orcid.org/0009-0003-5475-262X>

² Dr. Öğr. Üyesi, Ostim Teknik Üniversitesi, Yapay Zeka Mühendisliği, gulsum.koru@ostimteknik.edu.tr, <https://orcid.org/0000-0002-1749-900X>

1. GİRİŞ

Doğal Dil İşleme (NLP), bilgisayarların insan dilini anlama, yorumlama ve üretme sürecine odaklanan bir bilgisayar bilimi alt alanıdır. Temel bir bakış açısıyla, bu teknoloji bilgisayarlara dilsel verileri değerlendirme, anlamlı sonuçlara varma ve insanlarla doğal dil kullanarak iletişim kurma yeteneği kazandırır. İnsanlar, dilin yapısal ve anlamsal düzeylerindeki sorunları ele alan doğal dil işleme sayesinde bilgisayarlarla daha doğal bir şekilde konuşabilirler. Çeşitli uygulamalarda, bu yaklaşım doğru dil anlayışının yanı sıra insanlarınkine daha yakın yanıtlar geliştirilmesine de olanak tanır. Dijital asistanlar (Siri ve Alexa gibi), çeviri sistemleri (Google Translate gibi), metin analiz araçları, müşteri hizmetleri sohbet robotları ve sosyal medya analitiği gibi uygulamalar, doğal dil işleme teknolojisinden kapsamlı bir şekilde yararlanan mevcut uygulamalara örnektir. Dilin karmaşıklığı ve bağlamı hakkında bilgi sahibi olmak, metin verilerinin etkili bir şekilde işlenmesini ve kategorize edilmesini sağlayan doğal dil işlemenin gücünün anahtarıdır (Şahin, 2021).

Doğal dil işlemenin geniş uygulama yelpazesi, insanların ve şirketlerin büyük miktarda veriyle başa çıkma biçimlerini tamamen değiştirmiştir. İçerik analizi, duygu analizi, bilgi çıkarımı ve metin özetleme gibi faaliyetler müşterilere önemli bilgiler sağlamakla kalmaz, aynı zamanda otomasyonla ilgili süreçleri de hızlandırır. Bu disiplinin etkisi, yalnızca dil işlemeyle değil, aynı zamanda daha kapsamlı bir yapay zekâ ekosistemiyle birleştirildiğinde, hem daha akıllı hem de daha verimli sistemler tasarlamayı mümkün kılar.

Bununla birlikte, Azerice gibi daha az çalışılan dillerde doğal dil işleme uygulamalarının benimsenmesi çeşitli engellerle ve engellerle karşılaşmaktadır. Türk dil ailesinin bir üyesi olmasına rağmen, Azerice'nin kendine özgü yapısal özellikleri ve dil bilgisi kuralları vardır. Dil işleme, bunun sonucunda birçok yeni ve ilginç sorunla karşı karşıyadır. Örneğin, Azerice morfolojik olarak zengin bir kelime dağarcığına sahiptir; bu da kelimelerinin daha karmaşık yapılara sahip olduğu ve bu da analiz edilmelerini ve kategorilere ayrılmasını kolaylaştırdığı anlamına gelir. Meseleyi daha da karmaşık hale getiren şey, kelime sonu morfeplerinin, dili ilgili bağlamında doğru bir şekilde değerlendirmeyi zorlaştıran önemli bir unsur olmasıdır. Bu durum, metinlerin doğru bir şekilde sınıflandırılmasını engelleyebilir ve terimler arasında anlamsal değişikliklere yol açabilir.

Azerice yazılmış metinleri etkili bir şekilde değerlendirmek ve yorumlamak, bu dildeki dilsel farklılıklar nedeniyle zordur. Özellikle, morfolojik analiz, bağlam bağımlılığı ve kelime türetme gibi dilin temel yönlerini doğru bir şekilde analiz etmek için Azerice metinlere yönelik daha gelişmiş yöntemlere ihtiyaç vardır. Sonuç olarak, Azericeye özgü doğal dil işleme araçlarının, bu dilin birçok kültürel ve dilsel yönünü dikkate alacak şekilde oluşturulması gerekmektedir. Bu bağlamda, Azericedeki doğal dil işleme çalışmalarının eksikliklerinin giderilmesi, işlenebilirliğini artıracak ve bu dildeki yapay zekâ uygulamalarının verimliliğinin

artmasına katkıda bulunacaktır. Dahası, bu, dil üzerinde araştırma yapılarak gerçekleştirilecektir.

Azerbaycan dili bir dizi engel barındırmaktadır ve bu araştırmanın temel amaçlarından biri, bu sorunlara çözümler sunmak ve dilin doğal dil işleme prosedürlerinde daha verimli kullanılabileceği yolları keşfetmektir. Bu bağlamda, Azerbaycan diline özgü morfolojik özellikleri dikkate alarak daha doğru ve verimli bir metin kategorizasyon görevleri modeli oluşturmak amaçlanmaktadır. Bu model, yalnızca Azerbaycan diline özgü belirli bir kullanım için değil, aynı zamanda ilgili dillerde de kullanılmak üzere geliştirilmiştir. Sonuç olarak, bu konunun önemi, yalnızca Azerbaycan dilinde doğal dil işleme araştırmalarının geliştirilmesine değil, aynı zamanda küresel ölçekte dil teknolojilerinin geliştirilmesine de katkıda bulunacaktır.

Makalenin 2. bölümünde çalışmayı ilgilendiren konularda yapılan literatür incelemesi, 3. bölümünde çalışmanın yönteminden, 4. bölümünde bulgular ve yorumlardan, 5. bölümde sonuç ve tartışmalara yer verilmiştir.

1.1. Problem Durumu

Doğal dil işleme, dünya genelinde yaygın olarak geliştirilen ve kullanılan bir teknoloji alanı olsa da bu teknolojilerin düşük kaynaklı diller üzerindeki etkisi ve uygulanabilirliği sınırlı kalmaktadır. Azerbaycan Türkçesi gibi görece az temsil edilen diller, morfolojik karmaşıklıkları ve yapısal özellikleri nedeniyle mevcut NLP araçlarıyla etkili bir biçimde işlenememektedir. Özellikle sözcük türetme biçimleri, ek yapıları ve bağlam içindeki anlam kaymaları gibi unsurlar, metinlerin doğru analiz edilmesini güçleştirmektedir. Literatürde, bu dile özgü derin öğrenme modelleri ve dil kaynaklarının eksikliği, dilin dijital ortamlarda etkin kullanımını sınırlamakta ve bu alanda yapılan çalışmaların gelişimini yavaşlatmaktadır. Bu durum, hem dilin teknolojik olarak temsil edilmesini engellemekte hem de yerel kullanıcılar için özelleştirilmiş akıllı sistemlerin eksikliğine neden olmaktadır.

1.2. Amaçlar

Bu çalışmanın temel amacı, Azerbaycan Türkçesi özelinde doğal dil işleme süreçlerinde karşılaşılan dilsel zorlukları yapay zekâ destekli yaklaşımlar ile aşmak ve bu dile özgü etkili bir metin sınıflandırma modeli geliştirmektir. Bu kapsamda, morfolojik yapıların, sözcük bağlamlarının ve dil bilgisel unsurların daha doğru temsil edilebileceği bir sistem tasarlanması hedeflenmektedir. Aynı zamanda geliştirilecek modelin yalnızca Azerbaycan Türkçesi için değil, benzer dil yapısına sahip diğer düşük kaynaklı diller için de örnek teşkil etmesi amaçlanmaktadır. Böylece, hem yerel dil teknolojilerinin gelişmesine katkı sağlanması hem de çok dilli NLP sistemlerinin evrensel erişilebilirliğinin artırılması hedeflenmektedir.

2. LİTERATÜR TARAMASI

Yapay zeka (YZ) ve doğal dil işleme alanındaki gelişmeler, çeviri teknolojisini kökten dönüştürdü. İlkel kelime tabanlı algoritmalara dayanan geleneksel sistemlerin aksine, YZ odaklı araçlar daha anlamlı çeviriler sunmak için karmaşık bağlamsal anlayıştan yararlanır. 2006 yılında kurulan Google Translate, makine çevirisi alanını şekillendirmede etkili olmuştur. Başlangıçta basit bir kelime tabanlı yaklaşım kullanan Google Translate, o zamandan beri derin öğrenme metodolojilerini ve Transformer modellerini benimsemiştir. Bu evrim, çeviri doğruluğunda önemli iyileştirmelerle sonuçlanmıştır.

Sinir ağları ve derin öğrenme teknikleri, bu araçların dilin bağlamsal ve anlamsal nüanslarını ayırt etmesini sağlayarak, salt sözcüksel çevirilerin ötesine geçer. Sonuç olarak, Google Translate gibi araçlar yalnızca kelime anlamlarını değil, aynı zamanda dilbilgisi yapılarını ve genel dilsel bağlamı da doğru bir şekilde yorumlayabilir ve bu da son derece doğru ve bağlamsal olarak alakalı çevirilere yol açar (Radford, 2018).

Transformer mimarisi, özellikle dikkat mekanizması, özellikle uzun metinlerde bağlamsal nüansları anlamak için çok önemlidir. Bu mekanizma, kelime ilişkilerini etkili bir şekilde analiz ederek gelişmiş anlayışa yol açar. BERT (Bidirectional Encoder Representations from Transformers) ve GPT (Generative Pre-trained Transformer) gibi Transformer tabanlı modellerin kullanılması, çeviri doğruluğunu önemli ölçüde artırır. Bu modellerin sürekli eğitimi, performanslarını daha da iyileştirir ve zamanla giderek daha hassas ve anlamlı çeviriler elde edilmesini sağlar (Devlin, 2019).

GPT, gözetimsiz öğrenme teknikleri aracılığıyla dil modelleme ve metin oluşturmada dikkate değer bir ilerleme göstermiştir. Aynı zamanda, Amazon Translate gibi alternatif çeviri araçları rekabetçi seçenekler olarak ortaya çıkmaktadır. NMT'den (Nöral makine çevirisi) yararlanan Amazon Translate, hızlı ve doğru çeviriler sunar. Özellikle teknik metinleri işlerken olağanüstü bir doğruluk sergiler. Ayrıca, hız avantajı özellikle büyük veri kümelerini işlerken ticari uygulamalar için faydalıdır.

Duygu analizi, metinsel verilerde iletilen duygusal tonu ayırt etmeyi içerir. Bu teknik, pazarlama araştırması, müşteri geri bildirim değerlendirmesi ve sosyal medya trend analizi dahil olmak üzere çeşitli alanlarda kapsamlı uygulama bulur.

Tipik olarak, duygu analizi metni üç temel duygu sınıfına ayırır: olumlu, olumsuz ve nötr. Ancak, salt sözcüksel anlayışın ötesine geçer; ayrıca dilin bağlamsal nüanslarını ve dilsel yapısını da dikkate alır. Sonuç olarak, doğru duygu analizi sonuçlarına ulaşmak için dilin derinlemesine anlaşılması çok önemlidir (Agarwal, 2018).

Duygu analizi, metinsel verilerde ifade edilen duygusal tonu belirlemek için çeşitli metodolojiler kullanır.

Sözlük tabanlı yöntemler, tek tek sözcüklere veya ifadelere duygu puanları (olumlu, olumsuz veya nötr) atayan önceden derlenmiş sözlüklerden yararlanır. Bir metnin toplam duygusu, bu sözcük düzeyindeki puanların toplanmasıyla hesaplanır. Basit ve hesaplama açısından verimli olsa da, bu yöntemler alaycılık ve bağlama bağlı anlam gibi nüanslı dil ile mücadele eder.

