

**ANKARA ÜNİVERSİTESİ
FEN BİLİMLERİ ENSTİTÜSÜ**

DOKTORA TEZİ

**AZERBAYCAN TÜRKÇESİ İÇİN BİR DUYGU ANALİZİ
ÇERÇEVESİ GELİŞTİRME**

Yusuf Evren AYKAÇ

BİLGİSAYAR MÜHENDİSLİĞİ ANABİLİM DALI

**ANKARA
2026**

Her hakkı saklıdır

ÖZET

Doktora Tezi

AZERBAIJAN TÜRKÇESİ İÇİN BİR DUYGU ANALİZİ ÇERÇEVESİ GELİŞTİRME

Yusuf Evren AYKAÇ

Ankara Üniversitesi
Fen Bilimleri Enstitüsü
Bilgisayar Mühendisliği Anabilim Dalı

Danışman: Prof. Dr. Refik SAMET

Bu tez, Azerbaycan dili için duygu analizinin yalnızca sınıflandırma başarımına indirgenemeyeceği, dil kaynağı geliştirme, morfoloji-duyarlı modelleme, alanlar arası sağlamlık ve açıklanabilirliğin birlikte ele alınması gerektiği varsayımından hareket etmektedir. Bu amaçla, Azerbaycan dili için açıklanabilir, morfoloji-duyarlı ve alanlar arası sağlam bir duygu analizi çerçevesi geliştirilmiştir. Çalışmada problem alanı ve veri gereksinimleri tanımlanmış, veri toplama, temizleme ve ön işleme süreçleri yürütülmüş, Azerbaycan dili için *SentiAzNet* adlı kutupluluk sözlüğü geliştirilmiş, bu sözlük bilgisinin finans ve iş dünyası alanındaki aspekt-tabanlı duygu analizi görevine entegrasyonu incelenmiş, *MahSA* adlı morfoloji-duyarlı çok alanlı hibrit mimari tasarlanmış ve alan kaymasına dayanıklı cümle temsilleri öğrenilmiştir. Tüm bileşenler açıklanabilirlik, hata analizi ve karşılaştırmalı deneyler çerçevesinde birlikte değerlendirilmiştir. Bulgular, geliştirilen kutupluluk sözlüğünün güvenilir bir sembolik önbilgi kaynağı sunduğunu, sözlük bilgisinin bağlamsal modellere uygun biçimde aktarıldığında performansı ve yorumlanabilirliği güçlendirdiğini, hibrit mimarinin özellikle nötr sınıf ve zor alanlarda daha dengeli sonuçlar ürettiğini ve alanlar arası sağlam temsil öğrenmesinin genellenebilirliği artırdığını göstermektedir. Sonuç olarak bu tez, Azerbaycan dili için güvenilir, açıklanabilir ve yeniden kullanılabilir bir duygu analizi araştırma çerçevesi sunmakta; aynı zamanda düşük kaynaklı ve morfolojik açıdan zengin diğer diller için de yöntemsel bir temel önermektedir.

Haziran 2026, 100 sayfa

Anahtar Kelimeler: duygu analizi, Azerbaycan dili, düşük kaynaklı diller, morfoloji-duyarlı modelleme, açıklanabilirlik, alanlar arası sağlamlık

ABSTRACT

PhD Thesis

DEVELOPING A SENTIMENT ANALYSIS FRAMEWORK FOR AZERBAIJANI

Yusuf Evren AYKAÇ

Ankara University
Graduate School of Natural and Applied Science
Department of Computer Engineering

Supervisor: Prof. Dr. Refik SAMET

This dissertation treats sentiment analysis for Azerbaijani not as a standalone classification task, but as a broader research problem requiring resource construction, morphology-aware modeling, cross-domain robustness, and explainability. To this end, an explainable, morphology-aware, and domain-robust sentiment analysis framework was developed for Azerbaijani. The study defined the problem space and data requirements, carried out data collection, cleaning, and preprocessing, developed a polarity lexicon called *SentiAzNet*, investigated the integration of lexicon-based knowledge into aspect-based sentiment analysis in the finance and business domain, designed a multi-domain hybrid architecture called *MahSA*, and learned sentence representations that are more resistant to domain shift. All components were evaluated jointly through explainability analyses, error analyses, and comparative experiments. The findings show that the proposed lexicon provides reliable symbolic prior knowledge, that lexicon-guided contextual modeling strengthens both predictive performance and interpretability, that the hybrid architecture yields more balanced results especially for the neutral class and challenging domains, and that domain-robust representation learning improves generalization across domains. Overall, the dissertation offers a reliable, explainable, and reusable research framework for sentiment analysis in Azerbaijani and provides a methodological basis for other low-resource and morphologically rich languages.

June 2026, 100 pages

Keywords: sentiment analysis, Azerbaijani language, low-resource languages, morphology-aware modeling, explainability, cross-domain robustness

TEŞEKKÜR

Bu tez çalışmasının her aşamasında bilgisi, yönlendirmeleri ve yapıcı eleştirileriyle yanımda olan değerli danışmanım Prof. Dr. Refik SAMET'e en içten teşekkürlerimi sunarım. Araştırmanın bilimsel çerçevesinin şekillenmesinde, yöntemsel tercihlerin olgunlaşmasında ve bulguların bütüncül bir anlatıya kavuşmasında kendisinin yol göstericiliği belirleyici olmuştur.

Tez sürecinin her aşamasında değerli görüşleri, birikimi ve yapıcı katkılarıyla araştırmama yön veren Dr. Öğretim Üyesi Merve ÖZKAN'a ayrıca teşekkür ederim. Veri seti üzerinde yürütülen titiz çalışmalarda özverili emekleriyle yer alan Samir, Hüseyin, Ulviye, Revan ve Osman'a; ayrıca bu süreçte iş birliği içinde çalışma fırsatı bulduğum tüm araştırmacılara ve meslektaşlarıma içten teşekkürlerimi sunarım.

Değerli görüş ve önerileriyle tezin gelişimine katkı sağlayan Doç. Dr. Yılmaz AR ve Prof. Dr. Tansel ÖZYER hocalarıma; ayrıca akademik süreçlere ilişkin deneyim ve birikimini benimle paylaşarak yol gösteren Prof. Dr. Fatih Vehbi ÇELEBİ hocama teşekkürlerimi sunarım.

Bu tezin önemli bir bölümü, TÜBİTAK tarafından desteklenen 224N535 numaralı proje kapsamında geliştirilen araştırma altyapısından yararlanılarak yürütülmüştür. Sağladığı bu değerli destek için TÜBİTAK'a ve projede yer alan iş birliği kurumlarına teşekkürlerimi sunarım.

Araştırma sürecinin uzun ve çoğu zaman zorlu yolculuğunda maddi ve manevi desteklerini her an hissettiğim aileme sonsuz teşekkür borçluyum. Bu yolun her adımında yanımda olan, bana güç ve huzur veren sevgili eşime ve hayatımın en kıymetli armağanı biricik kızım Ada'ya en derin şükranlarımı sunarım. Son olarak, varlıklarıyla hayatıma anlam ve neşe katan kıymetli dostlarım Faruk, Gencer, Pelin, Seba ve Arda'ya da içtenlikle teşekkür ederim.

Yusuf Evren AYKAÇ

Ankara, Haziran 2026

İÇİNDEKİLER

TEZ ONAY SAYFASI	
ETİK	i
ÖZET	ii
ABSTRACT	iii
TEŞEKKÜR	iv
SİMGELER VE KISALTMALAR DİZİNİ	vii
ŞEKİLLER DİZİNİ	x
ÇİZELGELER DİZİNİ	xi
1. GİRİŞ	1
1.1 Problem Tanımı	2
1.2 Araştırmanın Önemi ve Kapsamı	3
1.3 Tezin Amacı	5
1.4 Tezin Temel Savı	6
1.5 Araştırma Soruları	6
1.6 Tezin Özgün Katkıları	7
1.7 Bu Tezin Literatürdeki Konumu	8
1.8 Tezin Organizasyonu	9
2. İLGİLİ ÇALIŞMALAR VE LİTERATÜR DEĞERLENDİRMESİ	10
2.1 Duygu Analizi Kavramı ve Düzeyleri	10
2.2 Düşük Kaynaklı ve Morfolojik Açından Zengin Dillerde Duygu Analizi	12
2.3 Kutupluluk Sözlükleri ve Sözlük Tabanlı Yaklaşımlar	14
2.4 Çok Dilli Temsiller, Bağlamsal Modeller ve Hibrit Yaklaşımlar	17
2.5 Aspekt-Tabanlı Duygu Analizi ve Finansal Duygu Analizi	21
2.6 Alan Uyarlaması, Alanlar Arası Sağlamlık ve Temsil Öğrenmesi	23
2.7 Açıklanabilir Yapay Zekâ ve Hata Analizi	25
2.8 Azerbaycan Dili için Duygu Analizi Çalışmaları	26
2.9 Literatürdeki Boşluk ve Bu Tezin Konumu	27
3. AZERBAJCANDİLİ İÇİN DUYGU ANALİZİ ÇERÇEVESİ GELİŞTİRME METODOLO- JİSİ	29
3.1 Yedi Aşamalı Metodoloji ve Genel İş Akışı	30
3.2 Çalışmada Kullanılan Derlem ve Veri Kümeleri	31
3.3 Veri Toplama, Temizleme ve Ön İşleme Süreci	33
3.4 SentiAzNet Kutupluluk Sözlüğünün Geliştirilmesi	33
3.4.1 Tohum sözlüklerin derlenmesi	34
3.4.2 Çeviri, geri çeviri ve anlamsal doğrulama	35
3.4.3 Manuel doğrulama ve sözlük kalitesinin güvence altına alınması	35
3.4.4 Özellik çıkarımı ve sözlük tabanlı doğrulama modeli	36
3.4.5 Dış geçerleme	37
3.5 Aspekt-Tabanlı Duygu Analizinde Sözlük Enjeksiyonu	38
3.5.1 Problem tanımı ve veri kümesinin özellikleri	39
3.5.2 Hedef-koşullu girdi gösterimi ve temel kodlayıcı	40
3.5.3 Token düzeyinde sözlük hizalama ve alan düzeltmeleri	41
3.5.4 Özellik birleştirme ile sözlük enjeksiyonu	42
3.5.5 Dikkat önyargısı ile sözlük enjeksiyonu	42

3.5.6	Eğitim ve değerlendirme protokolü	43
3.5.7	Açıklanabilirlik ve hata analizi	44
3.6	MahSA Çok Alanlı Hibrit Mimarisi	45
3.6.1	Tasarım problemi ve kurucu ilkeler	45
3.6.2	Çok alanlı derlemin oluşturulması ve etiketleme	46
3.6.3	Ön işleme ve morfoloji-duyarlı normalizasyon	47
3.6.4	Özellik blokları ve temsil notasyonu	48
3.6.5	Lexicon scoring ve morfoloji-duyarlı kutup düzenleme	50
3.6.6	Bağlama kutupluluk yayılımı: wPPMI	50
3.6.7	Karakter n-gramları, bağlamsal kodlayıcı ve morfoloji blokları	51
3.6.8	Füzyon katmanı ve duygu başlığı	52
3.6.9	Nötr sınıf için alan-duyarlı kalibrasyon	52
3.6.10	Eğitim ayrıntıları ve hiperparametreler	52
3.7	Alanlar Arası Sağlam Cümle Temsillerinin Öğrenilmesi	53
3.7.1	Problem tanımı ve veri kurgusu	54
3.7.2	Kaynak seçimi, filtreleme ve dil denetimi	56
3.7.3	Temel kodlayıcı ve cümle havuzlama	56
3.7.4	Alan-uyarlamalı ön eğitim ve alan-karşıtı katman	57
3.7.5	Kutupluluk ve morfoloji bilgisi ile temsillerin koşullandırılması	58
3.7.6	Denetimli karşıtsal düzenleştirme ve toplam amaç fonksiyonu	59
3.7.7	Deneyisel karşılaştırma aileleri ve değerlendirme protokolleri	60
3.8	Açıklanabilirlik, Hata Analizi ve Bütünleşik Değerlendirme Protokolü	61
4.	BULGULAR VE TARTIŞMA	63
4.1	Deneyisel Gerçekleme Ortamı ve Bulguların Sunum Çerçevesi	63
4.2	SentiAzNet Kutupluluk Sözlüğüne İlişkin Bulgular	65
4.3	Finans ve İş Dünyası ABSA Bulguları ve Tartışma	67
4.4	MahSA Çok Alanlı Hibrit Mimarisine İlişkin Bulgular	70
4.5	Alanlar Arası Sağlam Cümle Temsillerine İlişkin Bulgular	73
4.6	Bütünleşik Tartışma	76
5.	GENEL DEĞERLENDİRME	79
5.1	Tezin Genel Değerlendirmesi	79
5.2	Araştırma Sorularına Verilen Yanıtlar	81
5.3	Tezin Bilimsel ve Uygulamalı Katkılarının Sentezi	83
5.4	Sınırlılıklar	85
5.5	Gelecekte Yapılabilecek Çalışmalar ve Öneriler	86
6.	SONUÇ	88
	KAYNAKLAR	91
	ÖZGEÇMİŞ	99

SİMGELER VE KISALTMALAR DİZİNİ

Simgeler

w_i	: Girdi metnindeki i . belirteç
s_i	: Belirteç düzeyindeki kutupluluk skoru
x	: Girdi yorum/metin
s	: Cümle
a_t	: Hedef aspekt terimi
a_c	: Aspekt kategorisi
y	: Duygu etiketi
e	: Kaynak İngilizce lemma
$\mathcal{A}(e)$	: e için üretilen Azerbaycanca adaylar kümesi
$\mathcal{A}^*(e)$	: Anlamsal doğrulamadan geçen aday küme
$BT(a)$	: a adayının geri çevrilmiş İngilizce karşılığı
$Syn(\cdot)$	: Synset / eşanlam kümesi
z_i	: SentiAzNet doğrulama hattındaki bileşik özellik vektörü
$\hat{y}_i$	: Modelin öngördüğü sınıf etiketi
$e_i^{(emo)}$	: Duygu bayrakları vektörü
$m_i^{(morf)}$	: Morfolojik göstergeler vektörü
$u_i^{(LaBSE)}$	: LaBSE tabanlı indirgenmiş anlamsal temsil
$f_i^{(azWaC)}$	: azWaC derleminden türetilen log-frekans değeri
$t_i^{(char)}$	: Karakter n-gram TF-IDF vektörü
τ_{lex}	: Sözlük eşik değeri
γ_{neg}	: Olumsuzluk kuvvet katsayısı
γ_{int}	: Pekiştirme kuvvet katsayısı
α	: Dikkat önyargısı / wPPMI yumuşatma katsayısı
δ	: Küresel nötr kalibrasyon marjı
δ_d	: Alan-özü nötr kalibrasyon marjı
λ_L	: Sözlük özellik bloğu ağırlığı
λ_ψ	: wPPMI özellik bloğu ağırlığı
λ_C	: Karakter n-gram bloğu ağırlığı
λ_E	: Bağlamsal gömme bloğu ağırlığı
λ_M	: Morfoloji bloğu ağırlığı
p_y	: Sınıf olasılığı
w	: Bağlam penceresi genişliği
d_k	: Dikkat anahtar vektörü boyutu
N	: Toplam eğitim örneği sayısı
N_c	: İlgili sınıftaki örnek sayısı
N_d	: İlgili alandaki örnek sayısı

w_c	: Sınıf ağırlığı
w_d	: Alan-ağırlıklı örnekleme / kayıp katsayısı
$T_{\max}$	: Azami alt-sözcük uzunluğu
k	: Rastgele tohum / tekrar sayısı
X	: Hedef-koşullu girdi dizisi
$\text{tok}(\cdot)$	: Alt-sözcük parçalama işlemi
h_{cls}	: Kodlayıcının sınıflandırma temsili
W_o	: Çıkış katmanı ağırlık matrisi
b_o	: Çıkış katmanı bias vektörü
$F(x)$	: Füzyon sonrası birleşik gösterim
$L(x)$	: Sözlük tabanlı özellik vektörü
$C(x)$	: Karakter n-gram özellik vektörü
$E(x)$	: Bağlamsal gömme vektörü
$M(x)$	: Morfoloji-duyarlı özellik vektörü
$\psi(x)$	: wPPMI tabanlı bağlamsal kutupluluk vektörü
Δ_{rel}	: Göreli iyileşme oranı
S_{prop}	: Önerilen modelin skoru
S_{base}	: Taban modelin skoru
$\overline{F1}_{\text{LODO}}$	: Ortalama LODO Macro-F1 değeri
$F1_d$	: İlgili alan için Macro-F1 değeri
$\mathcal{D}$	: Alan kümesi

Kısaltmalar

ABSA	Aspect-Based Sentiment Analysis
AUC	Area Under the Curve
Az-RoBERTa	Azerbaijani RoBERTa
CE	Cross-Entropy
CLS	Classification Token
CNN	Convolutional Neural Network
DANN	Domain-Adversarial Neural Network
DAPT	Domain-Adaptive Pretraining
DL	Deep Learning
F1	F1 Skoru
GRU	Gated Recurrent Unit
HGBT	Histogram-based Gradient Boosting Tree
LaBSE	Language-agnostic BERT Sentence Embedding
LLM	Large Language Model
LODO	Leave-One-Domain-Out
LSTM	Long Short-Term Memory
MahSA	Morpheme-Aware Hybrid Sentiment Analysis
MCC	Matthews Correlation Coefficient
mBERT	Multilingual BERT
MLP	Multi-Layer Perceptron
NLP	Natural Language Processing
NMI	Normalized Mutual Information
NN	Nearest Neighbor
PCA	Principal Component Analysis
PEFT	Parameter-Efficient Fine-Tuning
PPMI	Positive Pointwise Mutual Information
QLoRA	Quantized Low-Rank Adaptation
RNN	Recurrent Neural Network
ROC-AUC	Receiver Operating Characteristic – Area Under the Curve
SHAP	SHapley Additive exPlanations
SentiAzNet	Sentiment Lexicon for Azerbaijani
SupCon	Supervised Contrastive Learning
SVD	Singular Value Decomposition
TF-IDF	Term Frequency-Inverse Document Frequency
wPPMI	Weighted Positive Pointwise Mutual Information
XAI	Explainable Artificial Intelligence
XLM-R	XLM-RoBERTa

ŞEKİLLER DİZİNİ

3.1	Tez kapsamında izlenen yedi aşamalı metodolojik iş akışı	30
3.2	SentiAzNet geliştirme sürecinin nicel olarak özetlenmiş akışı	39
3.3	Finans ve iş dünyası alanındaki ABSA alt probleminde izlenen sözlük enjeksiyonu ve açıklanabilirlik hattı	43
3.4	MahSA mimarisinin genel akışı ve temel bilgi blokları	49
3.5	Alanlar arası sağlam duygu temsilleri için izlenen genel iş akışı	58
4.1	Tezin deneysel kanıt zincirinin özet görünümü	65
4.2	XLM-R ince ayar ve MahSA için tez kapsamında elde edilen satır-normalize karışıklık matrisleri	71
4.3	Alanlar arası sağlam temsil öğrenmesi çalışmasında model ailesine göre alan-içi ve LODO ortalama Macro-F1 sonuçları	75
5.1	Tezin kavramsal ve deneysel katkılarının birbirini besleyen bütüncül akışı	80

ÇİZELGELER DİZİNİ

2.1	Duygu analizi düzeylerinin temel özellikleri	12
2.2	Kutupluluk sözlüğü temelli başlıca yaklaşım aileleri	16
2.3	Duygu analizinde öne çıkan başlıca temsil aileleri	20
2.4	Aspekt-tabanlı ve finansal duygu analizinde öne çıkan güçlükler	22
2.5	Alanlar arası sağlamlık için başlıca yöntem aileleri	24
3.1	Tez metodolojisinin yedi temel aşaması	31
3.2	Tez kapsamında kullanılan başlıca materyal ve veri kaynakları	32
3.3	<i>SentiAzNet</i> hattında kullanılan temel tohum sözlük kaynakları	34
3.4	<i>SentiAzNet</i> doğrulama hattında kullanılan özellik grupları	37
3.5	<i>SentiAzNet</i> oluşturma hattının nicel özeti	38
3.6	Finans ve iş dünyası ABSA veri kümesinin temsilî istatistikleri	40
3.7	Finans ve iş dünyası alanındaki ABSA alt probleminde karşılaştırılan sözlük en- jeksiyon varyantları	43
3.8	Finans ve iş dünyası alanındaki ABSA alt probleminde kullanılan deneysel protokol	44
3.9	MahSA çok alanlı derleminde alan dağılımı	46
3.10	MahSA derleminde sınıf dağılımı ve kısa derlem istatistikleri	47
3.11	MahSA etiketleme sürecinde sınıf bazlı ve genel uyum	47
3.12	MahSA'da kullanılan temel özellik blokları	49
3.13	MahSA için kullanılan temel eğitim ayarları	53
3.14	Alanlar arası sağlam temsil öğrenmesi aşamasında kullanılan derlemin alan dağılımı	55
3.15	Etiket dağılımı ve etiketleme güvenilirliği için temel göstergeler	55
3.16	Alanlar arası sağlam temsil öğrenmesi bölümünde karşılaştırılan model aileleri	61
4.1	<i>SentiAzNet</i> iç doğrulama sonuçları	66
4.2	<i>SentiAzNet</i> kararlılık ve dış geçerlik bulguları	66
4.3	Finans ve iş dünyası ABSA ana sonuçları	68
4.4	Finans ve iş dünyası ABSA ablasyon ve alternatif eğitim sonuçları	68
4.5	Finans ve iş dünyası ABSA'da gözlenen başlıca hata tipleri	69
4.6	MahSA ve temel karşılaştırma modelleri için alan bazlı Macro-F1 sonuçları	70
4.7	Seçilmiş güçlü modeller için sınıf bazlı ve genel sonuçlar	71
4.8	MahSA için ablasyon sonuçları	72
4.9	MahSA için LODO (bir alanı dışarıda bırakma) sonuçları	72
4.10	Alanlar arası sağlam temsil öğrenmesi çalışmasında alan-içi Macro-F1 sonuçları ..	74
4.11	Alanlar arası sağlam temsil öğrenmesi çalışmasında LODO sonuçları	74
4.12	Alanlar arası sağlam temsil öğrenmesi çalışmasında ablasyon ve gömme-probe sonuçları	75

1. GİRİŞ

Bu bölümde, Azerbaycan dili bağlamında duygu analizi probleminin temel boyutları ortaya konmakta; araştırmanın önemi, kapsamı, amacı, temel savı, araştırma soruları ve özgün katkıları sistematik biçimde sunulmaktadır. Son olarak tezin literatürdeki konumu ve bölüm yapısı özetlenmektedir.

Duygu analizi, kullanıcı yorumları, sosyal medya gönderileri, haber yorumları ve hizmet geri bildirimleri gibi metinlerde ifade edilen öznel yönelimlerin, değerlendirmelerin ve duygusal tutumların otomatik olarak belirlenmesini amaçlayan temel bir doğal dil işleme alanıdır. Görüş madenciliği ve kutupluluk sınıflandırması ile yakın ilişki içinde gelişen bu alan, müşteri deneyiminin izlenmesinden itibaren yönetimine, karar destek süreçlerinden kamusal eğilimlerin çözümlenmesine kadar geniş bir uygulama alanına sahiptir (Taboada vd. 2011; Shaik vd. 2023).

Yüksek kaynaklı diller için geliştirilen kutupluluk sözlükleri, büyük ölçekli veri kümeleri ve güçlü ön eğitilmiş dil modelleri sayesinde duygu analizi literatürü görece olgun bir araştırma çizgisine ulaşmıştır (Baccianella vd. 2010; Devlin vd. 2019; Conneau vd. 2020). Buna karşılık düşük kaynaklı ve morfolojik açıdan zengin dillerde aynı düzeyde kaynak, yöntem ve değerlendirme altyapısının oluşmadığı; veri kıtlığı, alan uyumsuzluğu, değerlendirme parçalanmışlığı ve çapraz dilli aktarım sınırlılıklarının hâlâ belirleyici olduğu görülmektedir (Aliyu vd. 2024).

Azerbaycan dili bu güçlükleri görünür kılan önemli örneklerden biridir. Ekleme yapı, alt-sözcük düzeyinde aşırı parçalanma riski, kullanıcı üretilen metinlerdeki yazım çeşitliliği ve dile özgü temel kaynakların sınırlılığı, genel amaçlı çok dilli modellerin tek başına yeterli bir çözüm sunmasını zorlaştırmaktadır (Alizada vd. 2024; Rust vd. 2021). Bu durum, Azerbaycan dili için duygu analizinin yalnızca daha güçlü bir sınıflandırıcı seçme meselesi olmadığını; aynı zamanda kaynak geliştirme, morfoloji-duyarlı modelleme, alanlar arası genelleme ve açıklanabilirlik eksenlerinde birlikte ele alınması gereken bütüncül bir araştırma problemi olduğunu göstermektedir.

Bu tez çalışmasında Azerbaycan dili için duygu analizi, yalnızca daha yüksek sınıflandırma başarımı elde etme problemi olarak değil; dil kaynağı geliştirme, morfoloji-duyarlı modelleme, alan kaymasına karşı dayanıklılık ve açıklanabilirlik boyutlarını birlikte içeren bütüncül bir araştırma

problemi olarak ele alınmaktadır. Bu çerçeve, kutupluluk sözlüğü geliştirme, alan-özü aspekt-tabanlı modelleme, çok alanlı hibrit duygu analizi ve alanlar arası sağlam temsil öğrenmesi katmanları üzerinden somutlaştırılmıştır (Aykaç vd. 2025, 2026b; Aykaç vd. 2026; Aykaç vd. 2026a).

1.1 Problem Tanımı

Bu tezde ele alınan duygu analizi problemi, Azerbaycan dili bağlamında dört temel güçlük ekseninden tanımlanmaktadır. Bu eksenler yalnızca yöntem seçimini değil, aynı zamanda veri tasarımını, değerlendirme stratejisini ve model kararlarının nasıl yorumlanacağını da doğrudan etkilemektedir.

İlk eksen, **kaynak yetersizliğidir**. Güvenilir bir kutupluluk sözlüğünün, geniş ölçekli ve elle etiketlenmiş veri kümelerinin ve alanlar arası karşılaştırmaya elverişli ortak değerlendirme düzeneklerinin sınırlı olması, geliştirilen modellerin hem doğruluğunu hem de yeniden üretilebilirliğini zayıflatmaktadır. Bu eksiklik, özellikle nötr sınıfın güvenilir biçimde modellenmesini ve farklı alanlara genellenebilir sonuçlar elde edilmesini güçleştirmektedir (Aliyu vd. 2024; Alizada vd. 2024; Aykaç vd. 2025; Aykaç vd. 2026).

İkinci eksen, **Azerbaycan dilinin eklemeli ve morfolojik açıdan zengin yapısıdır**. Kutupluluk çoğu durumda yalnızca sözcük kökü üzerinden değil; olumsuzluk, yoksunluk, pekiştirme ve biçimbirimsel değişim örüntüleri üzerinden kurulmaktadır. Bu nedenle yalnızca yüzeysel sözcük eşleşmelerine ya da genel amaçlı alt-sözcük temsillerine dayanan yaklaşımlar, gerçek duygu yönelimini kararlı biçimde yakalayamamaktadır. Özellikle kullanıcı üretimi metinlerde görülen yazım çeşitliliği, kısaltmalar, uzatmalar ve düzensiz biçimler bu sorunu daha da derinleştirmektedir (Dehkharghani vd. 2017; Altınel Girgin vd. 2024; Aykaç vd. 2026).

Üçüncü eksen, **alan bağımlılığı ve alan kaymasıdır**. Bir sözcüğün ya da ifadenin taşıdığı kutupluluk; finans, kamu hizmetleri, teknoloji, perakende veya sosyal yaşam gibi farklı alanlarda aynı biçimde ortaya çıkmamaktadır. Aynı cümle içinde birden fazla hedefe ilişkin farklı duygu yönelimlerinin bulunabilmesi, özellikle aspekt-tabanlı duygu analizi senaryolarında, cümle düzeyindeki tek etiketli yaklaşımların yetersiz kalmasına yol açmaktadır (Hua vd. 2024). Finansal metinlerde genel dildeki kutupluluk değerlerinin alan bağlamında farklı işleyebildiği de Loughran ve McDon-

ald (2011) tarafından gösterilmiştir. Düşük kaynaklı dillerde bu sorun daha görünür hale gelmekte ve alanlar arası genelleme yeteneği sınırlı kalmaktadır (Aliyu vd. 2024; Aykaç vd. 2026).

Dördüncü eksen ise **açıklanabilirlik gereksinimidir**. Güçlü bağlamsal modeller yüksek başarı sağlayabilse bile, hangi sözcüklerin, hangi morfolojik işaretlerin veya hangi sözlük sinyallerinin karara etki ettiği çoğu zaman açık değildir. Özellikle nötr-olumsuz belirsizliği, ironi, karma duygu, örtük değerlendirme ve kod değiştirimi gibi olgularda yalnızca çıktı başarımına bakmak yeterli olmamakta; model davranışının açıklanması ve hata türlerinin sistematik biçimde incelenmesi gerekmektedir. Nitekim dikkat ağırlıklarının tek başına yeterli açıklama sağlamayabileceği Jain ve Wallace (2019) tarafından gösterilmiştir; bu nedenle SHAP ve benzeri çoklu açıklama kanallarının birlikte değerlendirilmesi önem kazanmaktadır (Diwali vd. 2024; Lundberg ve Lee 2017).

Bu nedenle tezde ele alınan temel araştırma problemi, Azerbaycan dili için düşük kaynak koşullarında çalışabilen, morfoloji-duyarlı, sözlük destekli, açıklanabilir ve alanlar arası genellenebilir bir duygu analizi çerçevesinin nasıl geliştirilebileceğidir. Başka bir ifadeyle bu tez, yalnızca “hangi model daha yüksek skor üretir?” sorusunu değil; “hangi bilgi türleri, hangi koşullarda, hangi hata bölgelerini azaltır?” sorusunu da cevaplamayı amaçlamaktadır.

1.2 Araştırmanın Önemi ve Kapsamı

Bu tez çalışmasının önemi, Azerbaycan dili için duygu analizini tekil bir sınıflandırma uygulaması olmaktan çıkarıp; dil kaynağı geliştirme, modelleme, değerlendirme ve yorumlama katmanlarını birlikte ele alan sistematik bir araştırma alanı olarak konumlandırmasından kaynaklanmaktadır. Mevcut durumda Azerbaycan dili için kamuya açık kutupluluk sözlüklerinin, geniş ölçekli ve elle etiketlenmiş çok alanlı veri kümelerinin ve alan kaymasını görünür kılan ortak değerlendirme düzeneklerinin sınırlı olması, hem akademik ilerlemeyi hem de gerçek kullanım senaryolarına dönük uygulamaları yavaşlatmaktadır (Aliyu vd. 2024; Alizada vd. 2024; Aykaç vd. 2025; Aykaç vd. 2026). Bu nedenle tez, yalnızca daha yüksek doğruluk üretmeyi değil; hangi dil kaynağının eksik olduğunu, bu kaynağın nasıl üretileceğini ve üretilen bilginin çağdaş modellere nasıl aktarılacağını da açıklığa kavuşturmayı amaçlamaktadır.

Çalışmanın yöntemsel önemi de belirgindir. Azerbaycan dili; eklemeli yapısı, kutupluluğu

dönüştürebilen olumsuzluk ve yoksunluk biçimbirimleri, kullanıcı üretimli metinlerdeki yazım çeşitliliği ve alanlar arası sözlüksel kaymalar nedeniyle, yüksek kaynaklı diller için geliştirilmiş salt veri-merkezli yaklaşımların doğrudan aktarılmasıyla çözülebilecek bir problem sunmamaktadır. Özellikle nötr sınıfın güvenilir biçimde modellenmesi, finans ve iş dünyası gibi alanlarda aspekt düzeyinde kutupluluğun ayrıştırılması ve kamu hizmetleri gibi küçük fakat zor alanlarda kararlı sonuçlar elde edilmesi, sözlük bilgisi, morfoloji-duyarlı yüzey ipuçları ve bağlamsal temsillerin birlikte değerlendirilmesini gerekli kılmaktadır (Conneau vd. 2020; Rust vd. 2021; Aykaç vd. 2025; Aykaç vd. 2026). Bu bakımdan tez, sembolik ve sinirsel bilgi kaynaklarını birbirine alternatif değil, birbirini tamamlayan bileşenler olarak ele almaktadır.

Çalışmanın pratik önemi ise doğrudan uygulama alanlarından kaynaklanmaktadır. Sosyal medya izlemesi, kullanıcı yorumlarının otomatik çözümlemesi, finans ve iş dünyası metinlerinde hedefe bağlı duygu analizi, kamu hizmetlerine ilişkin geri bildirimlerin değerlendirilmesi ve pazar araştırması gibi senaryolar, Azerbaycan dili için güvenilir ve açıklanabilir duygu analizi araçlarına ihtiyaç duymaktadır (Shaik vd. 2023). Bu tezde geliştirilen çerçeve, bu tür uygulamalar için yalnızca bir model önerisi sunmamakta; aynı zamanda dil kaynağı, veri, yöntem ve açıklama katmanlarını birlikte sağlayan araştırma temelli bir altyapı oluşturmaktadır.

Bu tezin kapsamı, Azerbaycan dili için **metin tabanlı** duygu analizi problemleriyle sınırlandırılmıştır. Çalışma; kamuya açık veya araştırma amaçlı derlenen yazılı veriler üzerinde kutupluluk sözlüğü geliştirme, alan-özü aspekt-tabanlı duygu analizi, çok alanlı ve üç sınıflı yorum düzeyinde duygu analizi, alanlar arası sağlam temsil öğrenmesi ve açıklanabilirlik çözümlemelerini kapsamaktadır. Bu çerçevede ön işleme, normalizasyon, karakter ve alt-sözcük düzeyli temsiller, morfoloji-duyarlı yüzey ipuçları, kutupluluk sinyalleri, bağlamsal gömmeler ve hibrit füzyon yaklaşımları birlikte ele alınmaktadır (Conneau vd. 2020; Aykaç vd. 2025; Aykaç vd. 2026).

Buna karşılık tez; konuşma duygusu tanıma, görüntü veya çok kipli duygu analizi, diyalog sistemleri ya da genel amaçlı üretici yapay zekâ uygulamalarını doğrudan hedeflememektedir. Benzer biçimde çalışma, Azerbaycan dilinin tüm biçimbilimsel düzeylerini tam bir çözümleyici ile modelleme iddiası da taşımamaktadır; bunun yerine duygu yönelimine doğrudan etki eden olumsuzluk, yoksunluk, pekiştirme, uzatma ve benzeri işaretlere odaklanan pragmatik bir morfoloji-duyarlı yaklaşım benimsenmiştir (Aykaç vd. 2026). Ayrıca bu tez, tipolojik yakınlık nedeniyle diğer

Türk dilleri için yararlı bir yöntemsel perspektif sunsa da, bu dillere doğrudan deneysel genelleme yapma iddiasında değildir. Tezde yer alan uygulama ve platform bileşenleri de, bir üretim ortamı yazılımının tüm mühendislik ayrıntılarından ziyade, bilimsel katkıyı taşıyan dil kaynağı, modelleme ve değerlendirme çekirdeğini görünür kılan bir gerçekleştirim bağlamı olarak ele alınmıştır.

Bu sınırlar içinde tez, Azerbaycan dili için duygu analizini parçalı alt problemlere ayırmak yerine; kaynak geliştirme, model tasarımı, deneysel doğrulama ve açıklanabilirlik eksenlerini bir arada değerlendiren bütüncül bir araştırma programı önermektedir. Bu yönüyle kapsam, hem araştırma problemini sınırlandırmakta hem de savunmada beklenti yönetimini daha açık hale getirmektedir.

1.3 Tezin Amacı

Bu tez çalışmasının temel amacı, başta Azerbaycan dili olmak üzere kısıtlı kaynaklı diller için dilin morfolojik özelliklerini, alanlar arası kullanım farklılıklarını ve açıklanabilirlik gereksinimini birlikte dikkate alan bütüncül bir duygu analizi çerçevesi geliştirmektir. Bu amaç doğrultusunda tez kapsamında yalnızca bir sınıflandırma modeli önerilmemiş; aynı zamanda Azerbaycan dili için bir kutupluluk sözlüğü oluşturulmuş, alan-özü ve çok alanlı veri kümeleri üzerinde sözlük bilgisinin çağdaş bağlamsal modellerle nasıl bütünleştirilebileceği araştırılmış, morfoloji-duyarlı hibrit mimariler geliştirilmiş ve öğrenilen temsillerin alan kaymasına karşı dayanıklılığı incelenmiştir (Aykaç vd. 2025, 2026b; Aykaç vd. 2026; Aykaç vd. 2026a).

Bu amaç doğrultusunda benimsenen yaklaşım, metodolojik ayrıntılara burada girmeden ifade edilecek olursa, kaynak geliştirme ile model geliştirmeyi aynı omurga içinde birleştiren katmanlı bir çözüm stratejisine dayanmaktadır. Önce Azerbaycan dili için yeniden kullanılabilir sembolik ön bilgi üretilmekte; ardından bu bilginin alan-özü görevlerde bağlamsal modellere nasıl aktarılabilirliği sınanmakta; daha sonra çok alanlı ve morfoloji-duyarlı hibrit mimari kurulmakta; son aşamada ise alan kaymasına karşı daha sağlam temsiller ve açıklanabilirlik düzenekleri birlikte değerlendirilmektedir (Aykaç vd. 2025, 2026b; Aykaç vd. 2026; Aykaç vd. 2026a). Böylece çalışma, kaynak üretiminden modellemeye, açıklanabilirlikten alanlar arası genellemeye kadar uzanan katmanlı bir araştırma hattı sunmaktadır.

1.4 Tezin Temel Savı

Bu tezde savunulan temel görüş şudur: Azerbaycan dili gibi düşük kaynaklı, eklemeli ve morfolojik açıdan zengin bir dilde duygu analizi başarımı, yalnızca büyük ölçekli çok dilli dil modellerine dayanan yaklaşımlarla değil; kutupluluk sözlükleri gibi sembolik bilgi kaynaklarının, morfoloji-duyarlı yüzey ipuçlarının ve bağlamsal sinirsel temsillerin bütünleşik biçimde kullanılmasıyla daha güvenilir, daha açıklanabilir ve alanlar arası daha sağlam hale getirilebilir. Bu sav, tez boyunca kaynak geliştirme, alan-özü modelleme, çok alanlı hibrit tasarım ve temsil öğrenmesi basamaklarında kademeli olarak sınanmaktadır.

1.5 Araştırma Soruları

Bu problem çerçevesinde tezde aşağıdaki araştırma sorularına yanıt aranmıştır. Sorular; kaynak geliştirme, sembolik bilginin modele aktarımı, çok alanlı hibrit modelleme, nötr sınıf ve düşük veri rejimleri, temsil öğrenmesi ve açıklanabilirlik eksenlerini kapsamakta; bunlara verilen yanıtlar ilgili yöntem ve bulgu bölümlerinde sistematik olarak sınanmaktadır.

- AS1** Azerbaycan dili için, diğer dillere ait zengin kaynaklar, geri çeviri, anlamsal doğrulama ve uzman etiketlemesi birlikte kullanılarak güvenilir ve yeniden kullanılabilir bir kutupluluk sözlüğü oluşturulabilir mi?
- AS2** Oluşturulan kutupluluk sözlüğünden türetilen sinyaller, bağlamsal transformer tabanlı modellere eklendiğinde Azerbaycan dili için duygu analizi başarımına anlamlı katkı sağlayabilir mi?
- AS3** Kutupluluk sözlüğü bilgisi, karakter düzeyi ipuçları, bağlamsal gömme temsilleri ve morfoloji-duyarlı yüzey özellikleri birlikte kullanıldığında, çok alanlı ve üç sınıflı duygu analizi probleminde tek başına bağlamsal modellere göre daha güçlü sonuçlar elde edilebilir mi?
- AS4** Önerilen hibrit yaklaşım, özellikle nötr sınıfın modellenmesinde ve veri miktarı görece sınırlı alanlarda, mevcut güçlü taban modellere göre daha kararlı bir performans sergileyebilir mi?
- AS5** Alan-özü kutupluluk farklılıkları ve alan kayması dikkate alındığında, duyguya duyarlı ve alanlar arası daha sağlam cümle temsilleri öğrenmek mümkün müdür?
- AS6** SHAP, dikkat örüntüleri, sözlük eşleşmeleri ve hata çözümlemeleri birlikte kullanıldığında,

model kararlarının açıklanabilirliği artırılabilir ve başarımın hangi dilsel işaretlerden beslendiği sistematik biçimde ortaya konabilir mi?

1.6 Tezin Özgün Katkıları

Bu tez çalışmasının özgün katkıları, giriş bölümünde tanımlanan araştırma sorularıyla doğrudan ilişkili olacak biçimde aşağıda özetlenmiştir. Katkıları; kaynak/veri katkısı, yöntem katkısı, deneysel katkı ve açıklanabilirlik katkısı katmanlarını birlikte içermektedir.

K1 Azerbaycan dili için özgün bir kutupluluk sözlüğü geliştirilmiştir. Diğer dillere ait zengin sözlük kaynakları bir araya getirilmiş, 9.614 benzersiz İngilizce lemma üzerinden başlayan aday havuz; çeviri, geri çeviri, WordNet tabanlı anlamsal doğrulama ve insan etiketlemesi süreçleriyle daraltılarak 2.482 Azerbaycan dili girdisinden oluşan *SentiAzNet* kaynağı oluşturulmuştur (Aykaç vd. 2025). Böylece Azerbaycan dili için duygu analizi çalışmalarında kullanılabilecek, güvenilir ve yeniden kullanılabilir bir temel sembolik kaynak sunulmuştur.

K2 Kutupluluk sözlüğü bilgisinin bağlamsal modellere entegrasyonu için alan-özgü ve açıklanabilir bir yaklaşım önerilmiştir. Finans ve iş dünyası metinleri üzerinde gerçekleştirilen aspekt-tabanlı duygu analizi deneylerinde, *SentiAzNet* bilgisinin çok dilli transformer mimarisine iki hafif mekanizma ile eklenebildiği gösterilmiştir: özellik birleştirme ve dikkat ağırlıklarına sözlük tabanlı önyargı ekleme (Aykaç vd. 2026b). Ayrıca finans alanına özgü sınırlı kutup düzeltmeleri içeren bir sözlük uyarlaması geliştirilmiş ve model davranışı SHAP, dikkat örüntüleri ve sözlük eşleşmeleri üzerinden açıklanmıştır.

K3 Azerbaycan dili için büyük ölçekli, elle etiketlenmiş ve çok alanlı bir duygu analizi veri kümesi oluşturulmuştur. Teknoloji ve dijital hizmetler, finans ve iş dünyası, sosyal yaşam ve eğlence, perakende ve yaşam tarzı ile kamu hizmetleri olmak üzere beş farklı alandan toplanan 124.251 kullanıcı yorumu yerli konuşurlar tarafından etiketlenmiş ve böylece üç sınıflı duygu analizi için güçlü bir değerlendirme zemini hazırlanmıştır (Aykaç vd. 2026). Bu veri kümesi, Azerbaycan dili için hem nötr sınıf davranışının hem de alan kaymasının sistematik biçimde incelenmesine olanak tanımaktadır.

K4 Morfoloji-duyarlı, nöro-sembolik ve çok kaynaklı bir hibrit duygu analizi mimarisi geliştirilmiştir. Önerilen *MahSA* mimarisi; *SentiAzNet* tabanlı kutupluluk sinyallerini, bağlama kutupluluk yayan wPPMI bileşenini, karakter n-gram temsillerini, morfoloji-duyarlı

yüzey ipuçlarını ve transformer tabanlı bağlamsal gömmeleri tek bir füzyon yapısında birleştirmektedir (Aykaç vd. 2026). Bu yaklaşım, özellikle eklemeli yapılardan kaynaklanan kutup değişimlerini, alt-sözcük parçalanmasını ve nötr sınıf belirsizliğini hedeflemektedir.

K5 Alanlar arası sağlamlık ve temsil öğrenmesi boyutu tez kapsamına dahil edilmiştir.

Tez kapsamında tamamlanan ek bir çalışmada, kutupluluk bilgisi ile zenginleştirilmiş cümle gömme temsilleri ve alan kaymasına dayanıklı temsil öğrenme yaklaşımları ele alınmıştır (Aykaç vd. 2026a). Böylece katkı yalnızca son sınıflandırma kararının iyileştirilmesiyle sınırlı kalmamış; farklı alanlar arasında daha tutarlı, duyguya duyarlı ve genellenabilir temsillerin nasıl öğrenilebileceği de araştırılmıştır.

K6 Açıklanabilirlik, bu tezin kurucu bileşenlerinden biri olarak ele alınmıştır.

Çalışma boyunca SHAP tabanlı özellik atıfları, dikkat örüntüleri, sözlük eşleşmeleri ve hata analizi birlikte değerlendirilmiş; özellikle ironi, nötr-olumsuz belirsizliği, karma duygu, örtük değerlendirme ve kod değiştirme gibi zor örüntüler ayrı bir analitik katman olarak incelenmiştir (Aykaç vd. 2026b; Aykaç vd. 2026).

Bu bilimsel katkılara ek olarak, tez kapsamında geliştirilen yöntem ailesinin araştırma ve uygulama ortamlarında kullanılabilmesini kolaylaştıran bir gerçekleştirim bağlamı da oluşturulmuştur. Ancak bu unsur, ana bilimsel katkılar listesinde ayrı bir madde olarak değil, tezde geliştirilen kaynak ve yöntemlerin uygulamaya aktarılmasına yönelik tamamlayıcı çıktı olarak değerlendirilmiştir.

1.7 Bu Tezin Literatürdeki Konumu

Bu tezin özgünlüğü, duygu analizini yalnızca bir sınıflandırma problemi olarak değil, uçtan uca bir araştırma programı olarak ele almasından kaynaklanmaktadır. Literatürde özellikle dört boşluk belirginleşmektedir: dile özgü temel kaynak eksikliği, morfoloji-duyarlı modelleme azlığı, alanlar arası değerlendirme yetersizliği ve açıklanabilirlik/hata analizi eksikliği (Aliyu vd. 2024; Alizada vd. 2024; Hua vd. 2024; Diwali vd. 2024). Bu boşluklar, Azerbaycan dili gibi düşük kaynaklı ve morfolojik açıdan zengin dillerde daha görünür hale gelmektedir.

Bu tez; hangi dil kaynağının eksik olduğunu, bu kaynağın nasıl üretileceğini, üretilen kaynağın sinirsel modellere nasıl entegre edileceğini, morfolojik ipuçlarının hangi koşullarda kritik hale geldiğini, alan kaymasının model performansını nasıl etkilediğini ve bu tasarım kararlarının nasıl

açıklanabileceğini tek bir çerçeve altında yanıtlamaktadır. Bu yönüyle tez, Azerbaycan dili duygu analizi literatürüne artımsal bir iyileştirme sunmanın ötesine geçerek kaynak-merkezli ve model-merkezli katkıları bütünleştiren sistematik bir araştırma yaklaşımı ortaya koymaktadır (Aykaç vd. 2025, 2026b; Aykaç vd. 2026; Aykaç vd. 2026a).

1.8 Tezin Organizasyonu

Tez altı ana bölümden oluşmaktadır. **Birinci bölüm** olan bu giriş bölümünde, Azerbaycan dili bağlamında duygu analizi probleminin temel boyutları tanımlanmış; tezin amacı, temel savı, araştırma soruları ve özgün katkıları ortaya konmuştur. **İkinci bölümde**, tezin kuramsal arka planını oluşturan ilgili çalışmalar değerlendirilmektedir. Bu değerlendirme; duygu analizi kavramı ve düzeyleri, kutupluluk sözlükleri, düşük kaynaklı ve morfolojik açıdan zengin dillerdeki güçlükler, bağlamsal temsiller ve hibrit yaklaşımlar, aspekt-tabanlı ve finansal duygu analizi, alan uyarlaması ile açıklanabilir yapay zekâ eksenlerini kapsamakta ve Azerbaycan dili özelindeki literatür boşluğunu görünür kılmaktadır. **Üçüncü bölüm**, tezin yöntemsel omurgasını oluşturmaktadır. Bu bölümde Azerbaycan dili için duygu analizi çerçevesi geliştirme metodolojisi, yedi aşamalı bir iş akışı altında sunulmakta; veri kaynakları, *SentiAzNet* kutupluluk sözlüğünün geliştirilme süreci, aspekt-tabanlı sözlük enjeksiyonu, *MahSA* hibrit mimarisi, alanlar arası sağlam temsil öğrenmesi ve açıklanabilirlik protokolleri ayrıntılı biçimde açıklanmaktadır. **Dördüncü bölümde**, üçüncü bölümde tanımlanan her bir yöntem katmanına ilişkin deneysel bulgular artan karmaşıklık düzeni içinde sunulmakta; karşılaştırmalı değerlendirmeler, ablasyon deneyleri ve hata çözümlemeleri birlikte tartışılmaktadır. **Beşinci bölümde**, tezin genel değerlendirmesi yapılmakta; araştırma sorularına verilen yanıtlar, bilimsel ve uygulamalı katkıların sentezi, çalışmanın sınırlılıkları ve gelecekte yapılabilecek çalışmalar ele alınmaktadır. **Altıncı ve son bölümde** ise tezin temel çıktıları ve yöntemsel yol haritası kısaca özetlenerek çalışma sonuca bağlanmaktadır.