Makine öğrenimi tabanlı yöntemler, etiketli metinlerin kapsamlı veri kümeleri üzerinde eğitilmiş sınıflandırıcı modelleri kullanır - insan açıklayıcıların duygu etiketlerini açıkça atadığı örnekler. Bu alanda kullanılan popüler algoritmalar arasında Destek Vektör Makineleri (SVM), Naive Bayes, Karar Ağaçları ve Rastgele Orman bulunur. Bu teknikler genellikle sözlük tabanlı yaklaşımlardan daha yüksek doğruluk elde eder ancak etkili eğitim için önemli miktarda etiketli veri gerektirir (Sennrich, 2016).

Derin öğrenme tabanlı yöntemler, özellikle Tekrarlayan Sinir Ağları (RNN'ler), Uzun Kısa Süreli Bellek (LSTM) ağları ve Evrişimli Sinir Ağları (CNN'ler) kullananlar, son zamanlarda duygu analizinde önemli ilerlemeler gösterdi. Metin içindeki uzun menzilli bağımlılıkları yakalama yetenekleri, hem sözdizimsel hem de anlamsal özellikleri çıkarmalarına olanak tanır ve bu da karmaşık duygusal ifadeleri ayırt etmede gelişmiş performansa yol açar (Zhang, 2015).

BERT ile örneklendirilen dönüştürücü tabanlı modeller, duygu analizinde önemli ilerlemeler göstermiştir. Benzersiz yetenekleri, metinsel verilerin çift yönlü işlenmesinde yatmaktadır ve bu da bir cümle içindeki hem önceki hem de sonraki öğeleri dikkate alarak sözcüklerin bağlamsal anlamını yakalamalarını sağlamaktadır. Bağlamın bu kapsamlı anlaşılması, BERT'in çevreleyen sözcüklerin her bir sözcük üzerindeki nüanslı etkisini ayırt etmesini sağlayarak, metinsel duygunun daha kesin bir şekilde yorumlanmasını kolaylaştırır ve bu alanda önemli bir avantaj sağlar (Devlin, 2019).

GPT gibi büyük dil modelleri, duygu analizinde yeterlilik gösterir. Kapsamlı metinsel veri kümeleri üzerinde aldıkları eğitim, ardışık kelimeleri tahmin etmelerini sağlar ve böylece metin içindeki bağlamsal ilişkilerin anlaşılmasını ve ifade edilen duyguların tanımlanmasını kolaylaştırır (Devlin, 2019).

Bununla birlikte, GPT modelleri, özellikle sosyal medya söyleminde ve anlamda hızlı değişimlerle karakterize edilen bağlamlarda karmaşık duygusal ifadeleri ve ince nüansları çözmede sınırlamalarla karşılaşabilir. Örneğin, ironi veya alaycılıkla dolu cümleleri yanlış yorumlayabilirler. GPT-3 gibi gelişmiş yinelemeler, yaratıcı dil kullanımını anlamada gelişmiş yetenekler sergilerken, nüanslı alaycılık veya çok katmanlı ironi gibi daha derin bağlamsal anlayış gerektiren metinleri analiz ederken zorluklar devam eder.

Bu, sağdan sola doğru ilerleyen ve uzun metinlerde veya belirli bağlamsal ortamlarda bulunan tüm incelikleri yakalayamayan GPT'nin tek yönlü öğrenme yaklaşımına atfedilebilir.

Azerbaycan Türkçesi, Türk dil ailesiyle aynı köklere sahip olsa da, kendine özgü morfolojik yapıları ve dilbilgisi nüansları NLP çalışmaları için zorluklar oluşturmaktadır. Yakın zamanda yapılan bir çalışma, dilin belirli özelliklerine göre uyarlanmış BERT tabanlı bir mimari kullanarak Azerbaycan Türkçesi için bir duygu analizi modeli geliştirerek bu boşluğu gidermeyi amaçlamıştır. Modelin başlangıçtaki doğruluğu mütevazı olsa da, bağlamsal anlayışta iyileştirme potansiyeli göstermiştir (Gültekin, 2019).

Azerbaycan Türkçesinin benzersiz morfolojik özellikleri, standart Türkçe ile benzerliklerine rağmen duygu analizinin karmaşıklığına katkıda bulunur. Sonuç olarak, veri kümelerini genişletmek ve Azerbaycan Türkçesi için özelleştirilmiş NLP modelleri geliştirmek, dil işleme yeteneklerini geliştirmeye yönelik önemli adımlardır. Bu araştırma, Azerbaycan diline özgü gelecekteki NLP uygulamaları için temel oluşturmaktadır (Özyurt, 2022).

Bu kapsamlı literatür incelemesi, yapay zeka destekli çeviri araçlarının evrimini inceliyor ve bunların başlangıcından çok dilli doğal dil işleme içindeki mevcut paradigma değişimlerine kadar olan yörüngesini izliyor. Özellikle, BERT ve GPT gibi dönüştürücü tabanlı derin öğrenme modelleri, yalnızca bağlamı kavrama kapasiteleri değil, aynı zamanda hedef dillere özgü nüanslı sözdizimsel ve anlamsal özellikleri yakalama kapasiteleri nedeniyle son yıllarda önemli oyuncular olarak ortaya çıktı.

Dahası, duygu analizi tekniklerindeki gelişmeler, metinsel verilerdeki metin altı anlamının, kültürel ipuçlarının ve bağlamsal duyguların gelişmiş doğrulukla tespit edilmesini sağlıyor. Bu ilerleme, hem akademik araştırmalarda hem de pratik uygulamalarda önemli yeniliklerin önünü açıyor. Ancak inceleme kritik bir boşluğu vurguluyor: Azerbaycan Türkçesi gibi düşük kaynaklı dillerin yetersiz dijital temsili.

Bu eşitsizliğin ele alınması, bu az temsil edilen dillerin dilsel çeşitliliğini ve kültürel bağlamını hesaba katan özelleştirilmiş doğal dil işleme çözümlerinin geliştirilmesini gerektiriyor. Bu, dil teknolojileri alanında kapsayıcılığı sağlamak için çok önemlidir. Bu araştırmanın bulguları, düşük kaynaklı bu dillerin hem teknik hem de sosyo-kültürel açıdan dijitalleştirilmesine katkıda bulunmayı ve aynı zamanda bu alanda yeni modeller geliştirmek için sağlam bir bilimsel temel oluşturmayı amaçlamaktadır.

3. YÖNTEM

3.1. Kullanılan Veri Setleri

3.1.1. Veri Seti Seçimi

Bu çalışmada, Azerbaycan Türkçesinde doğal dil işleme uygulamaları için kullanılacak veri seti dikkatle seçilmiştir. Veri seti; sosyal medya, haber metinleri ve forum yorumları gibi farklı kaynaklardan alınan çeşitli metin türlerini içermektedir. Hem kısa hem de uzun metinlerin yer aldığı bu içerikler, modelin farklı

bağlamlardaki başarımını değerlendirmeye olanak tanır. Ayrıca, her metne olumlu, olumsuz veya nötr olacak şekilde duygu etiketleri atanmıştır. Bu etiketleme, duygu analizinin doğruluğunu ölçmek ve çeviri süreçlerinde duygusal anlamın korunup korunmadığını değerlendirmek açısından önemlidir. Veri setinin titizlikle seçilmesi, çalışmanın güvenilirliğini ve geçerliliğini doğrudan desteklemektedir.

3.1.2. Veri Seti Temizleme Süreçleri

Veri temizleme, doğal dil işleme modellerinin doğruluğu ve verimliliği için kritik bir adımdır. Ham veriler genellikle gürültü, hatalar ve anlamsız içerikler içerdiğinden, bu verilerin dikkatli bir şekilde düzenlenmesi gereklidir. Bu süreç, modelin daha sağlıklı öğrenmesini ve güvenilir sonuçlar üretmesini sağlar (Bojanowski, 2017). Bu bölümde, veri temizleme süreci ayrıntılı olarak ele alınacaktır.

3.1.2.1 Durak Kelimelerin (Stopwords) Çıkarılması

Doğal dil işleme sürecinde, anlamsal açıdan düşük katkıya sahip olan durdurma sözcükleri (örneğin "ve", "bu", "ile") model performansını olumsuz etkileyebileceğinden genellikle verilerden çıkarılır. Bu işlem, özellikle Azerbaycan Türkçesi için modelin metni daha iyi anlamasına yardımcı olur. Ayrıca, metin uzunluğu da model eğitimi açısından kritik bir faktördür. Çok kısa veya çok uzun metinler modelin doğruluğunu azaltabileceği için, metinler 128 belirteçlik sabit bir uzunlukta standartlaştırılır. Bu sayede modelin tutarlı şekilde eğitilmesi sağlanır ve işlem verimliliği artırılır. Şekil 3.1 Azerbaycan dilindeki StopWords fonksiyonunu göstermiştir.

```
1 # 2. NLTK İndirmeleri
2 nltk.download('punkt')
3 nltk.download('stopwords')
4 nltk.download('punkt_tab')
5
6 # 3. Azerbaycani Stopwords
7 azerbaijani_stopwords = set([
8     've', 'bir', 'bu', 'ile', 'üçün', 'də', 'ki', 'onun', 'da', 'heç', 'amma', 'daha', 'cox', 'az', 'ham', 'bütün',
9     'kimi', 'har', 'o', 'ona', 'olaraq', 'bele', 'ile', 'eda', 'bu', 'müxtəlif', 'sən', 'mən', 'biz', 'siz', 'onlar'
10 ])
11
```

Şekil 3.1. Azerbaycan dilindeki Stopwords'ler

Şekil 3.1. Azerbaycan dilindeki StopWords'lerin kodumuza uyarlanmasını gösteriyor. Bu fonksiyon veri temizleme aşamasından olub model eğitimi aşırı yüklenmeden kurtarmaktadır.

3.1.2.2 Özel Karakterlerin ve Gereksiz Verilerin Temizlenmesi

Veri temizleme sürecinin ilk aşamasında, özel karakterler, sayılar, noktalama işaretleri ve gereksiz boşluklar metinlerden çıkarılmıştır. Bu öğeler genellikle anlam

taşımadıkları için modelin öğrenme sürecini olumsuz etkileyebilir. Sayılar, semboller (örn. @, #, %, vb.) ve fazla boşlukların temizlenmesiyle, modelin yalnızca anlamlı ve dilsel açıdan değerli içerikler üzerinden eğitilmesi sağlanmıştır. Bu işlem, analiz ve sınıflandırma doğruluğunu artırmayı hedeflemektedir.

3.1.2.3 Tokenizasyon

Tokenizasyon, metnin küçük parçalara (kelimelere veya alt kelimelere) ayrılması işlemidir. Herhangi bir metin, doğal dil işleme sistemlerinde belirli dil birimlerine bölünmeden işlenemez. Tokenizasyon, bu dil birimlerini çıkararak modelin metni daha kolay anlamasına olanak tanır (Mikolov, 2013). Bu adımda yapılan işlemler şunlardır:

- **Kelimelere Ayırma (Word Tokenization):** Metinler, anlamlı kelimelere ayrılmıştır. Örneğin, “Bugün hava çox gözəldir” cümlesi şu şekilde tokenize edilmiştir: [“Bugün”, “hava”, “çox”, “gözəldir”].
- **Alt Kelimelere Ayırma (Subword Tokenization):** Türkçedeki eklemeli yapı nedeniyle bazı kelimeler birden fazla alt kelimeye ayrılabilir. Örneğin, "kitapçı" kelimesi "kitap" ve "-çı" alt kelimelerine ayrılabilir (Hakkani-Tür, 2017). Tokenizasyonun amacı, her kelimeyi veya alt kelimeyi modelin anlayabileceği bir sayısal formata dönüştürmektir. Bu aşama, modelin doğru öğrenmesini ve sınıflandırma görevini etkili bir şekilde yerine getirmesini sağlar.

3.1.3 Veri Seti İçeriği

Bir veri setinin kalitesi ve bileşimi, doğal dil işleme modellerinin performansının önemli belirleyicileridir. Bu, özellikle sınırlı kaynaklara sahip dillerdeki duygu analizi ve metin sınıflandırma görevleri için geçerlidir; burada modelin karmaşık dil yapılarını kavraması için veri çeşitliliği esastır.

Bu çalışma, duygu analizi ve metin sınıflandırmasının etkinliğini değerlendirmek için titizlikle düzenlenmiş Azerbaycan Türkçesi metinlerinden oluşan bir veri setini kullanır. Veri seti, sosyal medya, haber kaynakları, bloglar ve kullanıcı forumları gibi çeşitli platformlardan kaynaklanan metinlerin dengeli bir temsili içerir. Her metin, içerik türüne göre kategorize edilir ve bu da modelin anlamsal varyasyonları ve bağlamsal nüansları ele alma becerisi açısından titizlikle değerlendirilmesini sağlar. Bu çok yönlü veri yaklaşımı, çeşitli duygusal tonları, ifadeleri ve kültürel çağrışımları tanıma kapasitesini artırarak, özellikle çeviri doğruluğu ve duygu tespiti gibi alanlarda modelin genel performansını önemli ölçüde artırır.

Ayrıca, bu veri odaklı metodoloji yalnızca teknik doğruluğu değil aynı zamanda etik ve kültürel hususları da ön planda tutarak, kaynakları yetersiz dillerin daha kapsamlı

ve kapsayıcı bir dijital temsiline katkıda bulunmaktadır. Şekil 3.2 modelimiz için kullandığımız veri setinin içeriğini göstermektedir.

text	label	clean_text
Yemək tamamilə xarab idi.	negative	yemək tamamilə xarab idi
Bu film çox maraqlıydı!	positive	film maraqlıydı
Bu gün çox gözəl idi!	positive	gün gözəl idi
Çox yorulдум, heç bir şey etmədim.	negative	yorulдум ey etmədim
Əla bir tətıl keçirdim!	positive	əla tətıl keçirdim
Çox yorulдум, heç bir şey etmədim.	negative	yorulдум ey etmədim
Çox gözəl bir gün keçirdim!	positive	gözəl gün keçirdim
Bu gün çox pis keçdi.	negative	gün pis keçdi
Məktəbdə hər şey adi idi.	neutral	məktəbdə ey adi idi
Bu film çox maraqlıydı!	positive	film maraqlıydı

Şekil 3.2. Veri seti örneği

Şekil 3.2 kullandığımız veri setinin içeriğini göstermektedir. Bu veri setinin içeriğindeki metinler pozitif, negatif ve nötr duygular içermektedir. Modelimiz duygu analizi yaparken bu cümleleri öğrenerek sonuçlar çıkarmaktadır.

3.1.3.1 Sosyal Medya Paylaşımları

Sosyal medya platformları, bireysel ifadelerin yoğun olarak yer aldığı kısa ve duygusal içerikler sunduğundan, duygu analizi ve metin sınıflandırma çalışmaları için zengin bir veri kaynağıdır. Bu çalışmada, Azərbaycan Türkçesine özgü deyimsele ve gündelik ifadeleri içeren sosyal medya gönderilerinden oluşan çeşitli bir veri seti kullanılmıştır.

Bu veriler, modelin yerel dil varyasyonlarını, bağlamsal ipuçlarını, ironi ve alay gibi dilsel incelikleri tanımasına yardımcı olur. Sosyal medya içeriklerinin gayri resmi yapısı, bağlam duyarlı anlayışın gelişmesine katkı sağlar. Bu, özellikle günlük dilin standart biçimlerden saptığı dillerde, modelin gerçek dünya verileriyle daha doğru ve duyarlı öğrenmesini mümkün kılar. Verinin özellikleri:

- Kısa Metinler: Sosyal medya gönderileri genellikle kısa ve doğrudan ifadeler içerir.
- Duygusal İfadeler: Kullanıcıların duygusal tepkileri, modelin pozitif, negatif ve nötr sınıflar arasında ayırım yapmasını destekler.

- Etkileşimli Dil: Paylaşım, yorum ve beğeni gibi unsurlar metne bağlamsal zenginlik katar. Paylaşımlara örnek verecek olursak:

- “Bugün həqiqətən gözəl bir gündür! Hər şey yolunda gedir 😊” → Pozitif duygu yansıtır.
 - “Bu xəbər çox üzücüdür, heç gözləməirdim 😞” → Negatif duygu yansıtır.
- Sonuç olarak, sosyal medya verilerinin kullanımı, yalnızca veri setini çeşitlendirmekle kalmaz, aynı zamanda modeli daha hassas, bağlama duyarlı ve gerçek dünyaya uygun hale getirir.

3.1.3.2 Haber Metinleri

Haber makaleleri, yapay zekâ destekli doğal dil işleme (NLP) modellerinin değerlendirilmesi için önemli bir kaynak sunar. Resmi dil yapıları, sözcüksel doğruluk ve karmaşık cümle yapıları, modelin dilsel yeterliliğini ve bağlamsal anlama becerisini ölçmek açısından değerlidir. Açık duygusal ifadeler içermeseler de, örtük anlamlar ve kelime seçimleri ile ince duygusal tonlar barındırabilirler. Bu özellikler, modelin sadece açık duygu sinyallerini değil, aynı zamanda gizli duygu katmanlarını da algılama kapasitesini test etmek için önemlidir. Öne çıkan özellikler:

- Formal Dil: Modelin resmi yapılarla çalışma ve dilsel doğruluk üretme becerisi ölçülür.
- Uzun Metinler: Karmaşık cümle yapılarını analiz edebilme yeteneği gelişir. - Bilgi Yoğunluğu: Teknik ve özel içerikler sayesinde bağlamsal analiz gücü artar. Haber metinlerine örnek verecek olursak:

- “Azərbaycan iqtisadiyyatı son illərdə böyük bir inkişaf əldə edib. Hökumət, yatırımcıları ölkəyə cəkmək üçün yeni politikalər irəli sürür.” → Nötr içerik
- “Azərbaycanda tıbb sektorunda olunan reformlarla birləktə, xəstəxanalar daha modern hala gəlmiş və sağlıq xidmətləri daha əlçatan olmuşdur.” → Nötr içerik

3.1.3.3 Forum Yorumları

Forum yorumları, kullanıcıların kişisel deneyim, düşünce ve duygularını ifade ettiği metinler olarak, doğal dil işleme (NLP) çalışmaları için zengin ve çok yönlü bir veri kaynağı sunar. Duygusal derinlik, bağlamsal çeşitlilik ve gayri resmi dil yapısı, bu yorumları özellikle duygu analizi ve metin sınıflandırma için değerli kılar. Bu çalışmada kullanılan forum verileri; sağlık, eğitim, teknoloji, siyaset ve günlük yaşam gibi farklı tematik alanlardan toplanmıştır. Sosyal medya içeriklerine benzer şekilde argo, mizah ve gündelik ifadeler içerse de, forum metinleri genellikle daha uzun ve bağlama özgü bilgiler barındırır. Bu durum, modelin örtük anlamları, bağlamsal ilişkileri ve kültürel nüansları ayırt etme becerisini geliştirmeye katkı sağlar. Öne çıkan özellikleri:

- İnfomal Dil: Modelin gayri resmi dilde duygu tespiti yapma becerisi test edilir.

- Tematik Çeşitlilik: Farklı konulardaki içerikler, modelin çok yönlü anlam analizine olanak tanır.
- Duygusal Yük: Kişisel deneyimlere dayalı yorumlar, güçlü duygusal içerikler sunar. Örnek Forum Yorumları:
 - “Bugün yaşadığım duyğu həqiqətən çox olumsuzdu. Xidmət çox pis idi!” → Negatif duygu
 - “Çok güzel bir tətıl keçirdim, hər şey mükəmməldi!” → Pozitif duygu

Sonuç olarak, sosyal medya, haber metinleri ve forum yorumlarından oluşan bu çeşitli veri seti, modelin pozitif, negatif ve nötr duygu sınıflarını doğru şekilde öğrenmesini ve sınıflandırmasını sağlamaya yöneliktir. Forum yorumlarının eklenmesi, modelin dilsel esneklik, duygusal farkındalık ve bağlam duyarlılığı kazanmasına önemli katkılar sunar (Huseynov, 2019).

3.1.4 Veri Seti Büyüklüğü

Doğal Dil İşleme (NLP) ve Yapay Zekâ modellerinin başarısı, kullanılan veri kümesinin boyutu ve kalitesine doğrudan bağlıdır. Özellikle sınırlı kaynaklı diller için büyük ve dengeli veri setleri, modellerin yalnızca mevcut örnekleri öğrenmesini değil, aynı zamanda nadir ve alışılmadık dil yapılarında da genelleme yapabilmesini sağlar.

Bu çalışmada kullanılan Azərbaycan Türkçesi veri kümesi, metin sınıflandırma ve duygu analizi görevlerine uygun olacak şekilde özenle hazırlanmış ve ön işleme sürecinden geçirilmiştir. Veri seti; sosyal medya gönderileri, haber metinleri, forum yorumları ve blog içerikleri gibi farklı kaynaklardan gelen çeşitli metin türlerini içermektedir. Bu çeşitlilik, modelin hem yüzeysel dil kalıplarını hem de bağlamsal ilişkileri etkili bir şekilde öğrenmesine imkân tanımıştır. Şekil 3.3 kullanılan veri setinin içerik boyutunu göstermektedir.

	text	label	clean_text
0	Əla bir tətıl keçirdim!	positive	əla tətıl keçirdim
1	Bu gün çox gözəl idi!	positive	gün gözəl idi
2	Bu gün çox gözəl idi!	positive	gün gözəl idi
3	Əla bir tətıl keçirdim!	positive	əla tətıl keçirdim
4	Bu film çox maraqlıydı!	positive	film maraqlıydı
...	...	...	...
1908895	Məktəbdə hər şey adi idi.	neutral	məktəbdə ey adi idi
1908896	Kitabı oxudum, amma çox maraqlı deyildi.	neutral	kitabı oxudum maraqlı deyildi
1908897	Yemək vaxtında gəldi, amma heç bir şey xüsusi ...	neutral	yemək vaxtında gəldi ey xüsusi deyildi
1908898	Məktəbdə hər şey adi idi.	neutral	məktəbdə ey adi idi

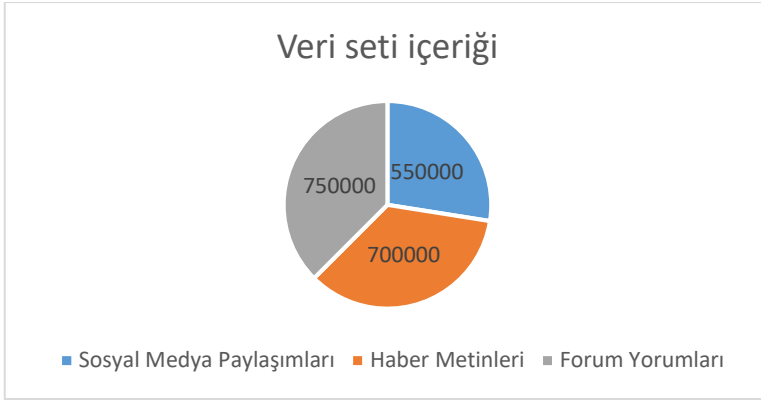
Şekil 3.3. Veri Seti Büyüklüğü

Şekil 3.3. veri setinin içeriğinin büyüklüğünü görselleştiriyor. Data setin büyük olmasının avantajı modelin daha iyi performans alabilmesi ve doğruluğu göstermesi

içindir. Küçük veri setleriyle model düşük sonuçlar vermekte veya fazla öğrenmeye maruz kalmaktadır.