Bu bölümde tez problemi, amaç, araştırma soruları ve özgün katkılar çerçevesi kurulmuştur. Bir sonraki bölümde ilgili çalışmalar ve literatürdeki boşluk ayrıntılı olarak tartışılarak, tezde önerilen metodolojik çerçevenin kuramsal dayanağı oluşturulacaktır.

2. İLGİLİ ÇALIŞMALAR VE LİTERATÜR DEĞERLENDİRMESİ

Bu bölümde, tez probleminin kuramsal arka planını oluşturan ilgili çalışmalar ve literatür çizgisi ele alınmaktadır. Öncelikle duygu analizi kavramı ve çözünürlük düzeyleri açıklanmakta, ardından düşük kaynaklı ve morfolojik açıdan zengin diller bağlamındaki temel güçlükler tartışılmaktadır. Daha sonra kutupluluk sözlükleri, bağlamsal temsiller, hibrit yaklaşımlar, aspekt-tabanlı duygu analizi, alan uyarlaması ve açıklanabilir yapay zekâ eksenindeki çalışmalar değerlendirilmektedir. Bölüm, Azerbaycan dili özelindeki çalışmaların ve bu tezin literatürdeki konumunun ele alınmasıyla tamamlanmaktadır.

Bu literatür değerlendirmesi, ilgili çalışmaları kronolojik biçimde sıralamaktan çok, alanın tarihsel dönüşümünü ve bugün ulaştığı sınırları görünür kılan bir anlatı olarak kurgulanmıştır. Bu nedenle bölüm, klasik sözlük tabanlı yaklaşımlardan makine öğrenmesine, derin öğrenmeden çok dilli ön eğitilmiş modellere ve oradan hibrit ile alan-uyarlamalı yapılara uzanan çizgiyi izlemekte; her aşamada hangi problemlerin çözüldüğü ve hangi boşlukların sürdüğü tartışılmaktadır (Pang vd. 2002; Maas vd. 2011; Zhang vd. 2018; Ligthart vd. 2021; Lu vd. 2023; Ali vd. 2025).

2.1 Duygu Analizi Kavramı ve Düzeyleri

Duygu analizi, metinlerde ifade edilen öznel yönelimlerin, değerlendirme tutumlarının ve duygusal eğilimlerin otomatik olarak belirlenmesini amaçlayan doğal dil işleme problemlerinden biridir. Literatürde bu alan; görüş madenciliği, öznel ifade çözümlemesi ve kutupluluk sınıflandırması gibi kavramlarla yakın ilişki içinde ele alınmaktadır. Özellikle kullanıcı yorumları, haber yorumları, sosyal medya gönderileri ve hizmet geri bildirimleri gibi veri kaynaklarında duygu analizi; kullanıcı memnuniyetinin izlenmesi, toplumsal eğilimlerin belirlenmesi, itibar yönetimi ve karar destek süreçlerinin desteklenmesi açısından önem taşımaktadır (Taboada vd. 2011; Liu 2012; Birjali vd. 2021; Wankhade vd. 2022; Ali vd. 2025).

Duygu analizinin amacı, bir metnin yalnızca olumlu ya da olumsuz yönünü belirlemekle sınırlı

değildir. Aynı zamanda bu yönelimin hangi dilsel işaretler üzerinden kurulduğunu, ne ölçüde açık ya da örtük olduğunu ve hangi bağlamsal koşullar altında değiştiğini incelemek de bu alanın kapsamına girmektedir. Bu nedenle duygu analizi, yalnızca sınıflandırma başarımı üzerinden değil; temsil edilen dilsel bilgi düzeyi, bağlam duyarlılığı ve açıklanabilirlik kapasitesi üzerinden de değerlendirilmektedir. Özellikle düşük kaynaklı dillerde, veri eksikliğinin sembolik ve bağlamsal bilgi kaynaklarının birlikte kullanılmasını daha önemli hale getirdiği görülmektedir (Aliyu vd. 2024).

Alan tarihsel olarak incelendiğinde, belge düzeyi duygu sınıflandırmasının ilk makine öğrenmesi uygulamalarından başlayarak daha zengin temsil ailelerine doğru ilerleyen belirgin bir paradigma değişimi görülmektedir. Erken dönem çalışmalar, kelime torbası ve lineer sınıflandırıcılar üzerinden duygu etiketlemenin mümkün olduğunu göstermiştir (Pang vd. 2002). Bunu izleyen dönemde duyguya duyarlı sözcük vektörleri, derin sinir ağları ve uzak denetimli temsil öğrenmesi öne çıkmış; böylece özellik mühendisliğinin yerini giderek temsil öğrenmesi almaya başlamıştır (Maas vd. 2011; Felbo vd. 2017; Xu vd. 2019; Zhang vd. 2018; Abdullah ve Ahmet 2022). Son aşamada ise ön eğitilmiş transformer modelleri ve geniş ölçekli benchmark'lar, duygu analizini daha yüksek doğruluk düzeylerine taşımış; fakat özellikle düşük kaynaklı dillerde veri ve dil kaynaklarının niteliği belirleyici olmaya devam etmiştir (Barbieri vd. 2020; Devlin vd. 2019; Liu vd. 2019; Conneau vd. 2020; Zhang vd. 2024).

Literatürde duygu analizi çoğunlukla üç farklı çözünürlük düzeyinde incelenmektedir: belge düzeyi, cümle düzeyi ve aspekt düzeyi. Belge düzeyinde tüm metne tek bir duygu etiketi atanırken, cümle düzeyinde her cümle ayrı bir kutupluluk birimi olarak ele alınmaktadır. Aspekt düzeyi ise aynı cümle içinde farklı hedeflere yönelik farklı duyguların bulunabileceğini kabul ederek daha ince taneli bir çözümleme sunmaktadır. Özellikle ürün yorumları, müşteri geri bildirimleri ve finans ya da hizmet metinlerinde aynı cümlede hem olumlu hem de olumsuz yönelimlerin birlikte bulunabilmesi, aspekt düzeyi yaklaşımları uygulama açısından daha anlamlı hale getirmektedir (Pontiki vd. 2016; D'aniello vd. 2022; Hua vd. 2024; Šmíd ve Král 2025).

Duygu analizinde en yaygın etiketleme düzenlerinden biri negatif, nötr ve pozitif olmak üzere üç sınıflı sınıflandırmadır. Bununla birlikte nötr sınıf, yalnızca duygusuz ya da bilgi verici içerikleri temsil etmemektedir. Zayıf öznel ifadeler, dengelenmiş olumlu–olumsuz görüşler, belirsiz değer-

lendirmeler, örtük tutum bildirimleri ve kimi zaman gelecek zamana dönük beklentiler de nötr sınıf altında yer alabilmektedir. Bu durum, üç sınıflı duygu analizinde asıl güçlüğün çoğu zaman negatif ile pozitif sınıfların ayırımından çok, nötr sınıfın güvenilir biçimde modellenmesinde ortaya çıktığını göstermektedir (Aliyu vd. 2024; Çakıcı vd. 2025; Aykaç vd. 2026).

Duygu analizinin çözünürlük düzeyleri arasındaki temel farklar Çizelge 2.1’de özetlenmiştir. Bu karşılaştırma, tez kapsamında neden hem genel yorum düzeyinde hem de belirli alanlarda aspekt düzeyinde değerlendirme yapılmasının gerekli görüldüğünü de ortaya koymaktadır.

Çizelge 2.1 Duygu analizi düzeylerinin temel özellikleri

Düzy	Temel birim	Başlıca sınırlılık / avantaj
Belge düzeyi	Tüm metin	Uygulaması görece kolaydır; ancak aynı metin içindeki farklı hedeflere yönelik karşıt duyguları ayırt etmekte zorlanır.
Cümle düzeyi	Tek cümle	Belge düzeyine göre daha ince tanelidir; buna karşılık bir cümle içinde birden fazla hedef varsa yönelimleri karıştırabilir.
Aspekt düzeyi	Hedef/aspekt	En açıklayıcı düzeydir; ancak veri etiketleme, hedef tanımı ve modelleme maliyeti daha yüksektir.

2.2 Düşük Kaynaklı ve Morfolojik Açıdan Zengin Dillerde Duygu Analizi

Duygu analizi çalışmalarının önemli bir bölümü İngilizce gibi yüksek kaynaklı dillere odaklanmış olsa da, düşük kaynaklı ve morfolojik açıdan zengin dillerde aynı başarı düzeyine ulaşmak daha güç görünmektedir. Bu güçlük yalnızca etiketli veri eksikliğinden kaynaklanmamaktadır. Yeterli büyüklükte derlemelerin, dil-özel sözlüklerin, alan etiketli veri kümelerinin, morfolojik çözümleyicilerin ve güçlü dil modellerinin sınırlı olması da performansı doğrudan etkilemektedir. Düşük kaynaklı ortamlar üzerine yapılan güncel derlemeler, veri kıtlığı, alan uyumsuzluğu, çapraz dilli aktarım hataları ve değerlendirme protokollerindeki parçalanmışlığın bu tür dillerde temel sorunlar arasında yer almaya devam ettiğini göstermektedir (Aliyu vd. 2024; Lamin ve Aziz 2025).

Bu tablo, NLP ekosistemindeki dilsel eşitsizlik sorunu ile de bağlantılıdır. Araştırma ve kaynak

yatırımlarının büyük bölümünün birkaç baskın dil üzerinde yoğunlaşması, düşük kaynaklı dillerin çoğu zaman çapraz dilli aktarım, kardeş dil yardımı ve çok dilli ortak temsil uzayları üzerinden ilerlemesine neden olmaktadır (Joshi vd. 2020; Haddow vd. 2022; Mabokela vd. 2022). Bu doğrultuda geliştirilen yöntemler; makine çevirisi tabanlı aktarım, çok dilli ortak derlemeler, kardeş dil transferi ve tipolojik yakınlıktan yararlanan ince ayar düzenekleri gibi farklı stratejiler içermektedir (Johnson vd. 2017; Pires vd. 2019; Demirtas ve Pechenizkiy 2013; Keung vd. 2020; Senel vd. 2024; Veitsman ve Hartmann 2025). Ural ve Türk dilli topluluklara yönelik daha yeni çalışmalar da, düşük kaynaklı dillerde başarılı duygu analizi için yalnızca model büyüklüğünün değil, veri niteliğinin ve dil-özel sezgilerin de önemli olduğunu göstermektedir (Alnajjar vd. 2023; Kuriyozov vd. 2022; Yeshpanov ve Varol 2024; Khan vd. 2022; Çöltekin 2020; Köksal ve Özgür 2021).

Morfolojik açıdan zengin dillerde bir sözcüğün duygu yönelimi çoğu zaman yalnızca kök biçim tarafından belirlenmez. Olumsuzluk, yoksunluk, derecelendirme, kiplik ve türetim süreçleri, aynı kökten türeyen yüzey biçimlerin duygu değerini anlamlı ölçüde değiştirebilmektedir. Eklmeli dillerde bu durum daha görünür hale gelmektedir; çünkü bir kök üzerine art arda gelen ekler, hem anlamsal hem de duygusal yönelimi dönüştürebilmektedir. Türkçe üzerine yapılan çalışmalar, zengin çekim ve türetim yapısının özellikle olumsuzluk ve bağlam kaynaklı kutup kaymalarını öne çıkardığını göstermektedir (Dehkharghani vd. 2017; Zeybek vd. 2021; Altinel Girgin vd. 2024). Azerbaycan dili de aynı dil ailesi ve benzer biçimbilimsel özellikler içinde değerlendirildiğinde, duygu analizinde morfoloji kaynaklı belirsizliklerin merkezi bir rol oynadığı söylenebilir.

Çok dilli transformer modelleri düşük kaynaklı dillere güçlü bir başlangıç noktası sunmakla birlikte, alt-sözcük tabanlı tokenizasyon stratejileri özellikle eklmeli dillerde aşırı parçalanmaya yol açabilmektedir. Bu durum kök ve eklerin kutupluluğa etkisini zayıflatabilmekte, dolayısıyla duygu sınıflandırmasında hata oranını artırabilmektedir. Tokenizer kalitesinin tek başına bile çok dilli modellerin başarımını belirgin biçimde etkileyebildiği gösterilmiştir (Rust vd. 2021). Benzer biçimde, gürültülü kullanıcı üretilmiş metinlerde yazım varyasyonları, kısaltmalar, uzatmalar ve kod değiştirimi gibi olgular da standart tokenizasyon ve sözcük temelli temsil biçimlerinin sınırlarını görünür kılmaktadır (Baldwin vd. 2015).

Bu nedenle düşük kaynaklı ve morfolojik açıdan zengin dillerde duygu analizi için yalnızca daha büyük model kullanmak yeterli bir strateji olarak görülmemektedir. Karakter n-gramları, yüzey

tabanlı morfem ipuçları, kutupluluk sözlükleri ve alan-özel sembolik sinyallerin bağlamsal modellerle birlikte kullanılması daha dengeli bir araştırma yönü sunmaktadır. Nitekim son yıllarda çok dilli gömmeler, karakter temsilleri, sözlük zenginleştirme ve hibrit tasarımları bir araya getiren çalışmaların artması da bu ihtiyacı yansıtmaktadır (Barnes vd. 2018; Joulin vd. 2017a; Mutinda vd. 2023; Aykaç vd. 2026).

Azerbaycan dili özelinde bu sorunlar daha da görünür hale gelmektedir. Kamuya açık büyük ölçekli veri kümelerinin, dile özgü kutupluluk sözlüklerinin ve geniş kapsamlı değerlendirme düzeneklerinin sınırlı olması, genel amaçlı çok dilli modellerin tek başına yeterli olmayabileceğini göstermektedir. Azerbaycan dili için bağlamsal gömme temsilleri üzerine yeni çalışmalar yayımlanmış olsa da, çok alanlı ve üç sınıflı duygu analizi bağlamında kaynak, yöntem ve değerlendirme boşluklarının sürdüğü görülmektedir (Alizada vd. 2024; Isbarov vd. 2024; Aykaç vd. 2025; Aykaç vd. 2026). Bu nedenle tez kapsamında benimsenen yaklaşım, düşük kaynaklı ve morfolojik açıdan zengin dil koşullarını tasarımın merkezine yerleştirmektedir.

2.3 Kutupluluk Sözlükleri ve Sözlük Tabanlı Yaklaşımlar

Kutupluluk sözlükleri, sözcükleri ya da çok sözcüklü ifadeleri belirli duygu değerleriyle eşleştiren yapılandırılmış dil kaynaklarıdır. Bu kaynaklar çoğu zaman olumlu, olumsuz ve nötr gibi ayrık etiketler; kimi zaman da bu etiketlerin şiddetini gösteren sürekli puanlar içerir. Duygu analizi literatüründe kutupluluk sözlükleri, özellikle etiketli verinin sınırlı olduğu koşullarda, hem başlangıç düzeyinde güçlü bir taban çizgisi hem de daha karmaşık modeller için insan tarafından yorumlanabilir bir ön bilgi kaynağı olarak kullanılmıştır (Baccianella vd. 2010; Mohammad ve Turney 2013; Hutto ve Gilbert 2014; Taboada vd. 2011; Shaik vd. 2023).

Sözlük oluşturma yaklaşımları genel olarak dört grupta toplanabilir. İlk grupta, uzman etiketlemesi ile hazırlanan ve yüksek kesinlik hedefleyen elle derlenmiş sözlükler yer alır. İkinci grupta, birlikte görülme istatistikleri, PMI türü ölçüler, dağıtımsal benzerlik veya gömme uzayları üzerinden otomatik ya da yarı otomatik biçimde genişletilen sözlükler bulunmaktadır. Üçüncü grupta, yüksek kaynaklı dillerden düşük kaynaklı dillere çeviri, geri çeviri, synset hizalama ve çapraz dilli gömme temelli aktarım yöntemleri öne çıkmaktadır. Dördüncü grupta ise sözlüklerin bağımsız son karar mekanizması olmaktan çıkarılıp sinirsel modellerin içine gömülü yardımcı sinyaller olarak

kullanıldığı sözlük-zenginleştirilmiş hibrit tasarımlar görülmektedir (Taboada vd. 2011; Hamilton vd. 2016; Barnes vd. 2018; Mutinda vd. 2023; Wu vd. 2025b; An vd. 2025).

Bu sınıflandırma, duygu sözlüklerinin durağan kaynaklar değil, farklı araştırma sorularına göre yeniden tasarlanabilen araçlar olduğunu göstermektedir. Örneğin bazı çalışmalar genel amaçlı sözlükleri alan verisi üzerinde yeniden tartarak teknik kullanımlara uyarlamış; bazıları ise gömme uzayı, PMI ve grafik yayılımı ile yeni kelimeleri sözlüğe dahil etmiştir (Hamilton vd. 2016; Yang vd. 2020; Zhou vd. 2020; Wu vd. 2025b). Çapraz dilli tarafta ise çeviri ve çok dilli gömme tabanlı aktarım, düşük kaynaklı diller için uygulanabilir bir başlangıç sağlamakla birlikte anlam kayması ve kültürel farklılıklar nedeniyle insan denetimini tamamen ikame edememektedir (Barnes vd. 2018; Demirtas ve Pechenizkiy 2013; Alnajjar vd. 2023).

Klasik sözlük tabanlı yöntemlerin önemli üstünlüklerinden biri, görece düşük veri ihtiyacı ile çalışabilmeleri ve karar sürecinin hangi sözcüklerden beslendiğini daha açık biçimde gösterebilmeleridir. Bu sayede, özellikle düşük kaynaklı dillerde, veriyle beslenen büyük modellerin henüz yeterince güçlü olmadığı araştırma evrelerinde önemli bir boşluğu doldurabilirler. Buna karşılık çok anlamlılık, alan bağımlılığı, ironi, örtük değerlendirme, kelime dışı bağlam etkileri ve sözdizimsel dönüşümler gibi olgular, yalnızca sözlük eşleştirmesine dayalı karar mekanizmalarının başarımını sınırlandırmaktadır. Başka bir deyişle, sözlükler çoğu zaman yüksek kesinlikli fakat sınırlı kapsamlı bilgi üretir; bu nedenle modern yaklaşımlarda giderek daha çok “tek başına sınıflandırıcı” yerine “yardımcı sembolik bilgi kaynağı” olarak değerlendirilmektedir (Taboada vd. 2011; Shaik vd. 2023; Mutinda vd. 2023).

Bu sınırlılıkların önemli bir bölümü alan değişimi ile daha görünür hale gelmektedir. Genel dilde olumsuz çağrışım taşıyan bir sözcük, finansal bir bağlamda teknik ve çoğu zaman nötr bir terim olarak işlev görebilir; aynı şekilde görünüşte nötr bir terim, belirli bir uzmanlık alanında olumsuz risk sinyali haline gelebilir. Bu nedenle alan-özel sözlükler veya genel sözlükler üzerinde yapılan alan-özel kutup düzeltmeleri, özellikle finansal duygu analizi ve aspekt-tabanlı duygu analizi gibi ince taneli görevlerde önem kazanmaktadır (Loughran ve McDonald 2011; Araci 2019; Wu vd. 2025b).

Morfolojik açıdan zengin dillerde sözlük tasarımının bir başka kritik boyutu, kutupluluğun yalnızca

kök biçime bağlanmamasıdır. Olumsuzluk, yoksunluk, yoğunlaştırma, küçültme ve türetim gibi biçimbilimsel işaretler, aynı kökten türeyen yüzey biçimlerin duygu yönelimini anlamlı biçimde değiştirebilmektedir. Türkçe üzerine yapılan çalışmalar, sözlük tabanlı yöntemlerin morfolojik çözümleme, karakter n-gramları veya ek-duyarlı sezgilerle desteklenmediğinde önemli kutup kaymalarını kaçırabildiğini göstermektedir (Dehkharghani vd. 2017; Erşahin vd. 2019; Altinel Girgin vd. 2024). Bu gözlem, Azerbaycan dili gibi eklemeli yapının güçlü olduğu diller için daha da önemlidir.

Sözlük tabanlı yaklaşımların literatürdeki başlıca biçimleri ve bunların Azerbaycan dili açısından anlamı Çizelge 2.2’de özetlenmiştir. Bu karşılaştırma, tez kapsamında neden yalnızca sözlük oluşturmanın değil, aynı zamanda bu sözlüğün çağdaş modellere nasıl entegre edildiğinin de ayrı bir araştırma problemi olarak ele alındığını göstermektedir.

Çizelge 2.2 Kutupluluk sözlüğü temelli başlıca yaklaşım aileleri

Yaklaşım	Üretim mantığı	Başlıca güçlü yön	Başlıca sınırlılık
Uzman derlemeli sözlük	Elle etiketleme, sözlükçülük, uzman doğrulaması	Yüksek kesinlik, yüksek yorumlanabilirlik	Ölçeklenmesi zordur; alan ve zaman değişimine yavaş uyum sağlar.
Derlem / dağıtımsal genişletme	PMI, birlikte görülme, gömme benzerliği, grafik yayılımı	Kapsamı hızlı artırır, yeni terimleri yakalayabilir	Gürültüye ve alan bağımlılığına açıktır; insan doğrulaması gerektirebilir.
Çapraz dilli aktarım	Çeviri, geri çeviri, synset hizalama, çok dilli gömmeler	Düşük kaynaklı diller için uygulanabilir bir başlangıç sağlar	Anlam kayması, yanlış eşleşme ve kültürel nüans kaybı riski taşır.
Sözlük-zenginleştirilmiş sinirsel model	Sözlük sinyalinin gömme, dikkat ya da istem düzeyinde modele eklenmesi	Bağlam ve sembolik bilgiyi birlikte kullanılır	Sözlük kapsamı düşükse katkı seyrek kalabilir; entegrasyon tasarımı kritiktir.

Azerbaycan dili özelinde bu literatür boşluğu, *SentiAzNet* çalışmasıyla doğrudan hedeflenmiştir. Söz konusu çalışma, İngilizce merkezli ve ilgili dillerden gelen çoklu sözlük kaynaklarını birleştirerek 9.614 benzersiz İngilizce lemma üzerinden başlayan bir aday havuzu oluşturmuş; çeviri, geri çeviri, anlamsal doğrulama ve insan etiketlemesi adımları sonucunda 2.482 Azerbaycan dili girişinden oluşan bir kutupluluk sözlüğü geliştirmiştir. Makalede etiketleyiciler arası uyum katsayısının Fleiss $\kappa = 0.87$ olduğu, sözlük oluşturma hattının ise iç doğrulamada %91.5 doğruluk ve 0.892 Macro-F1 ürettiği raporlanmaktadır (Aykaç vd. 2025). Bu sonuç, Azerbaycan dili için

sözlük düzeyinde güvenilir bir temel sembolik kaynağın üretilebildiğini göstermesi bakımından önemli bir adım olarak değerlendirilebilir.

Bununla birlikte *SentiAzNet* bu tez açısından tek başına nihai hedef değildir. Asıl önemli nokta, bu kaynağın daha sonra geliştirilen modellerde nasıl yeniden kullanıldığıdır. Finans ve iş dünyası alanındaki ABSA çalışmasında sözlük bilgisi, XLM-R tabanlı modele ya özellik birleştirme yoluyla ya da dikkat skorlarına eklenen bir önyargı terimi aracılığıyla aktarılmıştır. *MahSA* çalışmasında ise sözlük sinyalleri wPPMI, karakter n-gramları, morfoloji-duyarlı yüzey ipuçları ve bağlamsal gömmelerle birlikte nöro-sembolik bir füzyon yapısına dahil edilmiştir (Aykaç vd. 2026b; Aykaç vd. 2026). Dolayısıyla bu tez açısından kutupluluk sözlüğü, klasik anlamda bağımsız bir kaynak olmanın ötesinde, daha geniş bir hibrit duygu analizi mimarisinin kurucu bileşenlerinden biri olarak ele alınmaktadır.

Bu nedenle tezde kutupluluk sözlükleri literatürü iki düzeyde değerlendirilmektedir: ilk düzeyde Azerbaycan dili için böyle bir kaynağın neden gerekli olduğu ve nasıl üretildiği; ikinci düzeyde ise üretilen bu kaynağın çağdaş bağlamsal modellere ne şekilde enjekte edilerek daha yüksek performans ve daha güçlü açıklanabilirlik sağladığı tartışılmaktadır. Bu yaklaşım, sözlük tabanlı yöntemlerle sinirsel modelleme arasında keskin bir karşıtlık kurmak yerine, iki hattı birbirini tamamlayan araştırma yönleri olarak konumlandırmaktadır.