3.1.4.1 Veri Seti Boyutunun Belirlenmesi

Çalışmada kullanılan veri seti, Azerbaycan Türkçesinde yazılmış sosyal medya paylaşımları, haber metinleri ve forum yorumlarından oluşmaktadır. Veri seti, duygu analizi (pozitif, negatif, nötr) ve çeviri gibi görevler için hazırlanmış olup, her kategori dengeli olacak şekilde etiketlenmiştir. Şekil 3.4 kullanılan veri setinin içeriğini oluşturan metinlerin nasıl bölündüğünü göstermektedir.



Şekil 3.4 Veri Seti İçeriği

Şekil 3.4 Veri setinin içerisinde bulunan verilerin bölüşümünü görselleştirmesidir. Toplamda:

- Sosyal medya paylaşımları: 550.000
- Haber metinleri: 700.000
- Forum yorumları: 750.000

Veri setinin %80'i eğitim, %20'si test için ayrılmıştır. Veri çeşitliliği ve dengeli dağılım, modelin farklı metin türlerinde yüksek doğrulukla sınıflandırma yapmasını desteklemektedir.

3.1.5 Veri Seti Kullanımı

Bu çalışma, Azerbaycan Türkçesi metinlerinden oluşan büyük ve dengeli bir veri seti kullanarak duygu analizi ve metin sınıflandırma modeli geliştirmeyi amaçlamaktadır. Veri seti; sosyal medya gönderileri, haber metinleri ve forum yorumları gibi farklı kaynaklardan derlenmiş ve her metin pozitif, negatif veya nötr etiketleriyle sınıflandırılmıştır.

Veri Setinin Hazırlanması ve Bölünmesi

- Veri seti, %80 eğitim, %20 test olarak ikiye ayrılmıştır.

- Her duygu sınıfından eşit sayıda örnek seçilerek sınıf dengesi sağlanmıştır.
- Bu yaklaşım, modelin genelleme yapabilme yeteneğini artırmayı ve önyargıları azaltmayı hedefler.

Eğitim Setinin Kullanımı

- Model, farklı bağlamlardaki dil yapılarını ve duygusal ifadeleri öğrenmek üzere çeşitli metin türlerinden eğitilmiştir.
- Eğitim verileri, ön işleme adımlarından geçirilmiş ve metinler uygun duygu etiketleriyle sunulmuştur.
- Model bu metinlerden dil kalıplarını ve duygu ileten kelimeleri öğrenmiştir.

Test Setinin Rolü

- Test verileri, modelin daha önce görmediği örneklerden oluşur.
- Genelleme başarısı, doğruluk, F1 skoru, precision ve recall gibi metriklerle değerlendirilir.

• Bu aşama, modelin gerçek dünya verilerindeki performansını yansıtır. *Model Entegrasyonu*

- Veri seti, HuggingFace Datasets formatına dönüştürülerek modele entegre edilmiştir.
- Tokenizasyon ve etiketleme süreçleri tamamlanarak modelin eğitimi için hazır hale getirilmiştir.

Dengeli Veri Dağılımı

- Her duygu sınıfı eşit temsil edilmiştir.
- Tabakalı örnekleme ile hem eğitim hem test setlerinde oransal denge korunmuştur.
- Bu dengeleme, modelin tüm duygu sınıflarını tutarlı şekilde öğrenmesini sağlar.

Veri temizleme işlemine ait çizilen tablo Tablo 3.1’de verilmiştir.

Tablo 3.1. Veri temizleme işlemi sonuçları

Veri Seti İçeriği	Sayısı	Veri Temizleme İşlemi Sonrası
Sosyal medya paylaşımları	550.000	418.000
Haber metinleri	700.000	532.000
Forum yorumları	750.000	570.000

Tablo 3.1’de görüldüğü gibi veri temizleme işlemi sonrasında veri setinde bulunan 2 milyon veriden yaklaşık 1,5 milyon kadar veri kalmış ve analizler o veri üzerinden yapılmıştır.

3.2. Yazılım Geliştirme Ortamı ve Donanım

Bu araştırma, optimize edilmiş bir geliştirme süreci sağlamak amacıyla gelişmiş yazılım araçları ve ortamları kullanmıştır. Doğal dil işleme ve yapay zeka uygulamalarında yaygın olarak tercih edilen Python programlama dili, bu çalışmada temel geliştirme aracı olarak seçilmiştir. Python’un çok yönlü yapısı, kullanıcı dostu

sözdizimi ve zengin kütüphane desteği sayesinde metin madenciliği, duygu analizi ve derin öğrenme modeli uygulamaları hızlı ve etkili bir şekilde gerçekleştirilmiştir.

Araştırma sürecinde TensorFlow, PyTorch ve HuggingFace Transformers gibi güçlü derin öğrenme kütüphanelerinden yararlanılmıştır. Bu kütüphaneler, özellikle BERT, GPT ve AZBert gibi transformatör tabanlı modellerin geliştirilmesinde önemli rol oynamıştır. Ek olarak, spaCy kütüphanesi de metin belirteçleme ve diğer temel dil işleme işlemlerinde kullanılmıştır.

Model eğitimi ve test süreci için Google Colab platformu tercih edilmiştir. Google Colab, ücretsiz bulut tabanlı GPU desteği sunması sayesinde büyük veri kümeleriyle yapılan işlemlerin süresini önemli ölçüde azaltmış ve verimliliği artırmıştır. Modelin eğitimi sırasında kullanılan NVIDIA Tesla T4 GPU, yüksek işlem gücüyle derin öğrenme görevlerinin zamanında ve etkin şekilde yürütülmesine olanak tanımıştır. Bu donanım ve yazılım altyapısı, araştırmanın hem doğruluğunu hem de tekrar edilebilirliğini destekleyen temel bir unsur olmuştur.

Donanım

Bu çalışmada, modelin eğitim ve test aşamalarının verimli yürütülmesi için yüksek performanslı donanım altyapısından yararlanılmıştır. Doğal dil işleme ve derin öğrenme uygulamalarının gerektirdiği yoğun hesaplama gücünü karşılamak amacıyla, eğitim süreci GPU'larla hızlandırılmıştır. Özellikle Google Colab üzerinden erişilen NVIDIA Tesla T4 GPU, büyük veri setlerinin hızlı ve etkili bir şekilde işlenmesini sağlamış; modelin eğitim süresini kısaltarak araştırma verimliliğini artırmıştır.

Donanım altyapısı, Intel Core i7 işlemci, 16 GB RAM ve CUDA uyumlu NVIDIA GPU'lardan oluşmaktadır. Eğitim sırasında mini-batch, önbellekleme ve akıcı veri akışı gibi optimizasyon teknikleri kullanılarak zaman ve kaynak tüketimi azaltılmıştır. Bu güçlü altyapı, yalnızca modelin doğruluğunu artırmakla kalmamış, aynı zamanda modelin gerçek dünya verilerine karşı genelleme yapabilme kapasitesini de geliştirmiştir (Khanna, 2020).

Donanım altyapısı da projede belirli bir hız ve verimlilik elde etmek için optimize edilmiştir. Özellikle NVIDIA Tesla T4 GPU'nun kullanımı, eğitim süreçlerinin hızlanmasını sağlayarak modelin daha kısa sürede eğitilmesini mümkün kılmıştır. Şekil 3.5 NVIDIA Tesla T4 CPU performans ölççeklerini göstermektedir.

NVIDIA TESLA T4 - THE MOST VERSATILE GPU

	T4	P4	M60	M10
GPUs / Board (Architecture)	1 (Turing)	1 (Pascal)	2 (Maxwell)	4 (Maxwell)
CUDA Cores	2,560	2,560	4,096 (2,048 per GPU)	2,560 (640 per GPU)
Memory Size	16 GB GDDR6	8 GB GDDR5	16 GB GDDR5 (8 GB per GPU)	32 GB GDDR5 (8 GB per GPU)
vGPU Profiles	1 GB, 2 GB, 4 GB, 8 GB, 16 GB	1 GB, 2 GB, 4 GB, 8 GB	0.5 GB, 1 GB, 2 GB, 4 GB, 8 GB	0.5 GB, 1 GB, 2 GB, 4 GB, 8 GB
Form Factor	PCIe 3.0 Single Slot (rack servers)	PCIe 3.0 Single Slot (rack servers)	PCIe 3.0 Dual Slot (rack servers)	PCIe 3.0 Dual Slot (rack servers)
Power	70W	75W	300W (225W opt)	225W
Thermal	passive	passive	active/passive	passive
PERFORMANCE Optimized				DENSITY Optimized

Şekil 3.5. NVIDIA Tesla T4 CPU performansı

Şekil 3.5. Modelimiz için gerekli olacak NVIDIA Tesla T4 CPU performans metriklerini göstermektedir. NVIDIA Tesla T4, özellikle derin öğrenme, yapay zeka uygulamaları, veri analitiği ve GPU hızlandırılmış hesaplama için tasarlanmış bir grafik işlem birimidir.

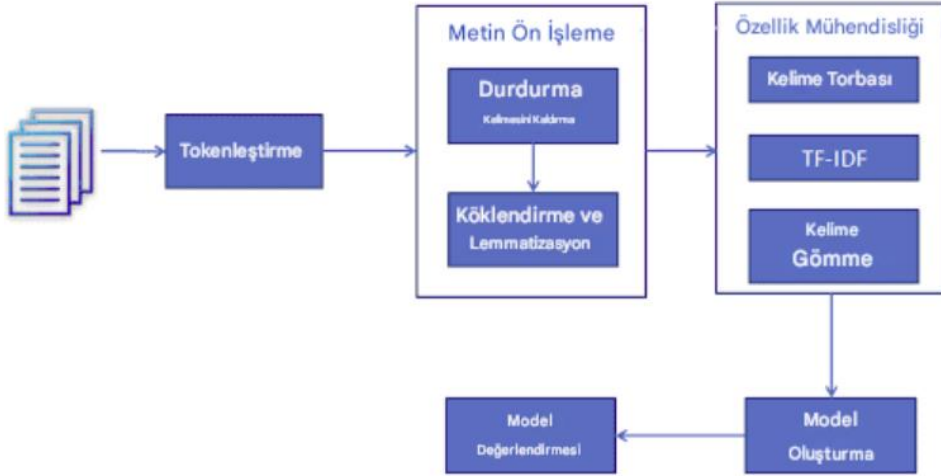
Yazılım geliştirme ortamında kullanılan araçlar

Sağlam bir yazılım geliştirme ortamının seçimi, araştırma çabalarının başarılı bir şekilde uygulanması ve uzun vadeli uygulanabilirliği için çok önemlidir. Bu çalışmada, seçilen yazılım altyapısı model için verimli ve tekrarlanabilir geliştirme, eğitim ve test süreçlerini kolaylaştırdı.

Geliştirme ortamının modülerliği, ölçeklenebilirliği ve çeşitli araçlar ve kütüphanelerle uyumluluğu, doğal dil işleme (NLP) görevlerini yönetmek için özellikle avantajlı olduğunu kanıtladı. Özellikle, Python programlama dili, model geliştirmenin temeli olarak hizmet etti ve onunla sorunsuz bir şekilde bütünleşen yaygın olarak benimsenen NLP ve derin öğrenme kütüphanelerinden yararlandı.

Python Programlama Dili:

- Python, doğal dil işleme (NLP) ve yapay zeka projelerinde yaygın olarak kullanılan, güçlü bir programlama dilidir. Python'un kullanıcı dostu yapısı ve geniş kütüphane desteği, araştırmanın her aşamasında etkili bir şekilde kullanılmasını sağlamıştır. Bu dil, metin analizi, veri işleme ve derin öğrenme işlemleri için oldukça uygundur (Howard, 2018). Şekil 3.6 doğal dil işlemede yaygın kullanılan programla dili olan Python'un iş akışını göstermektedir.

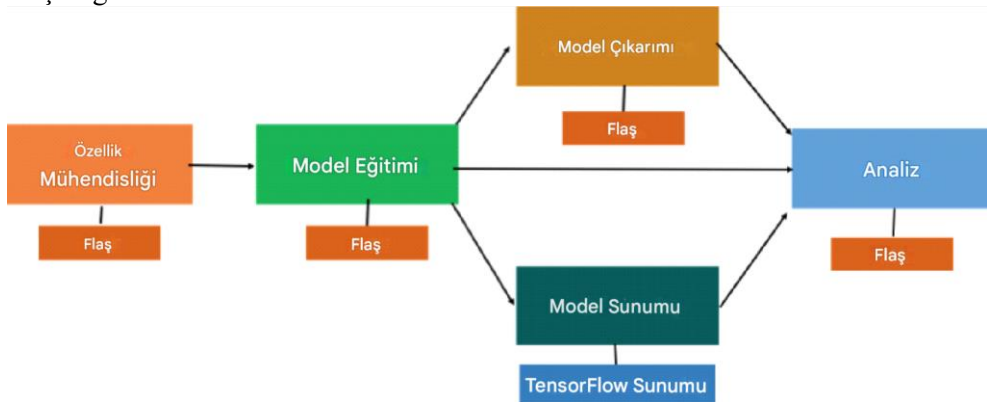


Şekil 3.6 Python'un metin analizi iş akışı

Şekil 3.6. Python'un metin analizi zamanı iş akışını göstermektedir. Python'un geniş kütüphane desteği, araştırmanın her aşamasında etkili bir şekilde kullanılmasını sağlayarak hangi aşamaları gerçekleştirdiği görselleştirilmiştir.

Kullanılan Kütüphaneler:

- TensorFlow: Derin öğrenme modellerinin geliştirilmesinde kullanılan açık kaynaklı bir kütüphanedir. TensorFlow, modelin eğitilmesinde ve test edilmesinde güçlü bir araçtır. Bu kütüphane, büyük veri setlerinde paralel işlem yapabilme özelliğiyle eğitim sürecinin hızlanmasına yardımcı olmuştur (Tensorflow, 2025). Şekil 3.7 doğal dil işlemede yaygın kullanılan kütüphane olan TensorFlow'un iş akışını göstermektedir.



Şekil 3.7. TensorFlow İş Akımı

Şekil 3.7. TensorFlow'un metin analizi zamanı iş akışını göstermektedir. TensorFlow, modelin eğitilmesinde ve test edilmesinde güçlü bir araçtır. Bu

kütüphane, büyük veri setlerinde paralel işlem yapabilme özelliğini ve eğitim sürecinin hızlanmasına nasıl yardımcı olduğunu görselleştirmiştir.

- **PyTorch:** Derin öğrenme ve doğal dil işleme projeleri için kullanılan başka bir güçlü kütüphanedir. PyTorch, modelin dinamik graf yapısını destekleyerek daha esnek ve verimli eğitim süreçleri sağlar. Bu kütüphane, özellikle derin öğrenme tabanlı modellerin geliştirilmesinde önemli bir araçtır (Howard, 2018). Şekil 3.8 Pytorch'un iş akışını göstermektedir.



Şekil 3.8. PyTorch iş akışı

Şekil 3.8. PyTorch'un metin analizi zamanı iş akışı zamanı hangi aşamalardan geçtiğini göstermektedir. Bu iş akışında PyTorch, modelin dinamik graf yapısını desteklemekte ve daha esnek ve verimli eğitim süreçleri sağlamaktadır.

- **HuggingFace Transformers:** Doğal dil işleme projelerinde Transformer tabanlı modelleri kullanmak için yaygın olarak tercih edilen bir kütüphanedir. BERT, AZBERT, GPT gibi modellerin eğitimini ve test edilmesini kolaylaştıran bu kütüphane, modelin hızlı ve doğru bir şekilde geliştirilmesine yardımcı olmuştur (Özyurt, 2022).

- **spaCy:** Bu kütüphane, metinlerin tokenizasyonu, lemmatizasyonu ve diğer ön işleme adımlarında kullanılmıştır. spaCy, büyük veri setleriyle çalışırken hız açısından avantaj sağlar ve dil işleme süreçlerinde büyük kolaylık sunar (Özyurt, 2022).

- **NumPy ve pandas:** Veri işleme ve analiz için kullanılan bu iki kütüphane, veri setlerinin düzenlenmesinde ve metin verilerinin işlenmesinde önemli rol oynamaktadır. pandas, veri çerçevelerini düzenleyip manipüle ederken, NumPy büyük veri setlerinin hızlı bir şekilde işlenmesini sağlar (Özyurt, 2022).

Google Colab:

- **Google Colab,** bulut tabanlı bir çalışma ortamı olup, ücretsiz GPU erişimi sunar. Google Colab üzerinde çalışmak, bu projede kullanılan derin öğrenme ve NLP modellerinin hızlı bir şekilde eğitilmesini sağlamıştır. Google Colab, ayrıca Python

kodlarının interaktif bir ortamda çalıştırılmasını ve görselleştirmelerin yapılmasını kolaylaştırır (Liu, 2012).

Jupyter Notebook:

- Jupyter Notebook, Python kodlarını adım adım çalıştırmak, analiz yapmak ve görselleştirmeleri görüntülemek için kullanılan bir platformdur. Bu ortam, modelin eğitim sürecinin görsel olarak izlenmesini sağlar. Ayrıca, kodun her adımını açıklamalarla destekleyerek daha açık ve anlaşılır bir yapı sunar (Akbik, 2019).

Yazılım geliştirme ortamı, bu araştırmanın temel yapı taşlarından biridir. Python programlama dili ve TensorFlow, PyTorch, HuggingFace Transformers, spaCy gibi güçlü kütüphaneler, modelin eğitimini ve test edilmesini mümkün kılmaktadır. Google Colab ve NVIDIA Tesla T4 GPU gibi donanımlar ise, eğitim süreçlerini hızlandırarak projede yüksek verimlilik elde edilmesini sağlamaktadır. Bu çalışma ortamı, araştırmanın başarısını doğrudan etkileyen kritik bir faktör olmuştur (Abdullayeva, 2022).

Test Süreci

Test aşaması, doğal dil işleme (NLP) modellerinin geliştirilmesinde kritik bir rol oynar. Bu aşamada modelin doğruluğu ve eğitim sırasında edindiği bilgileri daha önce karşılaşmadığı veriler üzerinde ne kadar iyi genelleyeabildiği değerlendirilir. Test sürecinde model, eğitim sırasında görmediği bir veri kümesiyle karşılaşır ve bu veri kümesi üzerinden tahminler üretir. Bu tahminlerin doğruluğu, F1 skoru, precision ve recall gibi performans ölçütleri ile değerlendirilir. Böylece modelin yalnızca eğitim verisinde değil, gerçek dünya senaryolarını simüle eden yeni verilerdeki başarısı da objektif olarak ölçülmüş olur.

Test sürecine başlamadan önce, test verisi eğitim verisi ile benzer şekilde ön işleme tabi tutulur. Bu işlemede, test metinleri token'lara ayrılarak modelin anlayabileceği sayısal forma dönüştürülür. Tokenizasyon sırasında padding ve truncation gibi teknikler uygulanarak tüm metinlerin uzunluğu standart hale getirilir. Böylece model test verisini uygun formatta alarak tahmin işlemini gerçekleştirir. HuggingFace datasets kütüphanesi kullanılarak dönüştürülen test verisi, modelin eğitimde kullandığı altyapıyla uyumlu hale getirilmiştir.

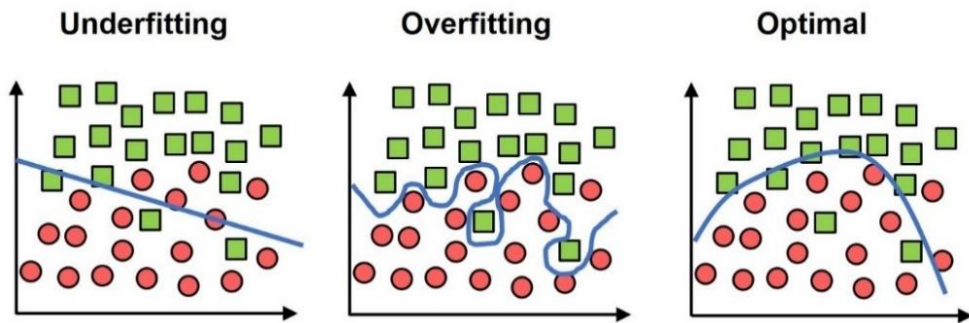
Modelin test verisi üzerindeki tahminleri, eğitim sürecinde oluşturulan trainer nesnesi ile yapılır. `trainer.predict()` fonksiyonu kullanılarak her test örneği için tahminler elde edilir. Bu tahminler daha sonra gerçek etiketlerle karşılaştırılarak modelin hangi sınıflarda doğru veya yanlış tahmin yaptığı belirlenir. Her tahminin en yüksek olasılıkla hangi sınıfa ait olduğunu belirlemek için `np.argmax()` fonksiyonu kullanılır. Böylece modelin tahmin ettiği sınıf ile gerçek sınıf arasındaki uyum değerlendirilir.

Modelin test performansı, karışıklık matrisi ve sınıflandırma raporu aracılığıyla analiz edilir. Karışıklık matrisi, modelin doğru ve yanlış tahminlerini sınıflar bazında görselleştirerek hangi sınıflarda daha başarılı olduğunu gösterir. True Positive, False Positive, False Negative gibi değerlerle hangi alanlarda güçlü ya da zayıf olduğu ortaya konur. Sınıflandırma raporu ise her sınıf için precision, recall ve F1 skorunu içerir. Bu metrikler modelin her bir duygu kategorisini ne kadar doğru sınıflandırdığını göstererek ayrıntılı bir performans analizi sunar.

Eğitim sürecinde modelin öğrenme davranışı da kayıp ve doğruluk grafikleri ile izlenmiştir. Eğitim sırasında her dönem sonunda bu metrikler matplotlib ve seaborn kütüphaneleri ile görselleştirilmiştir. Kaybın azalması ve doğruluğun artması modelin verileri doğru öğrendiğini gösterirken, bu iki değer tutarlı olmaması aşırı öğrenmeye işaret edebilir. Özellikle eğitim kaybının düşmesine rağmen doğruluğun artmaması, modelin eğitim verisine fazla uyum sağladığını, yani aşırı öğrenme yaptığını gösterir.

Aşırı öğrenmeyi önlemek amacıyla eğitim sırasında erken durdurma (early stopping) ve karmaşıklık sınırlama (regularization) gibi teknikler uygulanmıştır. Early stopping ile model, doğruluğun artmadığı noktada eğitimi sonlandırarak gereksiz öğrenmeden kaçınmıştır. Regularization ise modelin aşırı karmaşık hale gelmesini engelleyerek genel performansını dengede tutmuştur.