2.4 Çok Dilli Temsiller, Bağlamsal Modeller ve Hibrit Yaklaşımlar

Dağıtımsal temsil öğrenmesi ve bağlamsal dil modellerindeki gelişmeler, duygu analizi literatürünün yönünü önemli ölçüde değiştirmiştir. Word2Vec, GloVe ve fastText gibi statik gömme yöntemleri, sözcükleri büyük derlemeler üzerinden vektör uzayında temsil ederek benzer bağlamlarda görülen sözcüklerin anlamsal yakınlığını modellemiştir. Daha sonra BERT, XLM-RoBERTa ve benzeri transformer tabanlı bağlamsal modeller, bir sözcüğün anlamını sabit bir vektör olarak değil, içinde bulunduğu cümle bağlamına göre dinamik biçimde üretmeye başlamış ve duygu analizi dahil birçok görevde belirgin başarı artışları sağlamıştır (Devlin vd. 2019; Conneau vd. 2020). Cümle düzeyinde ise Sentence-BERT ve LaBSE gibi yaklaşımlar, özellikle benzerlik hesapları, temsil karşılaştırmaları ve gömme tabanlı aktarım senaryoları için güçlü bir altyapı sunmuştur (Reimers ve Gurevych 2019; Feng vd. 2022).

Bu dönüşüm yalnızca model ailesinin değişmesi anlamına gelmemektedir; aynı zamanda duygunun hangi temsil düzeyinde kodlandığına ilişkin bakışı da dönüştürmüştür. Statik sözcük gömmeleri ve alt-sözcük temelli hızlı modeller, özellikle veri-verimli ortamlarda ve yerel derlemler üzerinde hâlâ yararlı olmaya devam etmektedir (Bojanowski vd. 2017; Joulin vd. 2016, 2017b; Aydın vd. 2020). Buna karşılık RoBERTa, BERTurk, TurkishBERTweet ve benzeri kodlayıcılar, sosyal medya veya görev-özel ince ayar senaryolarında daha zengin bağlam duyarlılığı sağlamıştır (Liu vd. 2019; Schweter 2020; Najafi ve Varol 2024; Yildirim 2024). Bu durum, düşük kaynaklı diller için tek bir temsil ailesinin değil, görev ve veri koşullarına göre değişen çok katmanlı bir temsil ekosisteminin gerektiğine işaret etmektedir.

Düşük kaynaklı diller açısından çok dilli modellerin en önemli avantajı, yüksek kaynaklı dillerden dolaylı aktarım yapabilmeleridir. Aynı modelin ortak bir alt-sözcük sözlüğü ve ortak parametre uzayı içinde birçok dili birlikte işlemesi, etiketli verisi sınırlı dillere belirli bir başlangıç performansı kazandırmaktadır. Bununla birlikte bu aktarım her zaman sorunsuz değildir. Söz varlığı kapsamının düşük olması, alt-sözcük parçalanmasının aşırıya kaçması ve hedef dilin alan-özel örüntülerinin ön eğitim sırasında yeterince temsil edilmemesi, özellikle eklemeli dillerde önemli performans darboğazları yaratabilmektedir (Conneau vd. 2020; Rust vd. 2021; Keung vd. 2020).

Tokenizasyon kalitesi bu noktada merkezi bir değişken haline gelmektedir. Eklemeli dillerde bir kök üzerinde biriken çok sayıda ek, çok dilli tokenizer'lar tarafından fazla parçalanabildiği için, kutupluluğu dönüştüren biçimbirimler ile anlamsal çekirdeğin ilişkisi zayıflayabilmektedir. Gürültülü kullanıcı metinlerinde görülen yazım hataları, kasıtlı uzatmalar, karışık alfabe kullanımı ve kod değiştirimi de bu sorunu büyütmektedir. Noisy text üzerine yapılan çalışmalar, klasik NLP boru hatlarının ve temiz derlemlere göre eğitilmiş tokenizasyon şemalarının sosyal medya benzeri ortamlarda belirgin kırılma gösterdiğini ortaya koymaktadır (Baldwin vd. 2015; Rust vd. 2021).

Bu nedenle çok dilli bağlamsal model literatürü zaman içinde iki tamamlayıcı yöne ayrılmıştır. İlk yön, alan uyarlamalı ön eğitim, görev uyarlamalı ön eğitim, LoRA/QLoRA gibi parametre-verimli ince ayar teknikleri ve karşıtsal öğrenme ile mevcut büyük modelleri hedef veriye daha uygun hale getirmeye çalışmaktadır (Gururangan vd. 2020; Hu vd. 2022; Dettmers vd. 2023; Nwaiwu 2025). İkinci yön ise bağlamsal temsilleri karakter düzeyi bilgi, sözlük sinyali, dış bilgi tabanı

ya da morfoloji-duyarlı sezgilerle birleştiren hibrit ve nöro-sembolik modeller geliştirmektedir (Mutinda vd. 2023; An vd. 2025; Aykaç vd. 2026). Tez kapsamında izlenen yaklaşım, bu iki yönün birbirini dışlamadığını; aksine düşük kaynaklı bir dilde çoğu zaman birlikte değerlendirilmesinin daha verimli sonuçlar doğurabileceğini kabul etmektedir.

Azerbaycan dili için bağlamsal temsil araştırmaları henüz görece yenidir. Alizada ve arkadaşlarının çalışması, Azerbaycan dili için RoBERTa ve GPT-2 tarzı bağlamsal modellerle ilk sistematik girişimlerden birini sunarken, açık temel modeller üzerine yapılan daha yeni çalışmalar Azerbaycan dili için model ekosisteminin genişlemeye başladığını göstermektedir (Alizada vd. 2024; Isbarov vd. 2024). Bununla birlikte bu çalışmalar, çok alanlı ve üç sınıflı duygu analizi, nötr sınıfın güvenilir yönetimi ve alan kaymasına dayanıklı temsil öğrenmesi gibi sorunları doğrudan çözmemektedir. Bu nedenle Azerbaycan dili için yalnızca “bir bağlamsal kodlayıcı seçmek” yeterli bir araştırma stratejisi olarak görünmemektedir; kodlayıcının hangi ek bilgi kaynaklarıyla destekleneceği de aynı ölçüde önem taşımaktadır.

Temsil aileleri arasındaki temel farklar Çizelge 2.3’de özetlenmiştir. Çizelge, tez kapsamındaki modelleme tercihinin neden tek bir temsil ailesine yaslanmak yerine, farklı temsil düzeylerini birleştiren hibrit yapılar lehine şekillendiğini açıklamaktadır.

Hibrit yaklaşımın Azerbaycan dili açısından somutlaşmış iki önemli örneği tezden türeyen çalışmalarda görülmektedir. Finans ve iş dünyası alanındaki ABSA çalışmasında *SentiAzNet* sinyali, XLM-R modeline iki hafif mekanizma ile eklenmiştir: bağlamsal gösterime kutupluluk özelliği birleştirme ve dikkat skorlarına sözlük tabanlı önyargı terimi ekleme. Bu tasarım, sözlük kapsamının seyrek olmasına rağmen, sözlüğün mevcut olduğu örneklerde karar dağılımını faydalı biçimde yönlendirebildiğini göstermiştir (Aykaç vd. 2026b). *MahSA* çalışması ise bu fikri daha da genişleterek *SentiAzNet* sinyallerini wPPMI ile bağlama yaymış, karakter n-gramları ve morfoloji-duyarlı yüzey ipuçlarıyla birleştirmiş ve XLM-R tabanlı bağlamsal temsillerle füzyonlayarak çok alanlı bir nöro-sembolik mimari önermiştir (Aykaç vd. 2026).

Bu çizginin temsil öğrenmesi açısından doğal devamı, alan kaymasına dayanıklı ve duyguya duyarlı cümle temsilleri öğrenmektir. Düşük kaynaklı dillerde her alan için ayrı büyük veri kümeleri toplamak güç olduğundan, yalnızca sınıflandırma katmanını iyileştirmek yeterli değildir; farklı

Çizelge 2.3 Duygu analizinde öne çıkan başlıca temsil aileleri

Temsil ailesi	Tipik örnekler	Başlıca güçlü yön	Azerbaycan dili açısından temel sorun
Statik sözcük gömmeleri	Word2Vec, GloVe, fastText	Hesaplama maliyeti düşüktür; yerel derlemlerde kolay eğitilir	Çok anlamlılık ve bağlam değişimi zayıf temsil edilir; yüzey varyasyonları sorun yaratır.
Cümle gömme temsilleri	Sentence-BERT, LaBSE	Cümle düzeyinde anlamsal yakınlığı iyi taşır; benzerlik ve aktarım görevlerinde kullanışlıdır	Görev-özel ince ayar yapılmadığında duyguya özgü sinyaller seyrek kalabilir.
Çok dilli bağlamsal kodlayıcılar	mBERT, XLM-R	Çapraz dilli aktarım ve güçlü bağlam modelleme sağlar	Tokenizer parçalanması, alan uyumsuzluğu ve düşük kapsama özellikle eklemeli dillerde performansı sınırlandırabilir.
Hibrit / nöro-sembolik modeller	Sözlük+transformer, karakter+embedding füzyonu, MahSA	Sembolik bilgi ile bağlamsal gücü birleştirir; daha yorumlanabilir ve veri-verimli olabilir	Tasarım ve füzyon maliyeti artar; hangi bileşenin ne kadar katkı verdiği dikkatli analiz gerektirir.

alanlarda daha tutarlı davranan bir gömme uzayı öğrenmek de gereklidir. Bu nedenle güncel literatürde karşıtsal öğrenme, alan-uyarlamalı ön eğitim ve duyguya duyarlı temsil öğrenmesi gibi kavramlar daha fazla öne çıkmaktadır (Gururangan vd. 2020; Nwaiwu 2025). Tez kapsamında tamamlanan ve yayıma hazırlanan çalışma da bu ekseninde, kutupluluk enjeksiyonu ve alanlar arası sağlamlık odaklı cümle gömme temsilleri üzerine yoğunlaşmaktadır (Aykaç vd. 2026a).

Sonuç olarak, çok dilli bağlamsal modeller Azerbaycan dili için güçlü bir başlangıç noktası oluştursa da, tez kapsamında savunulan yaklaşım bu modellerin tek başına yeterli olmayabileceği yönündedir. Azerbaycan dili gibi düşük kaynaklı ve morfolojik açıdan zengin bir dilde daha güvenilir duygu analizi için, bağlamsal kodlayıcıların sözlük sinyalleri, karakter düzeyi bilgi, morfoloji-duyarlı yüzey ipuçları ve alan-özel düzenlemeler ile desteklenmesi gerekmektedir. Bu nedenle bu tez, modern temsil öğrenmesi literatürünü reddetmekten çok, onu sembolik ve dilbilgisel bilgiyle tamamlayan hibrit bir araştırma programı önermektedir.

2.5 Aspekt-Tabanlı Duygu Analizi ve Finansal Duygu Analizi

Aspekt-tabanlı duygu analizi (ABSA), bir cümle veya belge içinde geçen belirli bir hedefe, özelliğe ya da bileşene ilişkin duygunun ayrı ayrı belirlenmesini amaçlayan ince taneli bir duygu analizi yaklaşımıdır. Cümle düzeyindeki geleneksel sınıflandırma, tek bir ifadeye tek bir duygu etiketi atayarak pratik bir çözüm sunsa da, çok bileşenli değerlendirmelerin yer aldığı gerçek kullanıcı metinlerinde çoğu zaman yetersiz kalmaktadır. Bir ürünün fiyatı olumsuz, hizmet kalitesi olumlu; bir bankacılık uygulamasının arayüzü başarılı, ücret politikası ise sorunlu biçimde değerlendirilebilir. Bu nedenle ABSA, özellikle birden fazla hedefin aynı metin içinde birlikte değerlendirildiği uygulamalarda, belge ve cümle düzeyine göre daha açıklayıcı bir çerçeve sunmaktadır (Hua vd. 2024; Pontiki vd. 2016; D’aniello vd. 2022; Šmíd ve Král 2025).

Finans ve iş dünyası metinleri, aspekt-tabanlı duygu analizi açısından ayrıca özel bir güçlük kümesi içermektedir. Bu alanda birçok ifade açık biçimde olumlu ya da olumsuz görünmemekte; kimi zaman bilgi verici, düzenleyici ya da teknik içerik taşıyan cümleler, okuyucuda dolaylı bir değerlendirme etkisi oluşturmaktadır. Komisyon, faiz, limit, işlem süresi, onay süreci ve güvenlik gibi kavramlar, yalnızca sözlüksel kutuplulukla değil, ilgili aspektin bağlam içindeki işleviyle birlikte anlam kazanmaktadır. Bunun sonucu olarak, finansal ABSA problemlerinde nötr sınıf belirsizliği, örtük olumsuzluk, karma duygu ve aspektler arası çelişki daha belirgin hale gelmektedir. İngilizce literatürde finansal metinlerin genel dilden ayrı ele alınması gerektiğini gösteren Loughran–McDonald sözlüğü ve FinBERT gibi çalışmalar da bu noktaya işaret etmektedir (Loughran ve McDonald 2011; Araci 2019).

Güncel ABSA literatürü, sorunun yalnızca hedefe bağlı sınıflandırmadan ibaret olmadığını; hedef çıkarımı, çapraz dilli aktarım, alanlar arası uyarılama ve açıklanabilirlik boyutlarının da giderek daha fazla önem kazandığını göstermektedir (Hua vd. 2024; Šmíd ve Král 2025; Šmíd vd. 2025; Wu vd. 2025a; Xu ve Ibrahim 2022; Zhao vd. 2024). Finansal metinlere odaklanan güncel karşılaştırmalar ise, büyük dil modelleri ve ince ayarlı encoder ailelerinin umut verici sonuçlar ürettiğini; ancak alan sözlüğü, hedef sinyali ve hata analizi olmadan özellikle nötr ve örtük duygu vakalarında performansın kırılgan kalabildiğini göstermektedir (Mughal vd. 2024; Nath ve Dwivedi 2024; Nasiopoulos vd. 2025; Kreuzer vd. 2026; Sheetal ve Aithal 2025).

ABSA ve finansal duygu analizi bağlamında öne çıkan başlıca güçlükler Çizelge 2.4’de özetlenmiştir. Bu tablo, tez kapsamında neden genel cümle düzeyi sınıflandırmaya ek olarak finans ve iş dünyası alanında aspekt düzeyi bir incelemenin gerekli görüldüğünü de açıklamaktadır.

Çizelge 2.4 Aspekt-tabanlı ve finansal duygu analizinde öne çıkan güçlükler

Güçlük	Tipik görünüm	Duygu analizi açısından sonucu
Aynı metinde çoklu hedef	Ücret olumlu, hizmet olumsuz gibi eşzamanlı değerlendirme	Cümle düzeyinde tek etiket yetersiz kalır; hedefe bağlı modelleme gerekir.
Teknik ve alan-özel söz varlığı	Faiz, komisyon, limit, spread, gecikme, onay	Genel amaçlı sözlük ve modeller alan bağlamını kaçırabilir.
Nötr-ağırlıklı ifadeler	Bilgilendirici ya da raporlayıcı cümleler	Nötr sınıf, negatif ile karışmaya daha yatkın hale gelir.
Örtük ve karma duygu	Aynı cümlede hem memnuniyet hem şikâyet	Aspekt ayrımı yapılmadığında karar sınırları bulanıklaşır.
Kod değiştirme ve ödüncleme	İngilizce veya Türkçe finans terimlerinin karışımı	Tokenizasyon ve sözlük kapsamı sorunları artar.

Son yıllarda ABSA literatüründe dikkat mekanizmaları, hedef-koşullu transformer mimarileri ve çok dilli aktarım yaklaşımları yaygınlaşmıştır. Bununla birlikte düşük kaynaklı diller için temel sorun sürmektedir: pek çok yaklaşım yüksek kaynaklı dillerdeki veri bolluğuna ve hazır alan kaynaklarına dayanmaktadır. Düşük kaynaklı ortamlarda ise alan-özel sözlükler, harici bilgi kaynakları ve açıklanabilirlik araçları yeniden önem kazanmaktadır. Tez kapsamında finans ve iş dünyası Azerbaycan dili verisi üzerinde gerçekleştirilen çalışma da bu literatürde açık kalan boşluklardan birine karşılık gelmektedir. Bu çalışmada *SentiAzNet* kutupluluk sözlüğünden gelen sinyal, XLM-R tabanlı bir ABSA modeline hem özellik birleştirme hem de dikkat önyargısı yoluyla enjekte edilmiş; ayrıca sınırlı sayıda alan-özel düzeltme içeren *SentiAzNet-Fin* tanımlanmıştır (Aykaç vd. 2026b). Elde edilen sonuçlar, sözlük kapsamı seyrek olsa bile, doğru noktada kullanılan sembolik sinyallerin finansal ABSA performansını ve modelin açıklanabilirliğini iyileştirebildiğini göstermektedir. Bu nedenle finans ve iş dünyası ABSA, tezde yalnızca yan bir uygulama alanı değil; sözlük bilgisinin çağdaş bağlamsal modellere nasıl entegre edilebileceğini gösteren önemli bir ara katman olarak değerlendirilmektedir.

2.6 Alan Uyarlaması, Alanlar Arası Sağlamlık ve Temsil Öğrenmesi

Duygu analizi sistemleri çoğu zaman eğitim verisinin toplandığı alan ile gerçek kullanım alanı arasında bir dağılım farkı ile karşılaşmaktadır. Ürün yorumları, haber yorumları, sosyal medya içerikleri, kamu hizmeti şikâyetleri ve finansal geri bildirimler; sözcük dağarcığı, anlatım biçimi, duygu yoğunluğu ve nötr sınıf oranı bakımından birbirinden belirgin biçimde ayrışabilmektedir. Bu nedenle bir alanda iyi sonuç veren bir modelin başka bir alana doğrudan taşınması çoğu zaman performans kaybına yol açmaktadır. Literatürde bu sorun alan kayması, alan uyumsuzluğu veya alanlar arası genelleme problemi olarak ele alınmakta; özellikle düşük kaynaklı dillerde bu problem, veri kıtlığı nedeniyle daha da büyümektedir (Ganin ve Lempitsky 2015; Ramponi ve Plank 2020; Gururangan vd. 2020).

Alan uyarlaması literatürü, klasik yapılandırılmış karşıtlık öğrenmesinden güncel source-free ve meta-learning tabanlı yaklaşımlara uzanan geniş bir yöntem ailesi üretmiştir. Erken dönemde yapısal karşıtlık öğrenmesi ve alan-karşıtlı eğitim dikkat çekmiş; daha sonra alan-uyarlamalı ön eğitim, çok görevli karşıtlık, meta-genelleme ve kaynaksız uyarlama senaryoları öne çıkmıştır (Fang vd. 2014; Ganin vd. 2016; Liu vd. 2017; Wang vd. 2021; Badr vd. 2024; Fang vd. 2024). Duygu analizine özel tarafta ise sentiment-aware ön eğitim ve karşıtsal cümle gömme temsilleri, yalnızca son karar katmanını değil gömme geometrisini de düzenleyerek alanlar arası genellemeyi güçlendirmeyi amaçlamaktadır (Zhou vd. 2020; Gao vd. 2021; Khosla vd. 2020).

Alanlar arası sağlamlık için geliştirilen yöntemler kabaca birkaç grupta toplanabilir. İlk grup, eğitim örneklerini kaynak ve hedef alan arasındaki dağılım farkına göre yeniden ağırlıklandıran yaklaşımlardır. İkinci grup, alanlar arasında ortak çalışan özellik uzayları öğrenmeye odaklanan hizalama ve alana duyarsız ortak temsil yaklaşımlarını içermektedir. Üçüncü grup, alan-uyarlamalı ön eğitim (DAPT/TAPT) ile kodlayıcıyı hedef alana yakınlaştıran yöntemlerden oluşmaktadır. Dördüncü grup ise karşıtsal öğrenme, parametre-verimli ince ayar ve temsil uzayının geometrisini doğrudan düzenleyen daha yeni teknikleri kapsamaktadır. Bu stratejilerin genel görünümü Çizelge 2.5’da özetlenmiştir.

Düşük kaynaklı ve morfolojik açıdan zengin dillerde alan uyarlaması yalnızca sınıflandırıcıyı yeniden eğitime meselesi değildir. Aynı duygu sınıfına ait örneklerin farklı alanlarda benzer biçimde

Çizelge 2.5 Alanlar arası sağlamlık için başlıca yöntem aileleri

Yöntem ailesi	Temel fikir	Düşük kaynaklı diller açısından değerlendirme
Örnek ağırlıklandırma	Kaynak örneklerin etkisini hedef alana göre yeniden ayarlama	Hesaplama maliyeti düşüktür; ancak güçlü temsil farklarını tek başına kapatamaz.
Özellik hizalama / alana duyarsız temsil	Alan bilgisini bastırıp ortak bir temsil uzayı öğrenme	Alan kaymasına karşı etkilidir; ancak eğitim kararlılığı her zaman kolay değildir.
Alan-uyarlamalı ön eğitim	Kodlayıcıyı hedef alan metinleriyle yeniden ön eğitme	Büyük ek veri varsa faydalıdır; düşük kaynaklı senaryoda veri ve hesaplama maliyeti yükselebilir.
Karşıtsal öğrenme ve temsil düzenleme	Aynı sınıftaki örnekleri alanlar arasında yakınlaştırma	Temsil uzayını doğrudan iyileştirir; alanlar arası genelleme için güçlü bir seçenektir.
Parametre-verimli uyarlama	Az sayıda ek parametre ile alan uyarlaması yapma	Hesaplama verimlidir; ancak başarımlar, temsil kalitesine ve veri rejimine bağlıdır.

temsil edilmesi, başka bir deyişle duyguya duyarlı ve alan kaymasına karşı daha kararlı bir gömme uzayının öğrenilmesi gerekir. Bu nedenle temsil öğrenmesi son yıllarda duygu analizi araştırmalarının merkezine yerleşmiştir. Özellikle alan-uyarlamalı ön eğitim, karşıtsal düzenleme ve sentiment-aware cümle gömme temsilleri gibi yaklaşımlar, yalnızca son katmandaki karar sınırını değil, modelin iç temsillerini de iyileştirmeyi hedeflemektedir. Parametre-verimli uyarlama tekniklerinin yükselişi de bu bağlamda önemlidir; çünkü düşük kaynaklı ortamlarda her yeni alan için tam yeniden eğitim çoğu zaman pratik değildir (Ramponi ve Plank 2020; Gururangan vd. 2020; Nwaiwu 2025).

Tez kapsamında bu problem iki düzeyde ele alınmaktadır. İlk olarak *MahSA* çalışması, çok alanlı veri kümesi ve bir alanı dışarıda bırakma (LODO) değerlendirmeleriyle, sözlük ve morfoloji bilgisi içeren hibrit bir mimarinin alan kayması altında nasıl davrandığını göstermektedir (Aykaç vd. 2026). İkinci olarak tez kapsamında tamamlanan ve makale formuna dönüştürülen *Domain-Robust, Polarity-Injected Sentence Embeddings for Azerbaijani Sentiment Analysis* çalışması, problemi daha doğrudan temsil öğrenmesi düzeyine taşımaktadır (Aykaç vd. 2026a). Bu çalışmada kutupluluk enjeksiyonu, morfoloji-duyarlı yüzey ipuçları, alan uyarlaması ve karşıtsal düzenleme birlikte ele alınarak daha sağlam cümle gömme temsilleri öğrenilmektedir. Dolayısıyla tezde

alanlar arası sađlamlık, yalnızca ek bir deđerlendirme protokolü deđil; Azerbaycan dili için daha taşınabilir ve daha güvenilir duygu analizi sistemleri kurmanın temel koşullarından biri olarak konumlandırılmaktadır.

2.7 Açıklanabilir Yapay Zekâ ve Hata Analizi

Duygu analizinde performans artışı kadar, modelin hangi ipuçlarına dayanarak karar verdiğinin anlaşılması da önemlidir. Özellikle sosyal medya izleme, finansal karar destek ve kamuoyu analizi gibi uygulamalarda, modelin belirli bir metni neden negatif ya da pozitif olarak etiketlediğinin görünür olması güvenilirlik açısından kritik kabul edilmektedir. Bu nedenle açıklanabilir yapay zekâ yöntemleri, duygu analizinde yalnızca sonradan eklenen bir görselleştirme aracı deđil, model davranışını çözümlemenin sistematik bir bileşeni haline gelmiştir (Ribeiro vd. 2016; Lundberg ve Lee 2017).

Son yıllarda XAI literatüründe, açıklamaların sadakati, kararlılığı ve kullanıcı açısından anlamlılığı üzerine daha yoğun bir tartışma yürütölmektedir. Duygu analizi ve ABSA özelinde yayımlanan güncel derlemeler, SHAP, LIME, entegre gradyanlar, dikkat örüntüleri ve örnek-tabanlı açıklamaların birlikte deđerlendirilmesi gerektiğini vurgulamaktadır (Diwali vd. 2024; Perikos ve Diamantopoulos 2024; Mabokela vd. 2024; Chrysostomou ve Aletras 2021). Özellikle sarkazm, ironi ve örtük tutum gibi karmaşık örüntülerde açıklama kanallarının birbirini doğrulayıp doğrulamadığını görmek, salt yüksek F1 deđerinden daha öğretici olabilmektedir (Suhartono vd. 2024; Pokhriyal ve Jain 2025b,a; Yue vd. 2023; Dhall vd. 2025).