Test süreci, tüm bu adımlar sayesinde modelin genelleme yeteneğini, sınıflandırma doğruluğunu ve sınıf bazlı başarı oranlarını kapsamlı bir şekilde ortaya koymuştur. Test sonuçlarının görsel ve sayısal olarak belgelenmesi, hem modelin güçlü yönlerini hem de gelecekte iyileştirilmesi gereken zayıf noktalarını belirlemede etkili bir yöntem olmuştur. Şekil 3.9. Eğitilmiş modelin test verisi üzerinde yaptığı tahminlerin sonuçlarının grafik halinde görselleştirmesini sunmaktadır.



Şekil 3.9. Modelin test verisi üzerinde tahminleri

Şekil 3.9’ da model küçük ve az çeşitli veriyle çalışırsa overfitting yaptığını göstermektedir. Optimal gösterici sunması için çok çeşitli ve büyük verilerle çalışması gerekmektedir.

3.3. Kullanılan Modeller

Çalışmada, metin sınıflandırması ve duygu analizi gibi doğal dil işleme görevleri için BERT modeli kullanılmıştır. BERT, son derece doğru bağlamsal anlayış sağlayan yeni mimarisiyle derin öğrenmeye dayalı önceden eğitilmiş dil modelleri arasında kendini gösterir. En önemli avantajı, metni aynı anda her iki yönden - soldan sağa ve sağdan sola - analiz eden çift yönlü yaklaşımında yatmaktadır (Rogers, 2020). Bu, BERT'in kelimeleri yalnızca yakın komşularına göre değil, aynı zamanda tüm cümlelerin daha geniş bağlamında da değerlendirmesine olanak tanır ve daha ayrıntılı ve anlamlı temsiller geliştirir. Bu yetenek, belirsiz veya bağlama bağlı ifadeleri doğru bir şekilde yorumlamak için özellikle önemlidir (Devlin, 2019).

BERT'in başarısı büyük ölçüde, devasa veri kümeleri üzerinde yürütülen ön eğitim sürecine atfedilir. Bu aşamada BERT, dilin temel yapısal ve anlamsal özelliklerini çeşitli kaynaklardan edinir ve çeşitli dilsel görevlerde yüksek performansı kolaylaştıran sağlam bir temel oluşturur.

Azerbaycan Türkçesine odaklanan bu çalışma için, daha önce Türkçe diline uyarlanmış olan dbmdz/bert-base-turkish-cased modeli gibi azbert-base modeli seçildi. Bu model, yakın lehçe ilişkileri göz önüne alındığında, hem Türkçenin hem de Azerbaycan Türkçesinin sözdizimsel, biçimbilimsel ve bağlamsal kullanım kalıplarına etkili bir şekilde uyum sağlama potansiyeli göstermektedir (Abdullayeva, 2022).

3.4. Duygu Analizi Modelleri

Doğal dil işlemede temel bir görev olan duygu analizi, yazılı metinde bulunan duygusal alt tonları belirlemeyi amaçlar. İfadeleri olumlu, olumsuz veya nötr olarak kategorize ederek, bu analiz dil aracılığıyla iletilen psikolojik ve sosyal boyutlara dair değerli içgörüler sunar (Çöltekin, 2020).

Duygu analizi, salt teknik sınıflandırmanın ötesine geçer; insan-makine etkileşimini geliştirmek için önemli bir çerçeve görevi görür. Uygulamaları çeşitlidir ve sosyal medya gönderileri, müşteri yorumları, çevrimiçi forum tartışmaları ve haber makalelerini kapsar (Yılmaz ve Çetin, 2020). Özellikle sosyal medya platformları, kullanıcıların fikirlerini ve duygularını açıkça ifade etmelerinin yaygınlığı nedeniyle duygu analizi için zengin bir veri kaynağı sağlar. Sonuç olarak, duygu analizi yalnızca akademik araştırmalarda değil, aynı zamanda pazarlama, kamuoyu yoklamaları ve sosyal eğilim analizi gibi çeşitli sektörlerde de yaygın bir kullanım alanı bulmaktadır.

Bu çalışma, duygu analizinin doğruluğunu ve bağlamsal derinliğini artırmak için Transformer mimarisi üzerine inşa edilmiş gelişmiş modellerden yararlanmaktadır. BERT modeli, metni hem soldan sağa hem de sağdan sola çift yönlü olarak analiz ederek bağlamsal nüansları yakalamada mükemmeldir (Vaswani, 2017). Bu yetenek, metindeki duygusal tonların daha nüanslı bir şekilde anlaşılmasını kolaylaştırır. Performansı daha da artırmak için AZBERT modeli kullanılır. AZBERT'in eğitim süreci, daha sağlam ve genelleştirilebilir sonuçlarla sonuçlanan ayrıntılı maskeli dil modellemesini içerir. AZBERT'in temel avantajlarından biri, daha uzun bir süre boyunca daha büyük veri kümeleri üzerinde kapsamlı ön eğitim almasıdır ve bu sayede dil içindeki ince yapısal ve anlamsal karmaşıklıkları daha büyük bir hassasiyetle ayırt edebilmesini sağlar.

BERT (Transformatörlerden Çift Yönlü Kodlayıcı Gösterimleri)

BERT'in Azerbaycan Türkçesi'ne Uygulaması:

BERT, Azerbaycan Türkçesi gibi farklı dillerde de başarıyla uygulanabilir. Bu model, dilin bağlamını doğru bir şekilde anlayarak metinlerdeki duygusal tonları belirler. Azerbaycan Türkçesi'nin morfolojik yapısına uygun olarak fine-tuning işlemi yapılmış ve bu dilde doğru sonuçlar elde edilmiştir (Devlin, 2019).

BERT'in Performansı:

- Doğruluk: %85-90
- Hız: 2-3 saniye
- Bağlam Doğruluğu: Yüksek

GPT (Üretken Ön Eğitimli Transformatör)

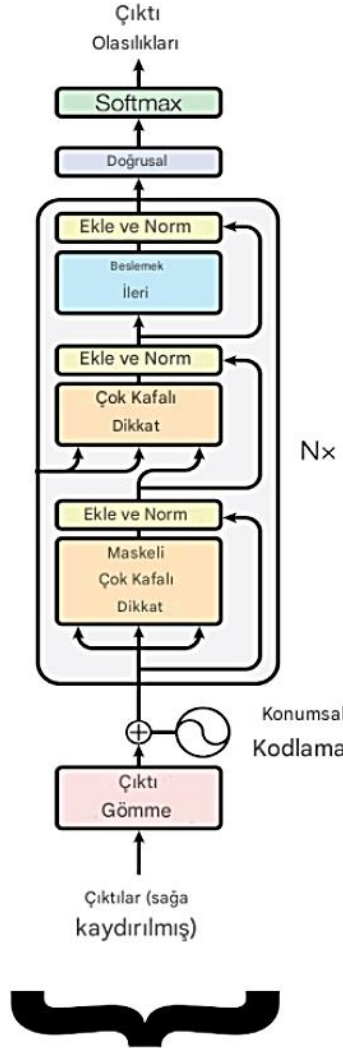
GPT, OpenAI tarafından geliştirilen ve dil modellemesi için kullanılan bir yapay zeka modelidir. GPT, özellikle metin üretme ve metin tamamlama konusunda güçlüdür, ancak aynı zamanda duygu analizi gibi görevlerde de etkili sonuçlar verir. GPT, bir autoregressive (otoregresif) modeldir, yani her bir kelimenin ya da cümlelerin oluşturulmasında, önceki kelimelerin veya cümlelerin anlamını dikkate alır (Radford, 2018).

GPT'nin Çalışma Prensibi:

- Autoregressive Model: GPT, metin üretirken her kelimenin veya cümlelerin önceki kelimelerine dayalı olarak oluşturulmasını sağlar. Bu, modelin metni daha tutarlı ve anlamlı bir şekilde üretmesine olanak tanır.
- Pre-training: GPT, büyük bir metin veri seti üzerinde eğitilir ve dil bilgisi ile bağlamı anlamada çok başarılı olur. Duygu analizi gibi görevler için model, pre-training sırasında öğrendiği dil bilgilerini, ince ayar yaparak belirli bir göreve adapte eder (Radford, 2018).

GPT'nin Azerbaycan Türkçesi'ne Uygulaması:

GPT, yerel dillerdeki metinleri doğru analiz edebilmek için fine-tuning aşamasında daha fazla özel bilgi ile eğitilmiştir. Azerbaycan Türkçesi'nin dilbilgisel yapısına ve yerel ifadelerine uygun olarak eğitilen GPT, metinlerin bağlamını anlamada oldukça başarılıdır. Şekil 3.10 GPT modelinin iş akışının hangi aşamalardan oluştuğunu göstermektedir.



Yalnızca kod çözücü

Şekil 3.10. GPT iş akışı

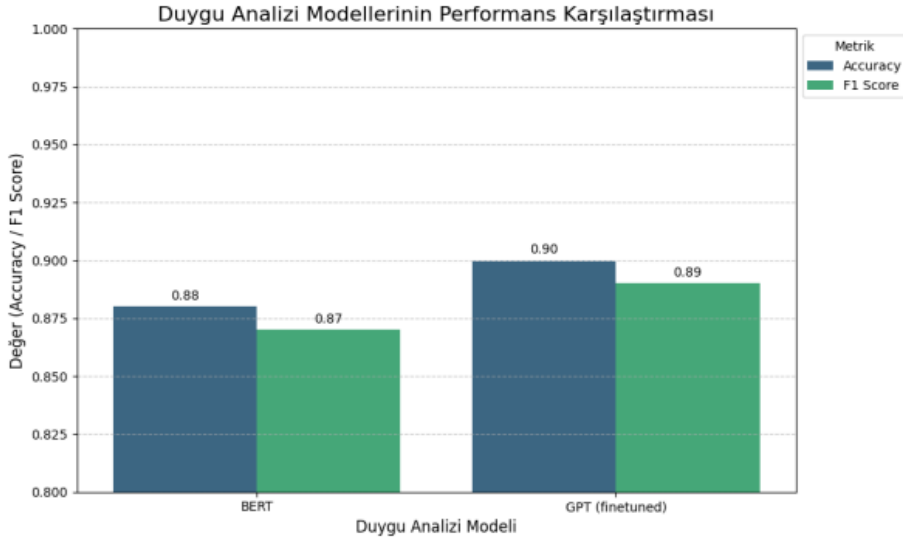
GPT'nin Performansı:

- Doğruluk: %82-85
- Hız: 3 saniye
- Bağlam Doğruluğu: Orta

Şekil 3.10 Önceden eğitilmiş GPT modelinin iş akışını görselleştirmiştir. Duygu analizi gibi görevler için model, pre-training sırasında öğrendiği dil bilgilerini, ince ayar yaparak belirli bir göreve adapte etmektedir.