Bununla birlikte açıklanabilirlik araçlarının nasıl yorumlanacağı konusunda dikkatli olunmalıdır. Özellikle dikkat görselleştirmeleri sıklıkla açıklama amacıyla kullanılmasına rağmen, dikkat ağırlıklarının tek başına nedensel açıklama sunmadığı gösterilmiştir. Bu nedenle SHAP, dikkat örüntüleri, sözlük eşleşmeleri ve hata kategorileri gibi birden fazla kanalı birlikte deđerlendirmek daha güvenli bir yaklaşım sunmaktadır (Jain ve Wallace 2019; Zhu vd. 2024). Duygu analizinde sık görölen hata türleri arasında ironi, sarkazm, karışık duygu, örtük deđerlendirme, nötr-negatif belirsizliği ve kod deđiştirme yer almaktadır. Bu hata kalıpları özellikle düşük kaynaklı ve gürültölü kullanıcı metinlerinde daha belirgin hale gelmektedir.

2.8 Azerbaycan Dili için Duygu Analizi Çalışmaları

Azerbaycan dili için duygu analizi literatürü, Türkçe ve İngilizce ile karşılaştırıldığında hâlen sınırlı görünmektedir. Son yıllarda haber ve sosyal medya verileri üzerinde kamuoyu eğilimini inceleyen çalışmalar ortaya çıkmış; ayrıca Azerbaycan dili için bağlamsal sözcük gömmeleri ve açık temel modeller üzerine ilk sistematik girişimler yayımlanmıştır. Bununla birlikte bu çalışmaların önemli bir kısmı ya dar veri kümeleri üzerinde kurulmuş ya da çok alanlı ve üç sınıflı değerlendirme ihtiyacını sınırlı ölçüde karşılamıştır (Guliyev vd. 2024; Alizada vd. 2024; Isbarov vd. 2024).

Bölgesel bağlam dikkate alındığında, Azerbaycan dili için duygu analizi araştırmalarının Türkçe, Kazakça, Özbekçe ve diğer akraba dillerdeki gelişmelerle birlikte okunması yararlıdır. Çünkü bu dillerde ortaya çıkan veri setleri, bağlamsal model uyarlamaları ve morfoloji-duyarlı yöntemler, Azerbaycan dili için doğrudan kopyalanmasa da güçlü karşılaştırma noktaları sunmaktadır (Köksal ve Özgür 2021; Mutlu ve Özgür 2022; Najafi ve Varol 2024; Kuriyozov vd. 2022; Yeshpanov ve Varol 2024; Veitsman ve Hartmann 2025). Bu nedenle Azerbaycan dili literatürünün görece darlığı, yalnızca eksiklik olarak değil, aynı zamanda tipolojik yakınlıktan yararlanan yöntemsel aktarım için bir fırsat olarak da okunabilir.

Azerbaycan dili için kamuya açık ve duygu analizine doğrudan hizmet eden temel sembolik kaynak eksikliği, bu literatürdeki en önemli boşluklardan biri olarak görünmektedir. Bu boşluğun giderilmesine yönelik son çalışmalardan biri olan *SentiAzNet*, Azerbaycan dili için diğer dillere ait zengin kaynaklardan türetilmiş ve manuel doğrulama ile güçlendirilmiş bir kutupluluk sözlüğü sunmaktadır. Devamındaki çalışmalarda ise bu sözlüğün yalnızca bağımsız bir kaynak olarak değil, çağdaş bağlamsal modelleri zenginleştiren yardımcı bir sinyal olarak kullanılabildiği gösterilmiştir. Finans ve iş dünyası metinlerinde sözlük destekli açıklanabilir ABSA yaklaşımı ile çok alanlı *MahSA* mimarisi, bu çizginin uygulama ve sistem düzeyindeki devamı niteliğindedir (Aykaç vd. 2025, 2026b; Aykaç vd. 2026).

Bu çalışmalar birlikte değerlendirildiğinde Azerbaycan dili için duygu analizi araştırmasının üç kademeli bir gelişim hattı izlediği söylenebilir. İlk kademe temel kaynağın, yani kutupluluk sözlüğünün oluşturulmasıdır. İkinci kademe, bu kaynağın alan-özel ve açıklanabilir modeller içinde kullanılmasının gösterilmesidir. Üçüncü kademe ise çok alanlı, morfoloji-duyarlı ve alan

uyarlamasını dikkate alan hibrit mimarilerin geliştirilmesidir. Mevcut tez çalışması bu hattı daha bütüncül bir çerçeve içinde birleştirmektedir.

2.9 Literatürdeki Boşluk ve Bu Tezin Konumu

Literatür bütüncül olarak incelendiğinde, Azerbaycan dili için duygu analizi alanında dört temel boşluk öne çıkmaktadır. Birincisi, dile özgü ve güvenilir temel kaynakların, özellikle kutupluluk sözlüklerinin ve çok alanlı etiketli veri kümelerinin sınırlı olmasıdır. İkincisi, morfolojik kutup değişimlerini ve eklemeli yapının etkisini açık biçimde modelleyen yöntemlerin azlığıdır. Üçüncüsü, alan-özel ve alanlar arası değerlendirmeyi birlikte ele alan, nötr sınıfları açıkça hedefleyen sistemlerin sınırlı olmasıdır. Dördüncüsü ise model kararlarını şeffaflaştıran açıklanabilirlik analizlerinin ve hata tiplerinin sistematik biçimde incelenmemesidir (Aliyu vd. 2024; Guliyev vd. 2024; Alizada vd. 2024; Aykaç vd. 2026).

Bu tez, söz konusu boşlukları birbirinden kopuk alt problemler olarak değil, tek bir araştırma programının parçaları olarak ele almaktadır. Buna göre ilk olarak Azerbaycan dili için yeniden kullanılabilir bir kutupluluk sözlüğü geliştirilmiş; ardından bu sözlük bilgisinin çağdaş bağlamsal modellere nasıl enjekte edilebileceği alan-özel bir ABSA senaryosunda incelenmiş; daha sonra çok alanlı ve morfoloji-duyarlı hibrit bir mimari ile genel duygu analizi problemi ele alınmış; son olarak ise alanlar arası sağlam temsil öğrenmesi yönünde ek bir genişleme gerçekleştirilmiştir. Bu konumlandırma, çalışmanın yalnızca tek bir model ya da tek bir veri kümesi üretmekten ibaret olmadığını; Azerbaycan dili için kaynak, yöntem, değerlendirme ve açıklanabilirlik katmanlarını bir araya getiren bütüncül bir duygu analizi çerçevesi önerdiğini göstermektedir.

Bu konumlandırma aynı zamanda literatürde görülen iki aşırı uçtan da bilinçli olarak uzak durmaktadır: bir tarafta salt sözlük tabanlı ama bağlamı sınırlı yaklaşımlar, diğer tarafta büyük modellerin tek başına yeterli olacağını varsayan veri-merkezli yaklaşımlar bulunmaktadır. Tezde önerilen hat, bu iki çizginin güçlü yönlerini birleştirmeyi; yani yeniden kullanılabilir sembolik kaynaklar, bağlamsal temsil gücü, morfoloji-duyarlı yüzey sezgileri ve alanlar arası değerlendirme protokollerini aynı araştırma programı içinde bir araya getirmeyi amaçlamaktadır (Taboada vd. 2011; Hamilton vd. 2016; Devlin vd. 2019; Conneau vd. 2020; Mutinda vd. 2023; Aykaç vd. 2025, 2026b; Aykaç vd. 2026; Aykaç vd. 2026a).

Bu bölümde sunulan literatür değerlendirmesi, bir sonraki bölümde ayrıntılandırılacak metodoloji bölümünün kuramsal gerekçesini oluşturmaktadır. Özellikle kutupluluk sözlüğü geliştirme süreci, sözlük destekli modelleme, morfoloji-duyarlı özellikler, alan uyarlaması ve açıklanabilirlik analizi, izleyen bölümde önerilen yaklaşımın temel bileşenleri olarak ele alınacaktır.

3. AZERBAIJAN DİLİ İÇİN DUYGU ANALİZİ ÇERÇEVESİ GELİŞTİRME METODOLOJİSİ

Bu bölümde Azerbaycan dili için önerilen duygu analizi çerçevesinin yöntemsel omurgası sunulmaktadır. Yedi aşamadan oluşan bu omurga, problem alanının ve kapsamın belirlenmesinden veri hazırlamaya, *SentiAzNet* geliştiriminden modelleme ve temsil öğrenmesine, oradan da açıklanabilirlik ve bütünselik değerlendirmeye uzanan sıralı bir araştırma akışını tanımlamaktadır.

İlk aşamada problem alanı ve kapsam belirlenmiş; ikinci aşamada veri toplama, temizleme ve ön işleme süreçleri yürütülmüştür. Üçüncü aşamada Azerbaycan dili için *SentiAzNet* kutupluluk sözlüğü geliştirilmiş; dördüncü aşamada bu bilginin finans ve iş dünyası alanındaki aspekt-tabanlı duygu analizi görevine nasıl entegre edilebileceği incelenmiştir. Beşinci aşamada *MahSA* adlı morfoloji-duyarlı çok alanlı hibrit mimari geliştirilmiş; altıncı aşamada alanlar arası sağlam cümle temsillerinin öğrenilmesine odaklanılmıştır. Yedinci ve son aşamada ise bütün yöntemler açıklanabilirlik, hata analizi ve karşılaştırmalı deneyler üzerinden bütünselik biçimde değerlendirilmiştir. Şekil 3.1, bu sıralı yapıyı ve özellikle üçüncü aşamada geliştirilen *SentiAzNet* bilgisinin dördüncü ile yedinci aşamalarda yeniden kullanıldığını göstermektedir.

Bu bölümde amaç, tek tek modellerin yalnızca teknik ayrıntılarını sıralamak değil; veri, dil kaynağı, hibrit modelleme, temsil öğrenmesi ve açıklanabilirlik bileşenlerini aynı yöntemsel omurga içinde birbirine bağlamaktır. Bu nedenle üçüncü bölüm, tez kapsamında yayımlanmış ve tamamlanmış çalışmaların tek bir bilimsel sav altında nasıl birleştiğini gösteren bütüncül bir çerçeve olarak kurgulanmıştır (Aykaç vd. 2025, 2026b; Aykaç vd. 2026; Aykaç vd. 2026a).

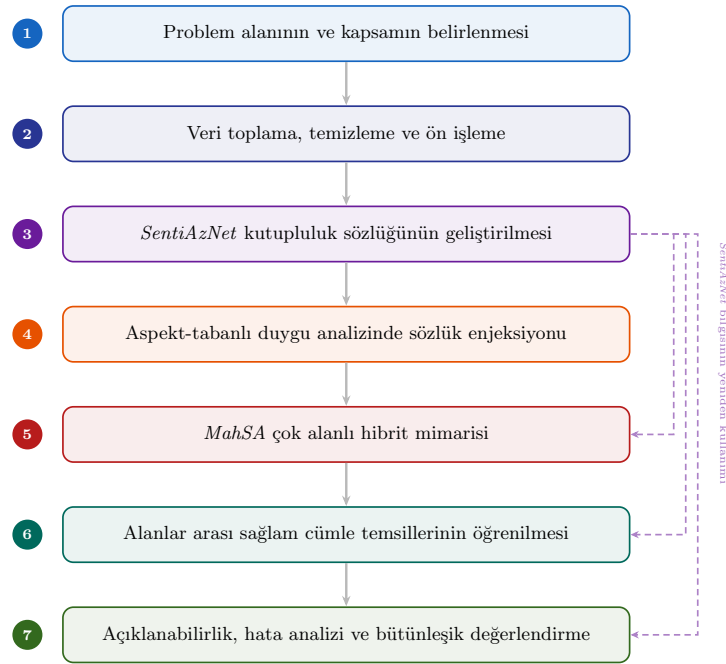
Bu tez, tasarım odaklı araştırma ile deneysel doğrulamayı birleştiren bir metodolojik yaklaşım izlemektedir. Önce Azerbaycan dili için duygu analizi problemi kaynak, morfoloji, alan kayması ve açıklanabilirlik eksenlerinde tanımlanmış; ardından bu probleme yanıt veren bir araştırma artefaktları ailesi tasarlanmıştır. Bu artefakt ailesi, *SentiAzNet* kutupluluk sözlüğünü, sözlük-enjekte edilmiş aspekt-tabanlı modeli, *MahSA* hibrit mimarisini ve alanlar arası sağlam temsil öğrenmesi hattını içermektedir. Son aşamada ise bu artefaktların katkısı ortak değerlendirme protokolleri, ablasyon deneyleri ve açıklanabilirlik çözümlemeleri üzerinden sınanmıştır. Böylece metodoloji bölümü yalnızca “nasıl” sorusunu yanıtlamamakta; aynı zamanda neden bu bileşenlerin

bu sırayla ele alındığını da görünür kılmaktadır.

3.1 Yedi Aşamalı Metodoloji ve Genel İş Akışı

Tezde benimsenen yöntem, birbirini izleyen ve belirli aşamalarda birbirini besleyen yedi aşamalı bir araştırma akışı üzerine kurulmuştur. Bu yapı, çalışmanın farklı yayın ve alt çıktılarının birbirinden bağımsız makaleler olarak değil, tek bir bilimsel savın kademeli olarak sınındığı bütünlük bir metodoloji olarak sunulmasına imkân vermektedir. Şekil 3.1 genel iş akışını, Çizelge 3.1 ise bu akışın bölüm içi karşılıklarını özetlemektedir.

Burada dikkat edilmesi gereken nokta, ikinci aşamanın metin içinde iki alt başlık altında ayrıştırılmasıdır: *Çalışmada Kullanılan Derlem ve Veri Kümeleri* (Bölüm 3.2) ile *Veri Toplama, Temizleme ve Ön İşleme Süreci* (Bölüm 3.3). Bu ayrışma, yedi aşamalı metodolojik şemanın yapısını bozmaz; aksine veri katmanını daha izlenebilir hale getirir.



Şekil 3.1 Tez kapsamında izlenen yedi aşamalı metodolojik iş akışı

Bu yedi aşama, tezin bilimsel katkılarıyla doğrudan hizalanmıştır. İlk üç aşama kaynak ve veri katkısını; dördüncü, beşinci ve altıncı aşamalar yöntem ve temsil öğrenmesi katkısını; yedinci aşama ise açıklanabilirlik ve hata çözümleme katkısını görünür kılmaktadır. Dolayısıyla bu

Çizelge 3.1 Tez metodolojisinin yedi temel aşaması

Aşama	Ana amaç	Tezdeki işlevi
1	Problem alanının ve kapsamın belirlenmesi	Duygu analizi problemini Azerbaycan dili için kaynak, morfoloji, alan kayması ve açıklanabilirlik eksenlerinde tanımlamak
2	Veri toplama, temizleme ve ön işleme	Derlem ve veri kümelerini oluşturmak; dil tanılama, tekrar ayıklama, normalizasyon ve alan eşleme süreçlerini yürütmek
3	<i>SentiAzNet</i> 'in geliştirilmesi	Azerbaycan dili için yeniden kullanılabilir kütüphane sözlüğü üretmek
4	Aspekt-tabanlı duygu analizinde sözlük enjeksiyonu	Sözlük bilgisinin bağlamsal modellere nasıl aktarılabilirliğini finans ve iş dünyası alanında sınamak
5	MahSA hibrit mimarisinin geliştirilmesi	Çok alanlı ve üç sınıflı duygu analizi için nöro-sembolik, morfoloji-duyarlı ana yöntemi kurmak
6	Alanlar arası sağlam temsil öğrenmesi	Duyguya duyarlı ve alan kaymasına karşı daha dayanıklı cümle temsilleri öğrenmek
7	Açıklanabilirlik, hata analizi ve bütünlük değerlendirme	Kararların hangi dilsel işaretlerden beslendiğini ve hangi hata bölgelerinin kaldığını birlikte analiz etmek

bölümün ilerleyen alt başlıkları, söz konusu metodolojinin ayrıntılı açılımı olarak okunmalıdır.

3.2 Çalışmada Kullanılan Derlem ve Veri Kümeleri

Tez kapsamında kullanılan materyaller tek bir veri kümesiyle sınırlı değildir. Çalışmanın farklı aşamalarında sözlük kaynakları, elle etiketlenmiş yorum derlemleri, alan-özgü aspekt düzeyinde veri kümeleri, büyük ölçekli etiketsiz derlemler ve açıklanabilirlik analizinde kullanılan yardımcı çıktılar birlikte değerlendirilmiştir. Bu tercih, Azerbaycan dili gibi düşük kaynaklı ve morfolojik açıdan zengin bir dilde duygu analizinin tek bir veri biçimine indirgenemeyeceği düşüncesine dayanmaktadır (Aykaç vd. 2026; Aykaç vd. 2026a).

Sözlük oluşturma aşamasında Bing Liu sözlüğü, NRC Emotion Lexicon, SenticNet, SentiWordNet, SentiTurkNet ve LexiPers gibi diğer dillere ait zengin kaynaklar kullanılmıştır (Aykaç vd. 2025). Alan-özgü modelleme aşamasında finans ve iş dünyası alanına ait yaklaşık 27.2 bin aspekt-düzeyle örnek içeren veri kümesi esas alınmıştır (Aykaç vd. 2026b). Çok alanlı genel modelleme aşamasında ise teknoloji ve dijital hizmetler, finans ve iş dünyası, sosyal yaşam ve eğlence, perakende ve yaşam tarzı ile kamu hizmetleri olmak üzere beş geniş alanı kapsayan 124.251 yorumluk etiketli derlem kullanılmıştır (Aykaç vd. 2026). Temsil öğrenmesi aşamasında bu alan taksonomisi daha büyük ölçekli bir yorum havuzuyla genişletilmiş ve alan kayması etkisi daha görünür hale getir-

ilmiştir (Aykaç vd. 2026a). Ayrıca bağlama kutupluluk yaymak için büyük ölçekli etiketsiz yorum derlemlerinden birlikte görülme istatistikleri hesaplanmıştır (Aykaç vd. 2026).

Kullanılan temel materyaller Çizelge 3.2’de özetlenmiştir. Bu çizelge, tezdeki her alt çalışmanın hangi veri veya kaynak üzerinde yükseldiğini açık hale getirmektedir.

Çizelge 3.2 Tez kapsamında kullanılan başlıca materyal ve veri kaynakları

Kaynak	Ölçek / tür	Kullanım amacı
Diğer dillere ait zengin tohum sözlükler	6 sözlük kaynağı	<i>SentiAzNet</i> ’in başlangıç aday havuzunu oluşturmak
SentiAzNet doğrulama verisi	2.482 lemma + dış yorum verisi	Sözlük kalitesini ve taşınabilirliğini sınamak
Finans ve iş dünyası ABSA verisi	Yaklaşık 27.2 bin aspekt örneği	Alan-özgü sözlük enjeksiyonu ve açıklanabilirlik incelemesi yapmak
MahSA çok alanlı derlem	124.251 kullanıcı yorumu	Üç sınıflı, çok alanlı genel duygu analizi modeli eğitmek ve değerlendirmek
Etiketsiz büyük derlem	Milyonlarca yorum / birlikte görülme istatistikleri	wPPMI ve bağlama kutupluluk yayılımı üretmek
Alanlar arası sağlam temsil derlemi	289.714 kullanıcı yorumu	Alan kaymasına dayanıklı cümle temsilleri öğrenmek

Bu materyallerin ortak özelliği, Azerbaycan dilinin hem günlük kullanıcı dilini hem de alan-özgü teknik kullanımlarını birlikte yansıtmasıdır. Tez kapsamında bu veriler, salt performans üretmek amacıyla değil; aynı zamanda morfolojik kutup kaymaları, nötr sınıf belirsizliği, kod değiştirimi ve alan bağımlılığı gibi güçlüklerin hangi koşullarda ortaya çıktığını görünür kılmak amacıyla kullanılmıştır. Bu nedenle veri seçimi, yöntem tasarımının ayrılmaz bir parçası olarak değerlendirilmiştir.

Elle etiketlenmiş veri kümelerinde etiketleme süreci tekil bir işlem olarak değil, görev-özel rehberler, pilot kontroller ve etiketleyiciler arası uyum hesaplarıyla desteklenen bir kalite güvence süreci olarak ele alınmıştır. Sözlük doğrulama hattındaki insan etiketlemesi ile çok alanlı *MahSA* derlemindeki etiketleme düzeni, ilgili alt bölümlerde ayrıntılı olarak açıklanmakta ve böylece veri kalitesi metodolojinin açık bir bileşeni haline getirilmektedir (Aykaç vd. 2025; Aykaç vd. 2026).

3.3 Veri Toplama, Temizleme ve Ön İşleme Süreci

Tez kapsamında kullanılan derlem ve veri kümeleri tek bir kaynaktan üretilmemiş; farklı alt problemler için farklı veri toplama ve hazırlama stratejileri izlenmiştir. Bununla birlikte tüm çalışmalarda ortak olan ön işleme mantığı aynıdır: önce dil tanılama yapılarak Azerbaycan dili dışındaki örnekler ayıklanmış, ardından yinelenen içerikler temizlenmiş, çok kısa ve bilgi taşımayan metinler elenmiş, sonrasında ise yazım çeşitliliğini azaltmaya yönelik hafif normalizasyon uygulanmıştır. URL, kullanıcı adı, etiket ve benzeri gürültü bileşenleri yer tutucu belirteçlere dönüştürülmüş; sosyal medya tarzı uzatmalar, olumsuzluk işaretleri ve yoğunlaştırıcı örüntüler ise duygu açısından taşıdığı bilgi korunacak biçimde işlenmiştir.

Bu ortak ön işleme hattı, bir yandan sözlük geliştirme aşamasında yüzey biçimlerin daha tutarlı eşleştirilmesine yardımcı olmuş, diğer yandan ABSA, MahSA ve alanlar arası sağlam temsil öğrenmesi deneylerinde bağlamsal modellerin gereksiz yazımsal gürültüye aşırı duyarlı hâle gelmesini engellemiştir. Daha ayrıntılı görev-özel ön işleme adımları ilgili alt bölümlerde ayrıca açıklanmaktadır.

Ön işleme tasarımıındaki temel ilke, gürültüyü azaltırken duygu taşıyan işaretleri korumaktır. Bu nedenle olumsuzluk yapıları, yoğunlaştırıcı örüntüler, tekrarlar ve duygu yönelimini etkileyebilecek yüzey ipuçları bütünüyle silinmemiş; tersine, ilgili görevlerde modelin yararlanabileceği biçimde korunmuş ya da yapılandırılmış özelliklere dönüştürülmüştür. Bu tercih, özellikle morfoloji-duyarlı modelleme çizgisinin sonraki aşamalarda tutarlı biçimde sürdürülmesini sağlamaktadır (Aykaç vd. 2026).

3.4 SentiAzNet Kutupluluk Sözlüğünün Geliştirilmesi

Yedi aşamalı metodolojinin üçüncü aşaması, Azerbaycan dili için özgün bir kutupluluk sözlüğü geliştirmeye ayrılmıştır. Buradaki amaç, yalnızca olumlu ve olumsuz sözcüklerin bir listesini üretmek değil; diğer dillere ait zengin kaynaklardan beslenen, anlamsal doğrulamadan geçen, insan etiketlemesi ile denetlenen ve daha sonra sinirsel modellere yardımcı bilgi olarak aktarılabilir yeniden kullanılabilir bir dil kaynağı oluşturmaktır. Bu nedenle *SentiAzNet* hattı, klasik sözlük üretiminden daha geniş bir mühendislik ve doğrulama çerçevesi içinde tasarlanmıştır (Aykaç vd.

2025). Ayrıca bu alt bölüm, tez planında erken aşamada öngörülen kutupluluk sözlüğü geliştirme, ön işleme, özellik çıkarımı ve model doğrulama adımlarını tek bir izlenebilir hat üzerinde birleştirmektedir.

SentiAzNet oluşturma sürecinde kullanılan çekirdek kaynaklar Çizelge 3.3’de özetlenmiştir. Kaynak seçimi yapılırken iki ilke benimsenmiştir: birincisi, İngilizce gibi yüksek kaynaklı dillerden gelen yerleşik duygu bilgisi mümkün olduğunca korunmalıdır; ikincisi ise Türkçe ve Farsça gibi tipolojik veya bölgesel yakınlığı olan dillerden gelen bilgi, hedef dildeki kültürel ve anlamsal kaymaları azaltmak için sürece dahil edilmelidir.

Çizelge 3.3 *SentiAzNet* hattında kullanılan temel tohum sözlük kaynakları

Kaynak	Düzy	Süreçteki işlevi
Bing Liu’s Opinion Lexicon	Sözcük	İkili olumlu/olumsuz kutupluluk için yüksek kesinlikli başlangıç adayları sağlamak
NRC Emotion Lexicon	Duygu kategorisi	Duygusal bayraklar ve kutuplulukla ilişkili yan anlamsal ipuçları üretmek
SenticNet 6	Kavram / kutup şiddeti	İnce taneli kutup şiddeti ve kavramsal duygu bilgisini sürece taşımak
SentiWordNet 3.0	Synset	WordNet tabanlı anlamsal hizalama ve synset denetimi için başvuru kaynağı olmak
SentiTurkNet 2.1	Türkçe synset	Tipolojik yakınlık sayesinde Azerbaycan dili için yardımcı kutupluluk bilgisi sağlamak
LexiPers	Farsça sözcük	Bölgesel ve kültürel yakınlık taşıyan kutupluluk örüntülerini desteklemek

3.4.1 Tohum sözlüklerin derlenmesi

İlk aşamada altı farklı duygu sözlüğü tek bir başlangıç havuzunda birleştirilmiştir (Aykaç vd. 2025). Bu aşama sonunda yinelenen girdiler ayıklanmış ve olumlu/olumsuz etiketler ortak bir gösterime dönüştürülmüştür. Sonuçta 9 614 benzersiz İngilizce lemma içeren bir tohum liste elde edilmiştir (Aykaç vd. 2025). Bu sayı, Azerbaycan dili için doğrudan elde bulunmayan kutupluluk bilgisinin çok kaynaklı biçimde toplanarak yeniden kullanılabilir hale getirildiğini göstermektedir. Tohum havuzun geniş tutulmasının temel nedeni, tek bir kaynağın kapsamadığı kültürel, anlamsal ve alan-özü boşlukları çoklu kaynak birleştirmesiyle azaltmaktır.

3.4.2 Çeviri, geri çeviri ve anlamsal doğrulama

Tohum havuzdaki İngilizce girdiler, otomatik çeviri araçları kullanılarak Azerbaycan diline aktarılmış ve her lemma için birden fazla aday karşılık tutulmuştur (Aykaç vd. 2025). Bu aşamada yalnızca doğrudan İngilizce-Azerbaycan dili eşlemesi yapılmamış; SentiTurkNet ve LexiPers kaynaklarındaki Türkçe ve Farsça kutupluluk girdileri de önce İngilizce tarafında hizalanmış, ardından Azerbaycan diline aktarılmıştır. Böylece hedef dilde tek yönlü çeviri kaynaklı kayıplar azaltılmaya çalışılmıştır.

Aday çevirilerin korunması için geri çeviri ve synset tabanlı anlamsal denetim birlikte kullanılmıştır. Kaynak İngilizce lemma e için üretilen Azerbaycanca adaylar kümesi $\mathcal{A}(e)$ içinden, geri çeviri sonrası WordNet eşanlam kümesi yeterince örtüşen adaylar tutulmuştur. Bu mantık aşağıdaki biçimde gösterilebilir:

$$\mathcal{A}^*(e) = \{a \in \mathcal{A}(e) \mid |\text{Syn}(\text{BT}(a)) \cap \text{Syn}(e)| \geq 2\} \quad (1)$$

Burada $\text{BT}(a)$, aday Azerbaycan dili teriminin geri çevrilmiş İngilizce karşılığını; $\text{Syn}(\cdot)$ ise ilgili synset kümesini göstermektedir. En az iki eşanlamlı örtüşmesi şartı, biçimsel olarak doğru görünen fakat hedef dilde farklı çağrışım taşıyan adayların elenmesini sağlamaktadır. Bu sürecin sonunda doğrulama ve eleme adımlarından geçmiş 2 482 Azerbaycan dili lemma içeren nihai aday liste elde edilmiştir (Aykaç vd. 2025).

3.4.3 Manuel doğrulama ve sözlük kalitesinin güvence altına alınması

Oluşturulan aday Azerbaycan dili sözlük listesi daha sonra iki yerli konuşur etiketleyici tarafından elle incelenmiştir (Aykaç vd. 2025). Doğrulama, yaklaşık %65’lik rastgele bir alt küme üzerinde bağımsız biçimde yürütülmüş; etiketleyiciler her girdi için üç değerli bir ölçek kullanmıştır: -1 olumsuz, 0 nötr ve $+1$ olumlu. Bu tercih, iki önemli amacı aynı anda yerine getirmektedir: bir yandan otomatik çeviri ve hizalama sürecinden sızan hataları temizlemekte, diğer yandan da kültürel bağlam, çok anlamlılık ve yerel kullanım farklarını doğrudan insan yargısıyla denetlemektedir.

Etiketleyiciler arası uyum $\kappa = 0,87$ olarak hesaplanmıştır; bu değer sözlüğün kutupluluk ata-

malarında yüksek düzeyde tutarlılık sağlandığını göstermektedir (Aykaç vd. 2025). Uyuşmazlık oluşturan girdiler ortak tartışma ile yeniden gözden geçirilmiş ve nihai etiketler bu aşamadan sonra sabitlenmiştir. Böylece SentiAzNet, yalnızca otomatik yöntemlerle derlenmiş bir liste olmaktan çıkarılarak dilsel açıdan denetlenmiş bir temel kaynak haline getirilmiştir.

3.4.4 Özellik çıkarımı ve sözlük tabanlı doğrulama modeli

SentiAzNet yalnızca derlenen bir sözlük olarak bırakılmamış; sözlüğün ayırt ediciliği ayrıca veri temelli bir doğrulama hattında sınanmıştır (Aykaç vd. 2025). Bu doğrulama tasarımında her lemma için duygusal bayraklar, morfolojik göstergeler, anlamsal gömmeler, derlem frekansı ve karakter n-gramları birlikte kullanılmıştır. Özellik uzayı şu şekilde özetlenebilir:

$$\mathbf{z}_i = \left[\mathbf{e}_i^{(emo)}; \mathbf{m}_i^{(morf)}; \mathbf{u}_i^{(LaBSE)}; f_i^{(azWaC)}; \mathbf{t}_i^{(char)} \right] \quad (2)$$

Burada $\mathbf{e}_i^{(emo)}$ sekiz boyutlu duygu bayraklarını, $\mathbf{m}_i^{(morf)}$ altı morfolojik göstergeden oluşan yüzey ipuçlarını, $\mathbf{u}_i^{(LaBSE)}$ 768 boyutlu LaBSE temsillerinin PCA ile 50 boyuta indirgenmiş sürümünü, $f_i^{(azWaC)}$ yaklaşık 94 milyon token içeren azWaC derleminden türetilen log-frekans değerini ve $\mathbf{t}_i^{(char)}$ 3–5 karakterlik n-gramlardan elde edilen TF-IDF vektörünü göstermektedir. Toplam boyut aşağıdaki gibidir:

$$\dim(\mathbf{z}_i) = 8 + 6 + 50 + 1 + 3000 = 3065 \quad (3)$$

Bu özellikler standartlaştırıldıktan sonra histogram tabanlı gradient boosting sınıflandırıcıya verilmiştir. Kullanılan karar fonksiyonu genel olarak şu biçimde gösterilebilir:

$$\hat{y}_i = g_{HGBT}(\mathbf{z}_i) \quad (4)$$

Burada $\hat{y}_i \in \{0, 1\}$, ilgili lemmaya atanan ikili kutupluluk kararını göstermektedir. Bu doğrulama hattında nötr terimler bilinçli olarak dışarıda bırakılmıştır; çünkü bu aşamanın amacı genel üç sınıflı bir sistem kurmak değil, yüksek kesinlikli bir olumlu-olumsuz kutupluluk çekirdeği üretmektir. Başka bir deyişle SentiAzNet, tezin ilerleyen bölümlerinde kullanılacak üç sınıflı modeller için karar verici son katman değil, sembolik ön bilgi kaynağı olarak düşünülmüştür.

Kullanılan başlıca özellik grupları ve boyutları Çizelge 3.4’de verilmektedir. Bu çizelge, ek-

lemeli morfolojinin yüzeyde bıraktığı ipuçları ile anlamsal gömme bilgisinin nasıl aynı doğrulama uzayında bir araya getirildiğini göstermektedir.

Çizelge 3.4 *SentiAzNet* doğrulama hattında kullanılan özellik grupları

Özellik grubu	Boyut	Amaç
Duygu bayrakları	8	NRC temelli temel duygusal eğilimleri açık biçimde kodlamak
Morfolojik göstergeler	6	Negasyon, yoğunlaştırma ve kutup değiştirici yapıları yüzeyden yakalamak
LaBSE + PCA	50	Lemma düzeyinde bağlamsal-anlamsal yakınlığı yoğun biçimde temsil etmek
azWaC log-frekansı	1	Lemmanın kullanım yaygınlığını ve derlem içi ağırlığını hesaba katmak
Karakter n-gram TF-IDF	3000	Kök-ek örüntülerini ve yazımsal düzenlilikleri temsil etmek
Toplam	3065	HGBT sınıflandırıcıya verilen bileşik özellik uzayı

3.4.5 Dış geçerleme

İç doğrulama sonuçları, geliştirilen sözlüğün denetimli bir test ortamında güçlü biçimde ayrışabildiğini göstermektedir (Aykaç vd. 2025). İkili sınıflandırma hattında doğruluk değeri 0,915, Macro-F1 değeri 0,892, Matthews korelasyon katsayısı 0,79 ve ROC-AUC değeri 0,98 olarak elde edilmiştir. Beş katlı çapraz doğrulamada Macro-F1 değeri yaklaşık $0,89 \pm 0,01$ aralığında kararlı kalmıştır. Bu sonuçlar, geliştirilen kutupluluk kaynağının yalnızca elle doğrulanmış değil, aynı zamanda veri temelli bir model altında da tutarlı olduğunu göstermektedir.

Sözlüğün gerçek kullanım ortamına taşınabilirliğini görmek için ayrıca yaklaşık 55 673 kullanıcı yorumundan oluşan dış bir veri üzerinde sına yapılmıştır (Aykaç vd. 2025). Google Play Store yorumları ve e-ticaret ürün değerlendirmeleri üzerinde yürütülen bu dış geçerleme adımında doğruluk 0,6324 olarak raporlanmıştır. Pozitif sınıf için kesinlik 0,77, negatif sınıf için duyarlık 0,78 düzeyindedir. İç deneylerle karşılaştırıldığında performansın düşmesi beklenen bir sonuçtur; çünkü dış veri daha gürültülü, bağlama daha bağımlı ve sözlük kapsamı dışında kalan çok sayıda yüzey biçim içermektedir. Bununla birlikte bu sonuç, *SentiAzNet*'in yalnızca laboratuvar koşullarında çalışan bir kaynak değil, doğal kullanıcı verisi üzerinde de anlamlı bir başlangıç

önbilgisi sunduğunu göstermektedir.

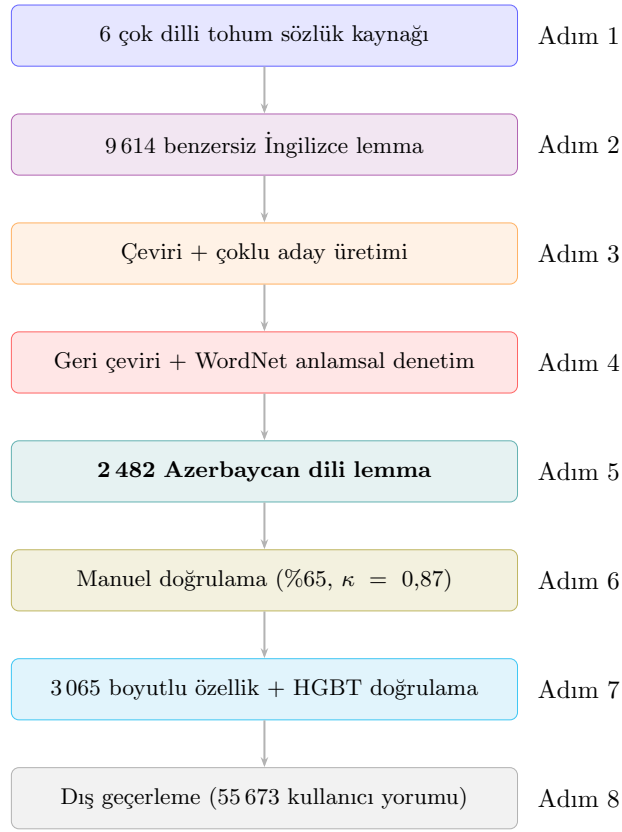
SentiAzNet oluşturma sürecinin nicel özeti Çizelge 3.5'te, çekirdek akış ise Şekil 3.2'te verilmiştir. Bu iki görsel birlikte okunduğunda, sözlük geliştirme hattının neden tezin geri kalanında kullanılan yöntemler için güvenilir bir başlangıç noktası ürettiği daha açık hale gelmektedir.

Çizelge 3.5 *SentiAzNet* oluşturma hattının nicel özeti

Bileşen / adım	Nicel bilgi	Yorum
Tohum havuz	9 614 İngilizce lemma	Diğer dillere ait altı zengin kaynaktan derlenen ve yinelenenleri ayıklanan başlangıç liste
Nihai aday sözlük	2 482 Azerbaycan dili lemma	Çeviri, geri çeviri ve synset tabanlı eleme sonrasında korunan girişler
Manuel doğrulama	%65 alt küme, 2 etiketleyici	Kültürel ve anlamsal uygunluğu denetlemek için kullanılan insan değerlendirmesi
Etiketleyiciler arası uyum	$\kappa = 0,87$	Güçlü tutarlılık düzeyi
Özellik uzayı	3065 boyut	Duygu, morfoloji, gömme, frekans ve karakter n-gram bileşimi
İç doğrulama	Acc.=0,915, Macro-F1=0,892	HGBT tabanlı ikili doğrulama hattı
Dış geçerleme	55 673 yorum, Acc.=0,6324	Gerçek kullanıcı verisi üzerindeki taşınabilirlik testi

3.5 Aspekt-Tabanlı Duygu Analizinde Sözlük Enjeksiyonu

Yedi aşamalı metodolojinin dördüncü aşamasında, genel cümle düzeyi kutupluluk sınıflandırmasının ötesine geçilerek belirli hedefe bağlı duygu çözümlemesi ele alınmaktadır. Aspekt-tabanlı duygu analizi literatürü, tek bir cümlede birden fazla hedefe ilişkin farklı duygu yönelimlerinin birlikte bulunabildiğini ve bu nedenle cümle düzeyi tek etiketli yaklaşımın çoğu zaman yetersiz kaldığını göstermektedir (Hua vd. 2024; Šmíd ve Král 2025). Finans ve iş dünyası metinlerinde bu sorun daha da görünür hale gelir; çünkü komisyon, faiz, hız, hizmet kalitesi ya da uygulama arayüzü gibi hedefler aynı yorum içinde karşıt biçimde değerlendirilebilir ve bazı terimlerin kutupluluğu alan bağlamına göre değişebilir (Loughran ve McDonald 2011; Araci 2019). Finans ve iş dünyası alanındaki ABSA alt problemi, bu nedenle tezde iki işlev üstlenmektedir: birincisi, *SentiAzNet* bilgisinin çağdaş bir bağlamsal modele nasıl aktarılabilceğini sınamak; ikincisi ise daha sonra *MahSA*'ya taşınacak nöro-sembolik tasarım ilkelerini daha kontrollü bir senaryoda doğrulamaktır.



Şekil 3.2 SentiAzNet geliştirme sürecinin nicel olarak özetlenmiş akışı

(Aykaç vd. 2026b, 2025).

3.5.1 Problem tanımı ve veri kümesinin özellikleri

Bu alt problemde her örnek, bir cümle, hedef aspekt, aspekt kategorisi ve duygu etiketinden oluşan bir dördümlü olarak ele alınmaktadır. Tez boyunca kullanılan gösterim aşağıdaki gibidir:

$$x = \langle s, a_t, a_c, y \rangle, \quad y \in \{\text{pozitif, nötr, negatif}\} \quad (5)$$

Burada s cümleyi, a_t hedef aspekt terimini, a_c aspekt kategorisini ve y ise hedefe bağlı duygu kutbunu göstermektedir. Bu tanım, aynı cümleden birden fazla aspekt-düzeyleli örnek üretilebilmesine izin vermektedir. Böylece bir yorumda hizmet olumlu, komisyon olumsuz ya da uygulama arayüzü nötr biçimde değerlendirildiğinde, model her hedef için ayrı karar üretmek zorunda kalmaktadır (Aykaç vd. 2026b).

Finans ve iş dünyası alanı için kullanılan veri kümesi yaklaşık 27.2 bin aspekt-düzeyle örnek içermektedir. Bu veri kümesinin temel istatistikleri Çizelge 3.6'de özetlenmiştir. Özellikle ortalama cümle uzunluğunun görece kısa olmasına rağmen, nötr sınıfın az temsil edilmesi ve alan terimlerinin kutupluluğunun bağlama göre değişmesi, bu problemi yüzeysel anahtar sözcük eşleştirmesinden daha zor hale getirmektedir (Aykaç vd. 2026b). Bu nedenle burada amaç, salt yüksek doğruluk üretmekten çok, sözlük bilgisinin hedefe bağlı bağlamsal kararlarla nasıl birleştirilebileceğini göstermektir.

Çizelge 3.6 Finans ve iş dünyası ABSA veri kümesinin temsili istatistikleri

İstatistik	Değer	Yöntemsel anlamı
Toplam aspekt örneği	~27.2 bin	Aynı cümleden birden fazla hedef-düzeyle örnek üretilebilmektedir
Eğitim / doğrulama / test	19 016 / 2 716 / 5 434	Temsili 70/10/20 ayrımı; esas raporlama tabakalı çapraz doğrulama ile yapılmaktadır
Sınıf dağılımı	Pozitif 11 826; Nötr 4 062; Negatif 11 278	Nötr sınıfın görece zor ve az temsil edilen bölge olduğunu göstermektedir
Ortalama cümle uzunluğu	18.4 token	Kısa fakat hedefe bağlı çok anlamlı yorumlar baskındır
Maksimum uzunluk	131 token	Uzun kuyrukta bağlamsal taşıma ve truncation riski vardır
Sözcük haznesi	42 850	Alan terimleri ve yazım çeşitliliği yüksektir
SentiAzNet token kapasitesi	%3.17	Sözlük sinyali seyrek; ana değil yardımcı kanal olarak düşünülmelidir
En az bir lexicon hit içeren örnek	~%38	Sözlük yalnızca belirli örneklerde aktif destek sağlamaktadır

3.5.2 Hedef-koşullu girdi gösterimi ve temel kodlayıcı

ABSA deneylerinde temel bağlamsal kodlayıcı olarak XLM-R kullanılmıştır (Conneau vd. 2020; Aykaç vd. 2026b). Girdi gösteriminde cümle ve hedef aspekt birlikte kodlanmakta; böylece kodlayıcı, aynı cümledeki farklı hedefler için farklı bağlamsal odaklar oluşturabilmektedir. Kullanılan dizi gösterimi şu biçimdedir:

$$\mathbf{X} = [<s>; \text{tok}(s); </s>; \text{tok}(a_t); </s>] \quad (6)$$

Burada $\text{tok}(\cdot)$, XLM-R tarafından kullanılan alt-sözcük parçalama işlemini göstermektedir. Kodlayıcı çıktısının ilk belirteç temsili, sınıflandırma başlığının girdisi olarak kullanılmaktadır. Nihai karar katmanını aşağıdaki biçimde özetlenebilir:

$$\hat{y} = \arg \max \text{softmax}(\mathbf{W}_o \mathbf{h}_{\text{cls}} + \mathbf{b}_o) \quad (7)$$

Bu kurulum, finans ve iş dünyası alanında cümle düzeyi tek etiketli yaklaşıma kıyasla daha ince taneli ve hedefe bağlı bir karar vermeyi mümkün kılmaktadır. Ayrıca tezin ilerleyen aşamalarında kullanılacak olan “bağlamsal kodlayıcı + sembolik yan kanal” mantığı da ilk kez bu alt problemde somutlaşmaktadır (Aykaç vd. 2026b).

3.5.3 Token düzeyinde sözlük hizalama ve alan düzeltmeleri

SentiAzNet sözlüğü yüzey sözcük düzeyinde tanımlandığından, çok dilli kodlayıcının alt-sözcük parçalarıyla doğrudan uyumlu değildir (Aykaç vd. 2025). Bu nedenle finans ve iş dünyası ABSA hattında önce yüzey sözcükleri sözlükle eşleştirilmiş, daha sonra bulunan kutupluluk değeri ilgili alt-sözcük parçalarına yayılmıştır (Aykaç vd. 2026b). Böylece sözlük eşleşmesi bilgisi, alt-sözcük tabanlı kodlayıcı içinde kaybolmadan temsil edilebilmiştir. Bu tasarım, sözlük kapsamı seyrek olduğundan özellikle önemlidir; çünkü burada amaç tüm token'lara kutupluluk atamak değil, yalnızca yüksek kesinlikli eşleşmeler bulunduğu anda modele kontrollü bir ön bilgi sağlamaktır.

Finans alanında bazı terimlerin genel dilden farklı kutupluluk davranışı göstermesi nedeniyle küçük ölçekli bir alan-özgü düzeltme listesi de kullanılmıştır. Bu düzeltme, tezde *SentiAzNet-Fin* olarak adlandırılan varyantı üretmektedir (Aykaç vd. 2026b). Alan düzeltmesi aşağıdaki parça-parça gösterimle ifade edilebilir:

$$\tilde{l}_i = \begin{cases} l_i^{\text{fin}}, & w_i \in \mathcal{V}_{\text{Fin}} \\ l_i, & \text{aksi halde} \end{cases} \quad (8)$$

Burada l_i genel sözlükten gelen kutupluluk değerini, l_i^{fin} finans alanı için düzeltilmiş kutbu, $\mathcal{V}_{\text{Fin}}$ ise alan düzeltmesi uygulanan sözcük kümesini göstermektedir. Bu mekanizma, tüm sözlüğü yeniden kurmadan yalnızca bariz alan uyumsuzluklarını gidermeyi amaçlamaktadır.

3.5.4 Özellik birleştirme ile sözlük enjeksiyonu

İlk enjeksiyon yaklaşımında sözlükten gelen kutupluluk sinyali, bağlamsal gösterime ek bir özellik olarak birleştirilmiştir (Aykaç vd. 2026b). Böylece kodlayıcının ürettiği temsil korunurken, yüksek kesinlikli sözlük eşleşmesi bilgisi karar katmanına açık biçimde aktarılmış olur. Bu birleşim, tezde aşağıdaki yardımcı dönüşümle gösterilmektedir:

$$\tilde{\mathbf{h}}_i = \mathbf{W}_{\text{cat}} [\mathbf{h}_i; \tilde{l}_i] + \mathbf{b}_{\text{cat}} \quad (9)$$

Burada $\mathbf{h}_i$ bağlamsal token temsilini, $\tilde{l}_i$ ise alan-düzeltilmiş kutupluluk sinyalini göstermektedir. Özellik birleştirme yaklaşımı mimari açıdan hafif, uygulanması kolay ve yorumlanabilir bir yan kanal sunduğu için başlangıç noktası olarak kullanılmıştır. Ancak bu yaklaşımda sözlük sinyali kodlayıcının iç dikkat dağılımını değiştirmez; yalnızca sonradan görülen ek bir öznitelik olarak davranır.