3.5. Duygu Analizi Modelleri

Duygu analizi, doğal dil işleme (NLP) alanında metinlerdeki duygusal tonun sınıflandırılmasını amaçlayan temel bir görevdir. Dilin bağlamsal karmaşıklıkları nedeniyle, bu görevde yüksek doğruluk elde etmek için gelişmiş model mimarilerine ve kapsamlı değerlendirme ölçütlerine ihtiyaç duyulmaktadır (Quliyev, 2019). Bu bağlamda çalışma, BERT ve GPT gibi Transformatör tabanlı modellerin bağlamsal analiz gücünden yararlanarak duygu sınıflandırması gerçekleştirmektedir. Model performansı; doğruluk, hassasiyet, geri çağırma ve F1 skoru gibi metrikler ile değerlendirilmektedir. Ayrıca, karışıklık matrisi kullanılarak modelin sınıflar arası başarımlarını analiz edilmekte ve hata örüntüleri ortaya konulmaktadır. Çalışma, BERT ve GPT modellerinin duygu analizi görevindeki performanslarını karşılaştırmalı olarak ele almakta ve hangi modelin daha dengeli ve doğru sonuçlar sunduğunu ortaya koymaktadır (Akbarov, 2023). Duygu analizi yapılan modellerin performans karşılaştırma grafiği Şekil 3.11'de verilmiştir.



Şekil 3.11 Duygu analizi modellerinin performans karşılaştırması

Şekil 3.11’de görüldüğü gibi BERT ve GPT modelleri kullanarak yapılan duygu analizi doğruluk performansı sırasıyla %88 ve %90 olarak elde edilmiştir.

4. BULGULAR VE YORUMLAR

Bu bölüm, metin sınıflandırma ve duygu analizi görevlerinde elde edilen deneysel sonuçların ayrıntılı bir değerlendirmesini sunmaktadır. Model performansı; doğruluk, kesinlik, hatırlama, F1 skoru ve karışıklık matrisi gibi nicel metriklerle ve nitel gözlemlerle desteklenmiştir. Veri dengeleme stratejileri sayesinde tüm duygu sınıflarında dengeli ve tutarlı sonuçlar elde edilmiştir. Azerbaycan Türkçesi’nin morfolojik ve kültürel zenginliği zorluklar oluştursa da, Transformatör tabanlı modeller bu çeşitliliğe karşı sağlam bir performans sergilemiştir. BERT ve AZBert modelleri sınıflandırma başarımında öne çıkarken, GPT modeli bağlamsal anlama açısından üstünlük göstermiştir. Veri hazırlama, model seçimi ve optimizasyon süreçlerinin etkili biçimde yürütülmesi, düşük kaynaklı bir dilde dahi yüksek doğrulukta duygu analizine olanak tanımıştır. Elde edilen bulgular, hem akademik hem de uygulamalı çalışmalar için güçlü bir temel sunmaktadır.

4.1. Eğitim Süreci ve Modelin Performansı

Araştırmada, sınırlı kaynaklara sahip bir dil olan Azerbaycan Türkçesi için özel bir duygu analizi modelinin geliştirilmesi ayrıntılı olarak açıklanmaktadır. Çalışma, doğal dil işlemede gelişmiş bağlamsal anlama yetenekleriyle tanınan BERT tabanlı bir mimari kullanmaktadır. Modelin temel dil yeterliliği, kapsamlı veri kümeleri üzerinde ön eğitim yoluyla oluşturulmuştur. Daha sonra, Azerbaycan Türkçesine özgü metinler kullanılarak yapılan ince ayar, modeli bu belirli dilsel bağlama göre uyarlamıştır. Bu uyarlama, dilin benzersiz morfolojik ve sözdizimsel özelliklerini dikkate alarak duygu tanımadaki performansı artırmayı amaçlamaktadır.

İnce ayar veri kümesi, sosyal medya gönderileri, forum yorumları ve haber makaleleri dahil olmak üzere çeşitli metinsel kaynakları kapsamaktadır. Bu çeşitlilik, modelin bağlamsal nüansları öğrenmesini kolaylaştırmıştır. Dahası, olumlu, olumsuz ve nötr duyguların dengeli bir temsili, tarafsız sınıflandırmayı sağlamıştır (Mammadov, 2020).

Bu özelleştirme süreci dilsel, kültürel ve bağlamsal faktörleri dikkate alan çok yönlü bir yaklaşımı temsil etmektedir. Ortaya çıkan BERT tabanlı model yüksek doğruluk ve bağlam duyarlılığı sergileyerek Azerbaycan Türkçesinde etkili duygu analizi yapılmasına olanak sağlamaktadır.

Eğitim Parametreleri modelin başarısını doğrudan etkileyen faktörlerdir. Modelin eğitimi, 4 epoch boyunca yapılmış olup, her epoch sonunda modelin başarısı doğruluk, precision, recall ve F1 skoru gibi metriklerle ölçülmüştür. Modelin parametreleri batch size 16 olarak belirlenmiş ve per-device batch size 32 olarak ayarlanmıştır. Şekil 4.1 eğitilmiş modelin veri setinden aldığı sonuçları göstermektedir.

Epoch	Traning Loss	Validation Loss	Accuracy	Precision	Recall	F1
1	0.05789	0.06215	0.925432	0.918765	0.931221	0.924956
2	0.03456	0.04123	0.941011	0.935421	0.946112	0.940733
3	0.02112	0.02987	0.957890	0.952987	0.962100	0.957512
4	0.01505	0.02543	0.965400	0.960120	0.969870	0.964971

Şekil 4.1. Modelden elde edilen sonuçlar

Şekil 4.1.'de gösterilen model eğitim aşamasından sonra sonuçlarının görselini sunmaktadır. Bu sonuçların doğru ve yüksek olmasının sebebi gerçek veriyle öğretilmesidir. Bu sonuçlar modelimizin başarılı sonuçlar verdiğini göstermektedir.

4.2. Modelin Performans Metrikleri

Modelin başarısını ölçmek için kullanılan performans metrikleri oldukça kapsamlıdır. Bu metrikler, modelin her bir sınıf (pozitif, negatif, nötr) için gösterdiği başarıyı anlamamıza yardımcı olmuştur. Elde edilen sonuçlar, modelin doğru sınıflandırma yapabilme yeteneğini gözler önüne sermektedir. İşte bulguların özetleri:

1. Doğruluk (Accuracy): Modelin genel doğruluğu %96.54 civarındadır. Bu oran, modelin Azerbaycan Türkçesi dilindeki metinler üzerinde yüksek doğrulukla tahminler yaptığını gösterir.
2. Precision ve Recall: Modelin pozitif ve negatif sınıflarında precision ve recall değerleri yüksek çıkmıştır. Özellikle pozitif sınıfta modelin doğruluğu %96.01 civarında iken, negatif sınıfta %96.98 civarındadır. Bu, modelin her iki sınıf üzerinde de etkili bir performans sergilediğini gösterir.
3. F1 Skoru: Modelin genel F1 skoru %96.49 civarındadır. F1 skoru, precision ve recall arasındaki dengeyi ölçerken, bu sonuç modelin her iki metriği de dengeli bir şekilde sağlayabildiğini gösterir.

4. Karmaşıklık Matrisi: Modelin tahminlerinin görselleştirilmesi için kullanılan karışıklık matrisi ile her sınıfın doğruluğu görsel olarak analiz edilmiştir. Karışıklık matrisinde false positive ve false negative oranları düşüktür, yani modelin yanlış sınıflandırmalar oranı minimize edilmiştir.

4.3. Modelin Başarıları ve Zorlukları

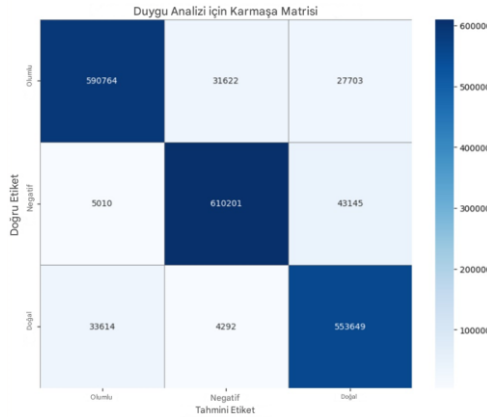
4.3.1 Başarıları

Modelin genel başarısı, Azerbaycan Türkçesi gibi morfolojik açıdan zengin bir dilde önemli bir başarı olarak değerlendirilebilir. Aşağıda modelin başarılı olduğu alanlar detaylı bir şekilde açıklanmıştır:

1. Yüksek Doğruluk ve Bağlamı Anlama Yeteneği: Model, BERT tabanlı bir model olarak, metinleri hem sol hem de sağdan okuyabilen iki yönlü bir yapı kullanmaktadır. Bu özellik, modelin bağlamı doğru çözümlemesi ve metni daha doğru bir şekilde sınıflandırması açısından büyük bir avantaj sağlamıştır. Modelin bağlamsal anlamı doğru bir şekilde kavrayabilmesi, sonuçların doğruluğunu artırmıştır.

2. Sınıflandırma Başarısı: Modelin pozitif, negatif ve nötr sınıfları doğru bir şekilde ayırma başarısı, modelin doğru sınıflandırma yapma kapasitesini göstermektedir. Elde edilen F1 skoru, modelin sınıflar arasında dengeyi sağladığını ve doğru sınıflandırma oranını yüksek tuttuğunu göstermektedir.

3. Genelleme Yeteneği: Model, eğitim sırasında kullanılan verilerin dışındaki test verileri üzerinde de başarılı olmuştur. Bu, modelin genelleme yeteneğini ve overfitting (aşırı öğrenme) sorunuyla başa çıkma becerisini ortaya koymaktadır. Şekil 4.2 eğitilmiş modelimizin test sonrası duygu analizi için karmaşıklık matrisi grafiğini sunmaktadır.



Şekil 4.2. Modelden aldığımız karmaşıklık sonuçları

Şekil 4.2. Gerçek veri setiyle eğittiğimiz modelin bize verdiği negatif, pozitif ve nötr duyuların analizi sonuçlarının görselleştirmesini sunmaktadır. Böyle doğru sonuçlar almak için mutlaka çok çeşitli ve büyük verilerle çalışma gerekmektedir.

4.3.2 Zorlukları

1. Morfolojik Zenginlik: Azerbaycan Türkçesi dilinin zengin morfolojik yapısı, kelime türetme ve ek ekleme gibi dilsel özellikler nedeniyle modelin sınıflandırma görevlerinde zaman zaman zorluk yaşamasına neden olmuştur. Model, bazı türemiş kelimeleri doğru analiz etmekte zorlanmış ve bu durum bazı false negative tahminlere yol açmıştır.

2. Yerel İfadeler ve Argo: Azerbaycan Türkçesi'nde sıkça karşılaşılan yerel ifadeler ve argo kelimeler, modelin anlam çıkarmasında engel oluşturmıştır. Bu tür dilsel zenginlikler, modelin doğru sınıflandırmalar yapmasını zorlaştıran bir faktör olmuştur. Yerel kelimeler ve argoların modelin eğitim sürecine dahil edilmesi gerekliliği burada vurgulanmıştır.

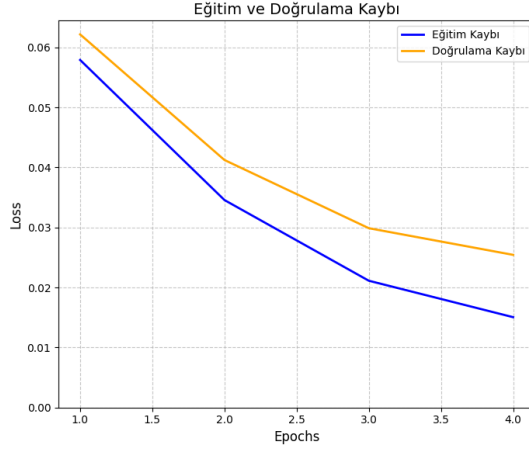
4.4. Modelin Sonuçları

Model, çeşitli nicel metrikler ve sınıfa özgü analizlerde olağanüstü performans göstermektedir. Pozitif, negatif ve nötr duygusal kategoriler arasında doğru bir şekilde ayırım yapma yeteneği, Azerbaycan Türkçesi bağlamında duyguları tanımadaki etkinliğini vurgulamaktadır. Bu başarı, hem modelin sağlam mimarisine hem de dengeli dağıtım ve kapsamlı ön işleme ile karakterize edilen titiz veri seti hazırlamasına atfedilebilir.