3.5.5 Dikkat önyargısı ile sözlük enjeksiyonu

İkinci yaklaşımda sözlük bilgisi, öz-dikkat skorlarına eklenen yumuşak bir önyargı terimi üzerinden kodlayıcının içine taşınmıştır (Aykaç vd. 2026b). Böylece sözlük eşleşmesi sözcükler yalnızca son katmanda değil, bağlam inşa edilirken de daha görünür hale gelir. Kullanılan temel dikkat kuralı şöyledir:

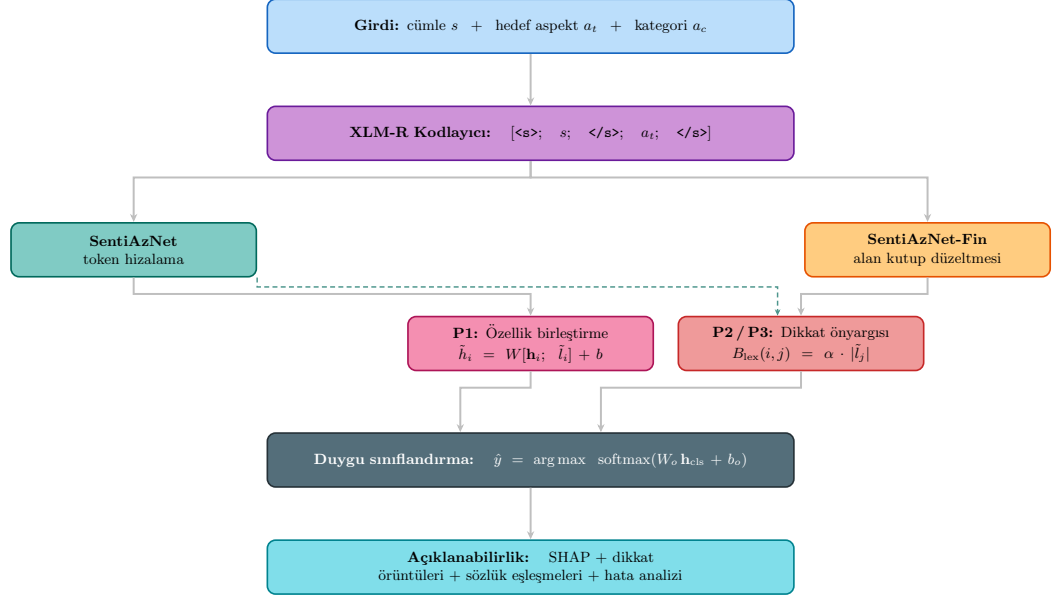
$$\text{Attn}(Q, K, V) = \text{softmax} \left(\frac{QK^\top}{\sqrt{d_k}} + B_{\text{lex}} \right) V \quad (10)$$

Bu denklemde yer alan sözlük önyargısı terimi ise aşağıdaki gibi tanımlanmıştır:

$$B_{\text{lex}}(i, j) = \alpha \cdot m_j, \quad m_j = |\tilde{l}_j| \quad (11)$$

Burada α öğrenilebilir ölçekleme katsayısını, m_j ise ilgili alt-sözcüğün sözlükten gelen kutupluluk büyüklüğünü göstermektedir. Bu sayede model, sözlük eşleşmesi sözcükleri tüm kararın taşıyıcısı yapmadan onlara dikkat ağırlığı düzeyinde yumuşak bir öncelik tanıyabilmektedir. Özellikle finans ve iş dünyası alanında kapsam seyrek olduğu için bu yaklaşım, sözlük bilgisini zorlayıcı değil tamamlayıcı bir sinyal olarak kullanmaktadır.

Şekil 3.3, finans ve iş dünyası ABSA alt probleminde izlenen tüm iş akışını şematik olarak göstermektedir. Çizelge 3.7 ise bu aşamada karşılaştırılan temel varyantları özetlemektedir. Bu iki görsel birlikte okunduğunda, alt problemin yalnızca bir model denemesi değil; *MahSA*'ya giden nöro-sembolik tasarım hattının ara doğrulama basamağı olduğu açık hale gelmektedir.



Şekil 3.3 Finans ve iş dünyası alanındaki ABSA alt probleminde izlenen sözlük enjeksiyonu ve açıklanabilirlik hattı

Çizelge 3.7 Finans ve iş dünyası alanındaki ABSA alt probleminde karşılaştırılan sözlük enjeksiyon varyantları

Varyant	Mimari değişiklik	Beklenen katkı
P1	Özellik birleştirme	Kutupluluk sinyalini bağlamsal vektöre hafif yardımcı özellik olarak eklemek
P2	Dikkat önyargısı	Sözlük eşleşmesi taşıyan sözcüklere dikkat düzeyinde yumuşak öncelik vermek
P3	P2 + <i>SentiAzNet-Fin</i>	Genel sözlüğü finans alanı için küçük kutup düzeltmeleriyle uyarlamak

3.5.6 Eğitim ve değerlendirme protokolü

Bu alt problemde bütün varyantlar aynı aspekt-koşullu giriş tasarımı ve aynı üç sınıflı karar başlığı altında karşılaştırılmıştır. Eğitim aşamasında çapraz entropi kaybı kullanılmış; erken durdurma ve doğrulama bölmesi üzerinden yapılan seçimlerle aşırı uyum riski sınırlandırılmıştır (Aykaç vd.

2026b). Nihai başarı ölçütü olarak Macro-F1'in seçilmesinin temel nedeni, nötr sınıfın hem daha zor hem de görece daha az temsil edilen bir sınıf olmasıdır. Buna ek olarak doğruluk, ağırlıklı F1 ve gerektiğinde istatistiksel anlamlılık testleri de değerlendirmeye dahil edilmiştir.

Finans ve iş dünyası ABSA hattında kullanılan deneysel protokol Çizelge 3.8'de özetlenmiştir. Burada özellikle tabakalı 5-katlı çapraz doğrulama tercihi önemlidir; çünkü veri bölmesine duyarlılığı azaltmakta ve sözlük enjeksiyonundan gelen kazançların tek bir rastlantısal ayırmadan kaynaklanmadığını göstermektedir. En iyi varyant ile temel model arasındaki farkların değerlendirilmesinde McNemar testi de kullanılmıştır (McNemar 1947; Aykaç vd. 2026b).

Çizelge 3.8 Finans ve iş dünyası alanındaki ABSA alt probleminde kullanılan deneysel protokol

Bileşen	Tercih	Amaç
Temel kodlayıcı	XLM-R	Çok dilli bağlamsal başlangıç sağlamak
Girdi kurgusu	Cümle + hedef aspekt	Duyguyu hedefe bağlamak
Eğitim kaybı	Çapraz entropi	Karşılaştırılabilir üç sınıflı karar başlığı kurmak
Ana başarı ölçütü	Macro-F1	Nötr sınıfı ve dengesizliği dengeli ölçmek
Doğrulama	Tabakalı 5-katlı çapraz doğrulama	Bölme duyarlılığını azaltmak
İstatistiksel test	McNemar	En iyi varyantın kazancını sınamak
Açıklanabilirlik	SHAP + dikkat örüntüleri + sözlük eşleşmesi	Kararı çoklu sinyalle yorumlamak

3.5.7 Açıklanabilirlik ve hata analizi

Bu alt çalışmada açıklanabilirlik, sonradan eklenen kozmetik bir unsur olarak değil, yöntemin davranışını denetleyen analitik bir katman olarak ele alınmıştır. SHAP atıfları, dikkat örüntüleri ve doğrudan sözlük eşleşmesi sinyalleri aynı örnek üzerinde birlikte incelenerek, modelin hangi durumlarda sözlüğe yaslandığı, hangi durumlarda bağlamsal temsilin baskınlaştığı ve bu iki kanalın ne ölçüde aynı sözcükleri öne çıkardığı araştırılmıştır (Lundberg ve Lee 2017; Jain ve Wallace 2019; Aykaç vd. 2026b). Bu amaçla top- k örtüşme ölçüleri de kullanılmıştır. Tezde kullanılan basit Jaccard-top- k tanımı şöyledir:

$$J@k(A, B) = \frac{|\text{Top}_k(A) \cap \text{Top}_k(B)|}{|\text{Top}_k(A) \cup \text{Top}_k(B)|} \quad (12)$$

Burada *A* ve *B*, örneğin SHAP ile dikkat örüntüsü tarafından en etkili görülen token kümelerini göstermektedir. Örtüşmenin yüksek olduğu örneklerde model kararının daha tutarlı davrandığı; düşük olduğu örneklerde ise ironi, nötr-negatif belirsizliği, karma duygu, örtük değerlendirme veya kod değiştirme gibi daha zor örüntülerin yoğunlaştığı gözlenmiştir (Aykaç vd. 2026b). Bu metodolojik bulgu, daha sonra *MahSA* bölümünde yapılacak hata analizi için doğrudan zemin hazırlamaktadır.

Sonuç olarak finans ve iş dünyası ABSA alt problemi, tezin genel savı açısından bir köprü işlevi görmektedir. Bu aşama sayesinde *SentiAzNet*'in çağdaş bir kodlayıcının içine yalnızca sonradan eklenen bir liste olarak değil, dikkat dağılımını ve temsil yapısını etkileyen yardımcı bir sinyal olarak yerleştirilebildiği gösterilmiştir. En önemlisi de burada elde edilen deneyim, *MahSA* mimarisinde sözlük, karakter, morfoloji ve bağlamsal temsil bloklarının neden birlikte düşünülmesi gerektiğini yöntemsel olarak meşrulaştırmıştır (Aykaç vd. 2026b; Aykaç vd. 2026).

3.6 MahSA Çok Alanlı Hibrit Mimarisi

Yedi aşamalı metodolojinin beşinci aşaması, tezin ana modelleme katkısı olan *MahSA* mimarisine ayrılmıştır. Bu mimari, Azerbaycan dili için yalnızca daha yüksek sınıflandırma başarımı hedeflememekte; aynı zamanda düşük kaynaklı ve morfolojik açıdan zengin bir dilde hangi bilgi kanallarının birlikte kullanıldığında daha dengeli, daha açıklanabilir ve alanlar arası daha dayanıklı sonuç verdiğini göstermeyi amaçlamaktadır. Bu nedenle *MahSA*, nöro-sembolik bir füzyon mimarisi olarak tasarlanmış; sözlük bilgisi, karakter düzeyi örüntüler, morfoloji-duyarlı yüzey işaretleri ve bağlamsal kodlayıcı temsilleri tek bir karar mekanizmasında birleştirilmiştir (Aykaç vd. 2026). Bu yapı, önceki aşamalarda sınanan *SentiAzNet* ve sözlük enjeksiyonu fikrini artık genel amaçlı ve çok alanlı bir duygu analizi çerçevesine taşımaktadır.

3.6.1 Tasarım problemi ve kurucu ilkeler

MahSA'nın çıkış noktası üçlü bir güçlük kümesidir. Birincisi, Azerbaycan dili için geniş ölçekli, elle etiketlenmiş ve alan çeşitliliği yüksek veri kümeleri sınırlıdır. İkincisi, dilin eklemeli yapısı nedeniyle kutupluluk çoğu zaman yalnızca kökte değil; olumsuzluk, yoksunluk, pekiştirme ve uzatma gibi biçimbirimsel veya yüzeysel işaretlerde ortaya çıkmaktadır. Üçüncüsü ise aynı mod-

elin teknoloji, finans, sosyal yaşam, perakende ve kamu hizmetleri gibi anlam dağılımları farklı alanlarda tutarlı davranması beklenmektedir. MahSA, bu üç sorunu aynı anda ele alan bir tasarım olarak düşünülmüştür (Aykaç vd. 2026).

Bu bağlamda mimarinin kurucu ilkeleri şu şekilde özetlenebilir: açık sembolik bilgi bütünüyle terk edilmemeli; bağlamsal kodlayıcı tek başına yeterli varsayılmamalı; morfoloji bilgisi tam çözümleme zorunluluğu olmadan da temsil edilebilmeli; ve nötr sınıf, ikincil bir artık sınıf değil, ayrı bir modelleme problemi olarak ele alınmalıdır. Bu nedenle MahSA'da *morpheme-aware* ifadesi, tam biçimbirim çözümleme kullanan ağır bir dil işleme boru hattını değil; gürültülü kullanıcı metninde daha dayanıklı çalışan yüzey sezgileri, karakter ipuçları ve kutup değiştirici örüntülerin sistematik olarak modellenmesini ifade etmektedir.

3.6.2 Çok alanlı derlemin oluşturulması ve etiketleme

MahSA için oluşturulan derlem, 2021–2023 arasında e-ticaret platformları, haber portalları, sosyal medya kaynakları ve hizmet değerlendirme sitelerinden toplanan kullanıcı yorumlarından oluşmaktadır. Ham veri; dil tanılama, yinelenen içerik ayıklama ve temel gürültü temizleme adımlarından geçirildikten sonra beş yüksek düzeyli alana eşlenmiştir: teknoloji ve dijital hizmetler, finans ve iş dünyası, sosyal yaşam ve eğlence, perakende ve yaşam tarzı ve kamu hizmetleri. Alan eşleme sürecinde kaynak üstverisi güvenilir olduğunda doğrudan kullanılmış; belirsiz durumlarda ise yüksek güvenli alan yeniden etiketleme adımlarıyla gürültü azaltılmıştır. Nihai derlem 124 251 yorum içermektedir (Aykaç vd. 2026).

Çizelge 3.9 MahSA çok alanlı derleminde alan dağılımı

Alan	Yorum sayısı	Pay
Teknoloji ve dijital hizmetler	39 867	%32,1
Finans ve iş dünyası	27 166	%21,9
Sosyal yaşam ve eğlence	31 308	%25,2
Perakende ve yaşam tarzı	16 548	%13,3
Kamu hizmetleri	9 362	%7,5
Toplam	124 251	%100,0

Her yorum negatif, nötr ve pozitif olmak üzere üç sınıftan biri ile etiketlenmiştir. Etiketleme süreci üç yerli konuşurun bağımsız etiketlemeleri üzerine kurulmuş; çoğunluk oyu ile karar ver-

ilemeyen durumlar kıdemli değerlendirici denetimiyle çözülmüştür. Özellikle nötr sınıfın, olgusal ifade, karışık değerlendirme ve zayıf öznel dil arasında kaldığı örneklerde daha sorunlu olduğu gözlenmiştir. Buna rağmen genel Fleiss katsayısı 0,68 düzeyine ulaşmış (Aykaç vd. 2026) ve bu değer, çok alanlı ve kullanıcı üretimli gürültülü veride kabul edilebilir düzeyde bir tutarlılığa işaret etmiştir. Sınıf bazlı uyumun en düşük olduğu bölgenin nötr sınıf olması, model tasarımında nötr kalibrasyonunun neden ayrı bir bileşen olarak ele alındığını da açıklamaktadır.

Çizelge 3.10 MahSA derleminde sınıf dağılımı ve kısa derlem istatistikleri

Alan	Neg. (%)	Nötr (%)	Poz. (%)	Ort. token
Teknoloji/dijital	38,4	27,8	33,8	19,4
Finans	35,2	25,6	39,2	21,1
Sosyal	26,7	42,4	30,9	24,7
Perakende	33,5	22,1	44,4	15,3
Kamu	46,3	36,2	17,5	26,8

Çizelge 3.11 MahSA etiketleme sürecinde sınıf bazlı ve genel uyum

Sınıf	Fleiss κ	Anlaşma (%)
Negatif	0,74	84,2
Nötr	0,58	71,3
Pozitif	0,71	82,6
Genel	0,68	79,4

MahSA hattında ayrıca bağlam kutupluluğu yaymak üzere etiketsiz büyük bir yorum havuzu da kullanılmıştır. Bu havuz, aynı kaynak ailesinden derlenen yaklaşık 2,1 milyon yorum ve yaklaşık 45 milyon token içermekte; (Aykaç vd. 2026) sözlükte doğrudan yer almayan ancak belirli kutuplu ifadelerle düzenli olarak birlikte görülen sözcükler için ek bağlamsal sinyal üretmektedir. Böylece yöntem, yalnızca 124 251 etiketli yorumdan değil, daha geniş bir kullanım evreninden de yararlanmaktadır.

3.6.3 Ön işleme ve morfoloji-duyarlı normalizasyon

Ön işleme boru hattı, klasik metin temizliğinden daha geniş tutulmuştur. Önce dil tanılama ile Azerbaycan diline ait olmayan veya çok kısa olduğu için güvenilir biçimde tanılanamayan örnekler ayıklanmıştır. Ardından spam ve boilerplate içerikler temizlenmiş; tam ve yakın kopya yorumlar karakter n-gram benzerliği kullanılarak veri havuzundan çıkarılmıştır. Bu adım, aynı şikâyetin veya aynı övgü kalıbının çok kez tekrar edilmesinden kaynaklanan yapay örnek yoğunluğunu azaltmak

için gereklidir.

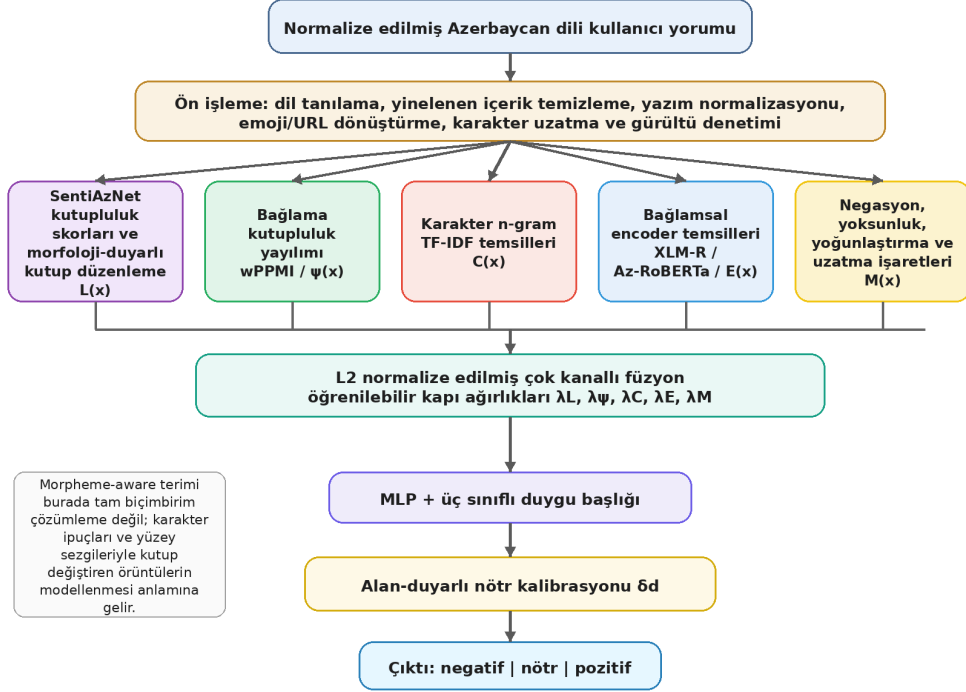
Yazı düzeyinde ise Latin harf envanteri korunmuş, karışık yazım biçimleri tekilleştirilmiş, diakritikler muhafaza edilmiş ve sosyal medya tarzı uzatmalar ayrı bir işaret olarak saklanacak biçimde sınırlandırılmıştır. Örneğin karakter tekrarı üçten fazla olduğunda yüzey biçim kısmen normalize edilmiş; fakat uzatma bilgisinin varlığı ek bir özellik olarak korunmuştur. URL, kullanıcı adı ve etiket gibi metinsel gürültüler yer tutucu tokenlara dönüştürülmüş; emoji kaba kutupluluk taşıyan özel belirteçler biçiminde temsil edilmiştir. Böylece hem kodlayıcı bileşeni hem de karakter blokları, çok farklı yüzey biçimlerinden kaynaklanan seyrekliği daha yönetilebilir hale getirmektedir.

Bu aşamada kritik tasarım kararı, tam morfolojik ayrıştırma kullanmak yerine yüzey tabanlı morfoloji sezgilerine dayanılmasıdır. Azerbaycan dili eklemeli bir dil olmakla birlikte, kullanıcı üretimi metinlerde yazım bozukluğu, eksik diakritik, kısaltma ve kod değiştirimi nedeniyle ağır biçimbirim çözümleyicilerin güvenilirliği düşebilmektedir. Bu nedenle MahSA; olumsuzluk ekleri, yoksunluk biçimleri, değılleme parçacıkları, pekiştiriciler ve karakter uzatmalarını açık biçimde işaretleyen daha hafif bir yaklaşımı benimsemiştir (Aykaç vd. 2026). Bu tercih, yöntemin hem taşınabilirliğini hem de gürültülü veri üzerindeki kararlılığını artırmaktadır.

3.6.4 Özellik blokları ve temsil notasyonu

MahSA, her yorum için beş tamamlayıcı bilgi bloğı üretmektedir: sözlük tabanlı kutupluluk özeti $L(x)$, bağlama kutupluluk yayılımı $\psi(x)$, karakter n-gram temsilleri $C(x)$, bağlamsal kodlayıcı temsilleri $E(x)$ ve morfoloji-duyarlı yüzey ipuçları $M(x)$. Bu blokların hiçbirisi tek başına yeterli varsayılmamaktadır. Bunun yerine, her bloğun eksik bıraktığı bilgi diğer bloklarla tamamlanmakta; örneğin sözlük yüksek kesinlikli ama düşük kapsamlı, karakter blokları gürültüye dayanıklı ama bağlama duyarsız, kodlayıcı ise bağlama güçlü ama kutup değıştiren ince yüzey işaretlerine karşı her zaman yeterince hassas değıldir.

MahSA çok alanlı hibrit duygu analizi mimari akışı



Şekil 3.4 MahSA mimarisinin genel akışı ve temel bilgi blokları

Çizelge 3.12 MahSA’da kullanılan temel özellik blokları

Blok	Boyut niteliği	İçerik ve işlev
$L(x)$	Küçük yoğun vektör	SentiAzNet eşleşmelerinden türetilen pozitif/negatif toplamalar, ortalamalar, en büyük mutlak kutup ve hit oranı gibi özetler
$\psi(x)$	3 boyut	Pozitif, negatif ve mutlak bağlam kutupluluğunu taşıyan wPPMI tabanlı bağlam vektörü
$C(x)$	Yaklaşık 42 000	3–5 karakter uzunluklarında TF-IDF n-gramları; gürültülü yazım ve alt-sözcük kırılmalarına karşı dayanıklı blok
$E(x)$	128 boyut	XLM-R tabanlı bağlamsal kodlayıcı çıktısının doğrusal izdüşüm ile sıkıştırılmış gösterimi
$M(x)$	Küçük yoğun vektör	Negasyon, değilleme, yoksunluk, yoğunlaştırma ve uzatma örüntülerini taşıyan sayısal/binary işaretler

3.6.5 Lexicon scoring ve morfoloji-duyarlı kutup düzenleme

MahSA içinde *SentiAzNet* sinyali doğrudan kullanılmakta; ancak bu sinyal olumsuzluk ve yoğunlaştırma gibi yüzey işaretleriyle yeniden düzenlenmektedir. Tez kapsamında kullanılan temel biçim aşağıdaki gibidir:

$$s_i = \begin{cases} 0, & |s_{\text{raw}}(w_i)| < \tau_{\text{lex}} \\ -\gamma_{\text{neg}} s_{\text{raw}}(w_i), & w_i \text{ olumsuzluk kapsamındaysa} \\ \gamma_{\text{int}} s_{\text{raw}}(w_i), & w_i \text{ bir yoğunlaştırıcı ile değiştirildiyse} \\ s_{\text{raw}}(w_i), & \text{diğer durumlarda} \end{cases} \quad (13)$$

Burada τ_{lex} çok zayıf kutupluluk değerlerini sıfırlayan eşik, γ_{neg} olumsuzluk etkisi, γ_{int} ise yoğunlaştırma katsayısıdır. Geliştirme deneylerinde $\tau_{\text{lex}} = 0,10$, $\gamma_{\text{neg}} = 1$ ve $\gamma_{\text{int}} = 1,5$ ayarları en kararlı sonuçları vermiştir. Bu tasarım sayesinde *yaxşı deyil*, *pis deyil*, *çooox yaxşı* gibi yüzey biçimleri yalnızca kök düzeyinde değil, bağlamda üstlendikleri kutupsal işlev bakımından da temsil edilebilmektedir.

3.6.6 Bağlama kutupluluk yayılımı: wPPMI

Sözlükte yer almayan fakat belirli kutuplu sözcüklerle aynı bağlamlarda düzenli olarak görülen sözcükleri değerlendirebilmek için, büyük ölçekli etiketsiz derlem üzerinde pencere-temelli birlikte görülme istatistikleri hesaplanmıştır. Bu amaçla kullanılan ağırlıklı pozitif PMI bileşeni aşağıdaki gibi tanımlanmıştır:

$$\text{wPPMI}_{\alpha}(t, c) = \max \left(\log \frac{p(t, c)}{p(t) p(c)^{\alpha}}, 0 \right) \quad (14)$$

Burada t sözlükte bulunan kutuplu terimi, c bağlam sözcüğünü, α ise seyrek bağlamları yumuşatan düzeltme katsayısını göstermektedir. Pratikte $\alpha = 0,75$ ve simetrik pencere boyutu $w = 5$ seçilmiştir (Aykaç vd. 2026). Yorum düzeyinde toplanan skalar bağlam sinyali şu biçimde hesaplanmaktadır:

$$\tilde{\psi}(x) = \sum_{i=1}^n \sum_{j \in N_w(i)} s(t_i) \text{wPPMI}_{\alpha}(t_i, w_j) \quad (15)$$

Bu skalar daha sonra pozitif, negatif ve mutlak bileşenlere ayrılarak küçük boyutlu bağlam vektörüne dönüştürülmektedir:

$$\psi(x) = [\psi^+(x), \psi^-(x), |\tilde{\psi}(x)|] \quad (16)$$

Bu mekanizma, örneğin doğrudan kutupluluk sözlüğünde yer almayan ama belirli şikâyet kalıplarıyla sık görülen *növb@*, *şikay@t*, *donub* gibi sözcüklerin dolaylı kutupsal yük taşımasını mümkün kılmaktadır. Dolayısıyla sözlük kapsamı dar olduğunda bile bağlam kanalının yardımıyla duygu sinyali daha geniş bir sözcük alanına yayılabilmektedir.