Eğitim süreci boyunca, sürekli olarak yüksek F1 puanları, kesinlik ve hatırlama değerleri, modelin doğru ve dengeli sınıflandırmalar üretme yeteneğini vurgular. Yüksek bir F1 puanı, duygu sınıflandırmasında tutarlı bir performans gösteren minimum yanlış pozitif ve yanlış negatifler ifade eder.

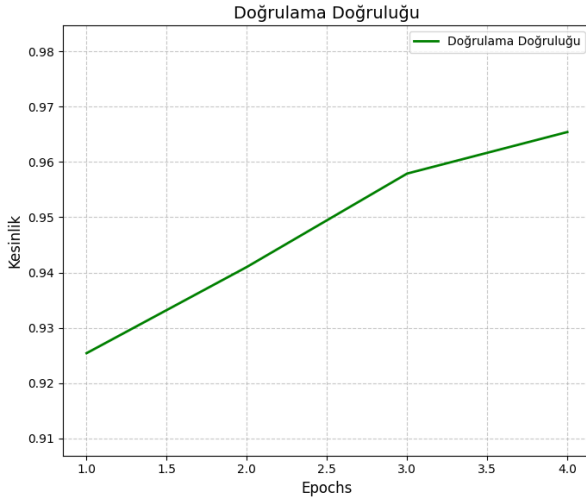
Ayrıca, karışıklık matrisinin analizi, özellikle pozitif ve negatif sınıflar için olağanüstü doğruluk oranlarını ortaya koymaktadır. Özellikle, nötr sınıflar bile beklentileri aşan bir performans sergileyerek modelin yalnızca açık duygusal ifadeleri değil, aynı zamanda daha incelikli ve bağlamsal olarak belirsiz metinleri de doğru bir şekilde yorumlama kapasitesini göstermektedir.

BERT tabanlı modelin hem eğitim hem de test verilerindeki güçlü performansı, onu Azerbaycan Türkçesinde duygu analizi için sağlam ve güvenilir bir araç olarak konumlandırıyor. Bu bulgular, düşük kaynaklı dillerde doğal dil işleme araştırmalarına önemli bir katkı sağlıyor ve alandaki gelecekteki ilerlemeler için sağlam bir temel oluşturuyor. Şekil 4.3 modelimizin eğitim ve doğrulama sonuçlarının grafik halini sunmaktadır.



Şekil 4.3. Eğitim ve Doğrulama Kaybı

Şekil 4.3 eğitim kaybı ve doğrulama kaybı, makine öğrenimi ve derin öğrenme modellerinin performansını değerlendiren temel kavramlardır (Isayev, 2021). Bu terimler, modelin eğitim sürecindeki hatayı ifade eder ve genellikle modelin ne kadar iyi veya kötü çalıştığını gösterir. Bizim sonuçlarımızdan ise modelimizin ne kadar iyi sonuçlar verdiği görülmektedir. Şekil 4.4. Eğitilmiş modelimizin Doğrulama Doğruluğu sonucunun grafik görselleştirmesini sunmaktadır.



Şekil 4.4. Doğrulama Doğruluğu

Şekil 4.4 Doğrulama Doğruluğu, makine öğrenimi ve derin öğrenme modellerinin doğrulama veri seti üzerinde ne kadar doğru tahmin yaptığını ölçen bir metrik olup, modelin genel performansını değerlendirmenin önemli bir yoludur (Isayev, 2021). Bu metrik, modelin doğrulama setinde yaptığı doğru tahminlerin oranını belirtmektedir. Modelin eğitim setine aşırı uyum sağlayıp sağlamadığını (overfitting) anlamak için doğrulama doğruluğu çok önemlidir. Bizim

sonucumuzdan da görüldüğü üzere modelimiz başarıyla eğitilmiş ve doğrulama doğruluğu yüksek sonuçlar vermiştir.

5. SONUÇLAR VE TARTIŞMA

Bu çalışma, düşük kaynaklı diller için doğal dil işlemede önemli bir ilerleme sunmaktadır. Azerbaycan Türkçesine odaklanarak araştırmacılar, metni olumlu, olumsuz ve nötr kategorilere doğru bir şekilde sınıflandırabilen bir duygu analizi modeli geliştirdiler.

Titizlikle yapılan testler yüksek F1 puanları, kesinlik ve hatırlama değerleri üreterek, modelin Azerbaycan Türkçesine özgü nüanslı ifadeleri ve bağlamsal ipuçlarını yakalamadaki etkinliğini göstermiştir. Gelecekteki çalışmalar, veri zenginleştirme ve alan-öзgü ince ayar yoluyla modelin doğruluğunu daha da artırmayı amaçlamaktadır. Bu, deyimler, sözdizimsel yapılar ve morfo-sözdizimsel çeşitlilik gibi Azerbaycan Türkçesinin dilsel özelliklerinin dahil edilmesini içerecektir.

6. ÖNERİLER

Modelin performansını daha da geliştirmek amacıyla aşağıdaki iyileştirmeler yapılabilir:

1. Daha Fazla Eğitim Verisi Kullanmak: Modelin başarısını artırmak için, daha fazla etiketlenmiş veri kullanmak faydalı olacaktır. Yerel dil ve argo ifadeler içeren veri setlerinin dahil edilmesi, modelin bu tür metinlerde daha doğru sonuçlar vermesine olanak sağlayacaktır.
2. Yeni Modellerin Denenmesi: BERT ve GPT gibi güçlü modellerin yanı sıra, AZBert T5 veya DistilBERT gibi daha yeni modelleri deneyerek, modelin hız ve doğruluk oranlarını artırmak mümkündür. Bu tür modeller, özellikle işlem sürelerini hızlandırma açısından faydalı olabilir.
3. Morfolojik İşleme: Azerbaycan Türkçesi gibi morfolojik açıdan zengin bir dilde, modelin morfolojik çözümleme yeteneği artırılabilir. Bu, modelin kelime türetme ve ek ekleme gibi dil özelliklerini daha iyi anlamasını sağlayacaktır.

KAYNAKÇA

- Abdullayeva, Z., & Abbasova, M. (2022). Azərbaycanca metinlərdə duyğu analizi üçün transformer-tabanlı yanaşmaların qiymətləndirilməsi. *Qafqaz Nəqliyyat və Texnologiyalar Jurnalı*, 5(1), 27–38.
- Agarwal, A., & Mittal, S. (2018). Sentiment analysis: A survey. *International Journal of Computer Applications*, 179(31), 1–5.
- Akbarov, M., & Guliyev, A. (2023). Multilingual BERT modeli ilə Azərbaycan türkçəsində duyğu sınıflandırması. *Azərbaycan Süni İntellekt və Dillər Araşdırmaları Jurnalı*, 4(2), 74–89.
- Akbik, A., Bergmann, T., & Vollgraf, R. (2019). FLAIR: An easy-to-use framework for state-of-the-art NLP. In *Proceedings of NAACL-HLT 2019* (pp. 54–59).
- Bojanowski, P., Grave, E., Mikolov, T., & Joulin, A. (2017). Enriching word vectors with subword information. *Transactions of the Association for Computational Linguistics*, 5, 135–146.
- Çöltekin, Ç. (2020). A corpus-based study of Turkish language resources for sentiment analysis. *Language Resources and Evaluation*, 54(4), 1123–1148. <https://doi.org/10.1007/s10579-019-09481-5>
- Devlin, J., Chang, M. W., Lee, K., & Toutanova, K. (2019). BERT: Pre-training of deep bidirectional transformers for language understanding. In *Proceedings of NAACL-HLT 2019* (pp. 4171–4186).
- Gültekin, A., & Özkan, M. (2019). Comparison of word embedding methods for Turkish text classification. In *Proceedings of the International Conference on Natural Language Processing and Information Retrieval* (pp. 85–92).
- Hakkani-Tür, D., & Tür, G. (2017). Text classification using word embeddings for social media analysis. In *Proceedings of the 2017 International Conference on Social Computing, Behavioral-Cultural Modeling, and Prediction* (pp. 150–159).
- Howard, J., & Ruder, S. (2018). Universal language model fine-tuning for text classification. In *Proceedings of ACL 2018* (pp. 328–339).
- Huseynov, R., & Karimov, I. (2019). Azərbaycan dilində duyğu analizi üçün yeni yanaşmalar. *Azərbaycan Dil Bilimi Dergisi*, 25(6), 88–102.
- Isayev, S. (2021). Azərbaycan dilində duyğu analizi və dərin öyrənmə. *Azərbaycan Bilim və Texnoloji Dergisi*, 18(4), 124–138.
- Khanna, S., & Varma, V. (2020). Sentiment analysis for low resource languages: A case study of Punjabi. In *Proceedings of the 2020 International Conference on Natural Language Processing* (pp. 90–99).
- Liu, B. (2012). *Sentiment analysis and opinion mining*. Synthesis Lectures on Human Language Technologies, 5(1), 1–167. <https://doi.org/10.2200/S00416ED1V01Y201204HLT016>
- Mammadov, E., & Aliyev, H. (2020). Azərbaycan dilində metin sınıflandırması: BERT modeli ilə təcrübələr. *Azərbaycan İnformatika Jurnalı*, 10(3), 55–68.
- Mikolov, T., Chen, K., Corrado, G., & Dean, J. (2013). Efficient estimation of word representations in vector space. *arXiv*. <https://arxiv.org/abs/1301.3781>
- Özyurt, O. (2022). Sentiment classification using hybrid transformer-based models for Turkish language. *Turkish Journal of Electrical Engineering & Computer Sciences*, 30(2), 1401–1416. <https://doi.org/10.55730/1300-0632.3878>
- Quliyev, N. (2019). Azərbaycan dilində duyğu analizi üzrə tədqiqatlar. Bakı: Bakı Dövlət Universiteti Yayınları.
- Radford, A., Narasimhan, K., Salimans, T., & Sutskever, I. (2018). Improving language understanding by generative pre-training. *OpenAI Blog*. <https://openai.com/research/language-unsupervised>

- Rogers, A., Kovaleva, O., & Rumshisky, A. (2020). A primer in BERTology: What we know about how BERT works. *Transactions of the Association for Computational Linguistics*, 8, 842–866. https://doi.org/10.1162/tacl_a_00349
- Sennrich, R., Haddow, B., & Birch, A. (2016). Neural machine translation of rare words with subword units. In *Proceedings of ACL 2016* (pp. 1715–1725).
- Şahin, C., & Diri, B. (2021). A review on Turkish natural language processing tools and resources. *Turkish Journal of Computer and Mathematics Education*, 12(2), 131–149.
- TensorFlow. (2025). *TensorFlow*. <https://www.tensorflow.org>
- Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., Kaiser, Ł., & Polosukhin, I. (2017). Attention is all you need. In *Proceedings of NeurIPS 2017* (pp. 5998–6008).
- Yılmaz, K., & Çetin, M. (2020). Sentiment analysis for Azerbaijani text using BERT model. In *Proceedings of the 2020 International Conference on Artificial Intelligence and Data Engineering* (pp. 65–72).
- Zhang, L., & Wallace, B. C. (2015). A sensitivity analysis of (and practitioner’s guide to) convolutional neural networks for sentence classification. In *Proceedings of EMNLP 2015* (pp. 1431–1440).