3.6.7 Karakter n-gramları, bağlamsal kodlayıcı ve morfoloji blokları

Karakter n-gram bileşeni, özellikle gürültülü kullanıcı metninde biçimsel varyasyonu dengelemek için tasarlanmıştır (Aykaç vd. 2026). $n \in \{3, 4, 5\}$ karakter uzunlukları kullanılmış; eğitim bölünmesinde çok seyrek görülen n-gramlar elendikten sonra yaklaşık 42 000 boyutlu bir TF-IDF uzayı elde edilmiştir. Bu blok; yazım bozukluğu, uzatma, eksik diakritik ve alt-sözcük kırılmaları durumunda bile kök ve ek örüntülerinin bir bölümünü koruyabildiği için, Azerbaycan dili gibi eklemeli bir dilde önemli bir dayanıklılık kanalı oluşturmaktadır.

Bağlamsal temsil tarafında temel kodlayıcı olarak XLM-R kullanılmıştır. Normalleştirilmiş yorum doğrudan SentencePiece tokenizer'a verilmiş; elde edilen ham cümle temsili daha sonra 128 boyutlu daha sıkı bir uzaya izdüşürülmüştür:

$$E(x) = \mathbf{W}_E E_{\text{raw}}(x) \quad (17)$$

Bu izdüşüm, yüksek boyutlu bağlamsal vektörün füzyon katmanında karakter ve sözlük blokları ile daha dengeli biçimde birleştirilmesine imkân tanımaktadır. Morfoloji bloğu ise her yorum için negasyon ek sayısı, değilleme parçacığı, yoğunlaştırıcı, karakter uzatma ve yoksunluk biçimleri gibi sayısal ya da ikili işaretler üretmektedir. Bu blok, tam morfolojik çözümleyicinin başarısına bağımlı olmadan, kutup değiştirici örüntülerin doğrudan karar mekanizmasına görünür olmasını sağlar.

3.6.8 Füzyon katmanı ve duygu başlığı

MahSA'nın temel gücü, yukarıdaki bilgi bloklarını tek bir birleşik temsilde toplamasından gelmektedir. Önce tüm bloklar L2-normalize edilmekte, ardından öğrenilebilir ve negatif olmayan kapı ağırlıkları ile ölçeklenerek aşağıdaki birleşik gösterim oluşturulmaktadır:

$$F(x) = [\lambda_L L(x); \lambda_\psi \psi(x); \lambda_C C(x); \lambda_E E(x); \lambda_M M(x)] \quad (18)$$

Burada λ katsayıları, her blok için göreceli güven düzeyini temsil etmektedir. Başka bir deyişle model, her örnekte aynı bilgi kaynağına aynı ölçüde yaslanmaz; alanın, yorumun ve yüzey örüntülerinin niteliğine göre sözlük, karakter, bağlam veya morfoloji kanalını daha baskın hale getirebilir. Birleşik vektör daha sonra küçük bir çok katmanlı algılayıcıya ve üç-yollu bir duygu başlığına verilerek negatif, nötr ve pozitif kararları üretilmektedir.

3.6.9 Nötr sınıf için alan-duyarlı kalibrasyon

Üç sınıflı duygu analizinde en sorunlu bölgenin çoğu zaman nötr sınıf olması nedeniyle, *MahSA* içinde alan-duyarlı bir kalibrasyon adımı da kullanılmıştır. Bu mekanizma aşağıdaki karar kuralı ile özetlenebilir:

$$\hat{y}(x, d) = \begin{cases} \text{neutral}, & p_{\text{neu}} + \delta_d > \max(p_{\text{neg}}, p_{\text{pos}}) \\ \arg \max_{y \in \{\text{neg}, \text{pos}\}} p_y, & \text{aksi halde} \end{cases} \quad (19)$$

Burada δ_d , ilgili alan için nötr sınıf eşikini ayarlayan alan-özgü marjı göstermektedir. Geliştirme aşamasında önce küresel bir δ taranmış, daha sonra alan bazında daha küçük ayarlamalar yapılmıştır. Bu tasarım özellikle sosyal ve kamu hizmetleri gibi nötr örnek oranının yüksek olduğu alanlarda, sistemin gereksiz biçimde kutuplu alarm üretmesini azaltmak için etkilidir.

3.6.10 Eğitim ayrıntıları ve hiperparametreler

MahSA eğitiminde varsayılan amaç fonksiyonu standart üç sınıflı çapraz entropidir. Sınıf dengesizliğini azaltmak için sınıf ağırlıkları ters frekans mantığıyla hesaplanmıştır:

$$w_c = \frac{N}{3 \cdot N_c} \quad (20)$$

Burada N toplam eğitim örneği sayısını, N_c ise ilgili sınıftaki örnek sayısını göstermektedir. Bunun yanında, büyük alanların küçük alanları baskılamasını azaltmak için alan-ağırlıklı örnekleme ya da kayıp ağırlıklandırması da kullanılmıştır; bu yaklaşım kabaca $w_d \propto 1/N_d$ mantığına dayanır. Eğitim süreci AdamW, erken durdurma ve birden çok rastgele tohum ile kararlı hale getirilmiştir (Aykaç vd. 2026).

Çizelge 3.13 MahSA için kullanılan temel eğitim ayarları

Hiperparametre	Değer
Kodlayıcı öğrenme oranı	2×10^{-5}
Füzyon / MLP öğrenme oranı	1×10^{-3}
Batch boyutu	32
Azami alt-sözcük uzunluğu $T_{\max}$	128
MLP gizli katman boyutu	256
Dropout	0,20
γ_{neg}	1
γ_{int}	1,5
wPPMI pencere boyutu w	5
wPPMI yumuşatma katsayısı α	0,75
Küresel nötr marj δ	0,05
Rastgele tohum sayısı	$k = 5$

Bu yapı sayesinde MahSA, bir yandan sözlük ve morfoloji gibi insan tarafından yorumlanabilir bilgi kaynaklarını görünür kılmakta; diğer yandan çok dilli bağlamsal temsillerin esnekliğini korumaktadır (Aykaç vd. 2026; Aykaç vd. 2025). Tezin deneysel bölümünde gösterileceği üzere, özellikle küçük ve nötr-ağırlıklı alanlarda elde edilen kazanımlar, bu çok kanallı tasarımın yalnızca sayısal değil yönetsel olarak da gerekli olduğunu göstermektedir.

3.7 Alanlar Arası Sağlam Cümle Temsillerinin Öğrenilmesi

Yedi aşamalı metodolojinin altıncı aşamasında, sınıflandırma kararını doğrudan iyileştirmekten çok, duyguya duyarlı cümle temsillerinin alanlar arasında ne ölçüde kararlı kalabildiği incelenmektedir. Burada amaç, aynı duyguyu taşıyan örneklerin farklı alanlardan gelseler bile gömme uzayında birbirine yakın konumlanmasını sağlamak; buna karşılık yalnızca alana özgü ama duygu açısından taşınabilir olmayan örüntülerin temsili baskılamasını engellemektir. Böylece bu alt

çalışma, *MahSA*'nın çok kaynaklı füzyon mantığını temsil öğrenmesi düzeyine taşıyarak tezin teorik katkısını daha genel bir zeminde sınamaktadır (Aykaç vd. 2026a; Aykaç vd. 2026; Aykaç vd. 2025).

Bu aşamada kullanılan temel varsayım şudur: Azerbaycan dili gibi düşük kaynaklı ve morfolojik açıdan zengin bir dilde alan kayması, yalnızca karar sınırlarını değil, gömme uzayının geometrisini de bozmaktadır. Başka bir ifadeyle aynı sınıfa ait örnekler farklı alanlarda birbirinden uzaklaşmakta, özellikle nötr sınıf kutupsal bölgeler arasına dağılmakta ve kodlayıcı tabanlı modeller alan-özgü yüzey kısayollarına gereğinden fazla yaslanabilmektedir. Alanlar arası sağlam cümle temsilleri bölümü, bu problemi azaltmak için alan-uyarlamalı ön eğitim, alan-karşıtı düzenleme, kutupluluk ve morfoloji enjeksiyonu ile denetimli karşıtsal düzenlileştirmeyi tek bir temsil-merkezli çerçevede bir araya getirmektedir (Gururangan vd. 2020; Ganin vd. 2016; Khosla vd. 2020; Ramponi ve Plank 2020; Aykaç vd. 2026a).

3.7.1 Problem tanımı ve veri kurgusu

Bu alt problemde amaç, her yorum x ve alan etiketi d için üç sınıflı duygu etiketi y üretmek ve aynı zamanda bu kararın dayandığı cümle gömme uzayını alan kaymasına karşı dayanıklı hale getirmektir. Görev, tez boyunca aşağıdaki biçimde gösterilmektedir:

$$f : (x, d) \mapsto y, \quad y \in \{\text{neg, neu, pos}\}, \quad d \in \mathcal{D}, \quad |\mathcal{D}| = 5 \quad (21)$$

Burada $\mathcal{D}$ kümesi teknoloji ve dijital hizmetler, finans ve iş dünyası, sosyal yaşam ve eğlence, perakende ve yaşam tarzı ile kamu hizmetleri alanlarından oluşmaktadır. Bu veri kurgusu, bir önceki MahSA çalışmasında kurulan alan taksonomisinin korunmasını; ancak daha büyük ölçekli bir yorum havuzu üzerinde temsil geometrisinin ayrıca incelenmesini mümkün kılmaktadır (Aykaç vd. 2026a; Aykaç vd. 2026). Böylece model yalnızca “hangi etiket doğru?” sorusuna değil, “aynı etiket farklı alanlarda nasıl temsil edilmeli?” sorusuna da yanıt vermek zorunda kalmaktadır.

Alanlar arası sağlam temsil öğrenmesi aşamasında kullanılan yorum derlemi toplam 289 714 örnekten oluşmaktadır. Bu derlemin alanlara göre dağılımı Çizelge 3.14’de verilmiştir. Teknoloji ve finans alanları veri hacmi bakımından baskın görünse de, kamu hizmetleri ve sosyal yaşam alanları nötr sınıfın daha yüksek olduğu ve alan kaymasının daha belirgin hissedildiği bölgeler

olarak özel önem taşımaktadır. Bu nedenle veri kümesi yalnızca büyük olmakla değil, alan çeşitliliği ve etiket dağılımı bakımından da temsili bir yapı sunmaktadır (Joshi vd. 2020; Isbarov vd. 2024).

Çizelge 3.14 Alanlar arası sağlam temsil öğrenmesi aşamasında kullanılan derlemin alan dağılımı

Alan	Yorum sayısı	Pay
Teknoloji ve dijital hizmetler	85 847	29.63%
Finans ve iş dünyası	70 712	24.41%
Sosyal yaşam ve eğlence	59 024	20.37%
Perakende ve yaşam tarzı	38 319	13.23%
Kamu hizmetleri	35 812	12.36%
Toplam	289 714	100.00%

Toplam etiket dağılımı negatif 72 138, nötr 121 453 ve pozitif 96 123 örnek biçimindedir. Nötr sınıfın toplam veri içindeki payının yüksek olması, alanlar arası sağlamlık problemini özellikle bu sınıf üzerinde görünür hale getirmektedir. Etiket güvenilirliğini denetlemek için 2 847 örneklik tabakalı bir alt kümede üçlü etiketleme yapılmış; genel Fleiss κ değeri 0.73, nötr sınıf için ise 0.62 olarak ölçülmüştür. Bu sonuç, nötr sınıfın yalnızca modelleme açısından değil, insan etiketlemesi açısından da en belirsiz bölge olduğunu göstermektedir (Fleiss 1971; Li vd. 2023). Çizelge 3.15, bu bölümün temsil öğrenmesi mantığını belirleyen temel veri kalitesi göstergelerini özetlemektedir.

Çizelge 3.15 Etiket dağılımı ve etiketleme güvenilirliği için temel göstergeler

Gösterge	Değer	Yorum
Toplam etiket dağılımı	Neg.: 72 138; Nötr: 121 453; Poz.: 96 123	Nötr sınıfın baskın ve ayrıştırılması güç olduğunu göstermektedir
Üçlü etiketleme alt kümesi	2 847 yorum	Etiket güvenilirliğini ayrı bir kalite denetimi ile sınamak için kullanılmıştır
Fleiss κ (genel)	0.73	Genel düzeyde güçlü tutarlılık
Fleiss κ (nötr)	0.62	Nötr sınıfın sınırlarının daha belirsiz olduğuna işaret etmektedir

3.7.2 Kaynak seçimi, filtreleme ve dil denetimi

Bu aşamada veri kaynakları üç ana gruptan derlenmiştir: YouTube üzerindeki kullanıcı yorumları, Google Play Store uygulama değerlendirmeleri ve yerel e-ticaret sitelerindeki ürün yorumları. Böylece teknoloji, finans, sosyal yaşam, perakende ve kamu hizmetleri alanlarının her biri için hem kısa hem uzun, hem daha formel hem daha gündelik dil örnekleri toplanabilmiştir. Özellikle video tabanlı kaynaklarda, alan-özgü anahtar sözcüklerle aday içerikler bulunmuş; ardından yorumların Azerbaycan dili yoğunluğu ölçülerek yalnızca uygun içerikler tam kazıma sürecine alınmıştır. Kullanılan yoğunluk oranı aşağıdaki gibi tanımlanmıştır:

$$r_{az} = \frac{n_{az}}{n_{sample}} \quad (22)$$

Burada n_{sample} incelenen örnek yorum sayısını, n_{az} ise dil tanılaması ve ek sezgiler sonrasında Azerbaycan dili olarak kabul edilen yorum sayısını göstermektedir. Bu ölçüt, özellikle Azerbaycan dili ile Türkçenin karışabildiği sosyal medya ve video platformlarında alan temsiliyetini korurken dil saflığını artırmak için önemlidir. Bu kaynak seçimi mantığı, farklı alanlardan gelen verileri ortak bir temsil uzayında karşılaştırılabilir hâle getirmeyi amaçlamaktadır (Aykaç vd. 2026a; Aykaç vd. 2026).

Filtreleme aşamasında hızlı dil tanılama modeli ile Azerbaycan dili olmayan örnekler ayıklanmış; ancak Azerbaycan diline özgü harfler veya biçimbirimler taşıdığı halde otomatik olarak Türkçe olarak etiketlenen yorumlar ek sezgilerle geri kazanılmıştır. Çok kısa yorumlar elenmiş, yinelenen ya da neredeyse yinelenen içerikler karakter tabanlı benzerlik ölçüleriyle azaltılmış ve en az 10 token koşulu uygulanmıştır. Bu tercih, temsil öğrenmesi aşamasında cümle gömme uzayının çok kısa ve bağlamsız örnekler tarafından gürültülenmesini engellemektedir (Bojanowski vd. 2017; Aykaç vd. 2026).

3.7.3 Temel kodlayıcı ve cümle havuzlama

Alanlar arası sağlam temsil öğrenmesi hattında temel kodlayıcı olarak çok dilli transformer ailesi kullanılmış; deneylerin ana omurgasında XLM-R tercih edilmiştir. Kodlayıcı her yorum için alt-sözcük düzeyinde bağlamsal temsiller üretmekte, daha sonra bu token temsilleri yorumun tekil cümle gömmesine indirgenmektedir. Kullanılan havuzlama işlemi aşağıdaki genel gösterimle ifade

edilebilir:

$$\mathbf{s} = \text{Pool}(\tilde{\mathbf{h}}_1, \tilde{\mathbf{h}}_2, \dots, \tilde{\mathbf{h}}_T) \quad (23)$$

Burada $\tilde{\mathbf{h}}_t$, kutupluluk ve morfoloji bilgisi ile koşullandırılmış token temsilini; $\mathbf{s}$ ise duygu sınıflandırma başlığına ve probe deneylerine verilen son cümle gömmesini göstermektedir. Havuzlama katmanı pratikte ortalama havuzlama ya da ilk belirteç temsili ile uygulanabilmektedir (Reimers ve Gurevych 2019; Gao vd. 2021). Nihai duygu kararı ise aşağıdaki sınıflandırıcı üzerinden elde edilmektedir:

$$\hat{y} = \arg \max \text{softmax}(\mathbf{W}_c \mathbf{s} + \mathbf{b}_c) \quad (24)$$

Bu çerçevede asıl yenilik, sınıflandırıcıyı değiştirmekten çok $\mathbf{s}$ vektörünün geometrisini alan kaymasına daha dayanıklı hale getirmektir. Başka bir deyişle, sınıflandırma başlığı bu alt çalışmada nihai amaç değil; öğrenilen temsilin pratikte işe yarayıp yaramadığını gösteren ana denetim katmanıdır.

3.7.4 Alan-uyarlamalı ön eğitim ve alan-karşıtı katman

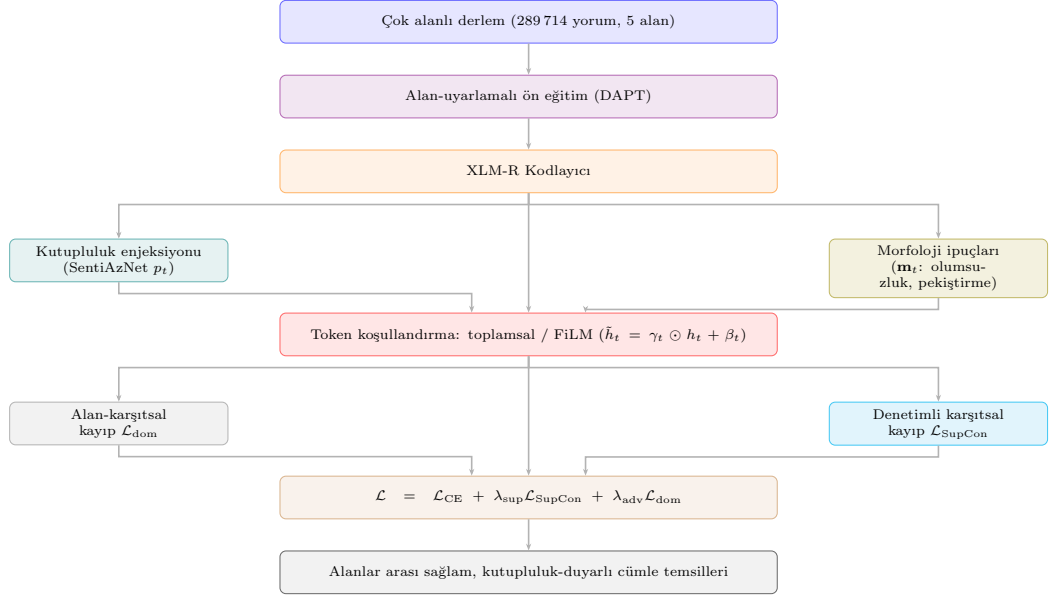
Kodlayıcının yalnızca genel çok dilli ön eğitim bilgisiyle sınırlı kalmaması için alan-uyarlamalı ön eğitim uygulanmıştır. Bu aşamada amaç, modelin Azerbaycan dili kullanıcı metinlerindeki alan-özgü sözcük dağılımına, yazım alışkanlıklarına ve yorum tarzına daha iyi uyum sağlamasıdır. Alan-uyarlamalı ön eğitim, yeni alan terimlerinin alt-sözcük parçalanması sonucu anlamsal olarak zayıf temsil edilmesi riskini azaltmakta ve sonraki denetimli eğitim için daha elverişli bir başlangıç noktası oluşturmaktadır (Gururangan vd. 2020; Conneau vd. 2020; Devlin vd. 2019).

Bu ön eğitimi tamamlayıcı ikinci bileşen, alan-karşıtı düzenlemedir. Buradaki amaç, duygu kararına yardımcı olmayan fakat modele alanı ezberleten yüzey ipuçlarını bastırmaktır. Kullanılan min-maks amaç aşağıdaki biçimde özetlenebilir:

$$\min_{\theta_f, \theta_y} \max_{\theta_d} \mathcal{L}_{\text{cls}}(\theta_f, \theta_y) - \lambda_{\text{adv}} \mathcal{L}_{\text{dom}}(\theta_f, \theta_d) \quad (25)$$

Burada θ_f özellik çıkarıcı kodlayıcıyı, θ_y duygu başlığını, θ_d ise alan ayırt etmeye çalışan domain sınıflandırıcısını göstermektedir. Gradyan tersleme katmanı üzerinden uygulanan bu düzenleme,

kodlayıcının duygu için yararlı ama alan için olabildiğince ayırt edici olmayan temsiller üretmesine yardımcı olur. Alan-uyarlamalı ön eğitim ile adversarial düzenleme birlikte düşünüldüğünde, biri sözcük dağılımı ve yazım kapsamını genişletmekte, diğeri ise alan bağımlı kestirmeleri zayıflatmaktadır (Ganin vd. 2016; Ramponi ve Plank 2020). Şekil 3.5 bu iş akışını özetlemektedir.



Şekil 3.5 Alanlar arası sağlam duygu temsilleri için izlenen genel iş akışı

3.7.5 Kutupluluk ve morfoloji bilgisi ile temsillerin koşullandırılması

Bu alt çalışmada SentiAzNet’ten gelen kutupluluk sinyalleri ile morfoloji-duyarlı yüzey işaretleri, token temsillerine doğrudan enjekte edilmektedir. Her token için kutupluluk ve morfoloji bileşenlerinden oluşan yardımcı özellik vektörü şu şekilde tanımlanmaktadır:

$$\mathbf{f}_t = [p_t; \mathbf{m}_t] \quad (26)$$

Burada $p_t \in [-1, 1]$ ilgili sözcük için SentiAzNet’ten gelen kutupluluk değerini, $\mathbf{m}_t$ ise olumsuzluk, yoğunlaştırma, yoksunluk, değilleme ve benzeri yüzey işaretlerini göstermektedir (Aykaç vd. 2025; Aykaç vd. 2026). Alt-sözcük parçalama kullanılan kodlayıcı yapısında bu sözcük-düzeyi sinyaller, ilgili alt-sözcük parçalarına yayılmaktadır. Böylece kutupluluk bilgisinin tokenizasyon sırasında kaybolması engellenmektedir.

İlk koşullandırma biçimi toplamsal enjeksiyondur:

$$\tilde{\mathbf{h}}_t = \mathbf{h}_t + \mathbf{W}\mathbf{f}_t \quad (27)$$

Bu form, kodlayıcı çıktısını korurken kutupluluk ve morfoloji bilgisini hafif ama doğrudan bir yan kanal olarak ekler. İkinci ve daha esnek koşullandırma biçimi ise FiLM tarzı ölçekleme-kaydırma mekanizmasıdır:

$$[\gamma_t; \beta_t] = \mathbf{W}_{\text{film}}\mathbf{f}_t + \mathbf{b}_{\text{film}} \quad (28)$$

$$\tilde{\mathbf{h}}_t = \gamma_t \odot \mathbf{h}_t + \beta_t \quad (29)$$

Burada $\odot$ eleman bazlı çarpımı göstermektedir. FiLM yaklaşımı, aynı kutupluluk sinyalinin farklı bağlamlarda aynı etkiyi yapmasını zorunlu kılmadığı için alanlar arası sağlamlık problemi açısından daha esnek bir seçenek sunmaktadır. Özellikle aynı sözcüğün teknoloji alanında olumsuz, kamu hizmetlerinde ise nesnel ya da nötr kullanılabildiği durumlarda, bağlamla birlikte ayarlanan ölçekleme-kaydırma parametreleri daha kararlı temsiller üretebilmektedir.

3.7.6 Denetimli karşıtsal düzenileştirme ve toplam amaç fonksiyonu

Temsil uzayının gerçekten alanlar arası sağlam hale gelmesi için yalnızca çapraz entropi kaybına güvenilmemiştir; ayrıca aynı sınıfa ait örnekleri alanlar arasında birbirine yaklaştıran denetimli karşıtsal düzenileştirme kullanılmıştır. Bu amaçla mini-batch içindeki her örnek için, aynı duygu etiketine sahip ve mümkün olduğunda farklı alanlardan gelen örnekler pozitif çiftler olarak değerlendirilmiştir. Kullanılan SupCon kaybı aşağıdaki biçimdedir:

$$\mathcal{L}_{\text{SupCon}} = \sum_{i=1}^N -\frac{1}{|\mathcal{P}(i)|} \sum_{p \in \mathcal{P}(i)} \log \frac{\exp(\text{sim}(\mathbf{z}_i, \mathbf{z}_p)/\tau)}{\sum_{a \in \mathcal{A}(i)} \exp(\text{sim}(\mathbf{z}_i, \mathbf{z}_a)/\tau)} \quad (30)$$

Burada $\mathbf{z}_i$ normalize edilmiş cümle gömmesini, $\mathcal{P}(i)$ örnek i için pozitif kümesini, $\mathcal{A}(i)$ ise kendisi dışındaki tüm adayları göstermektedir. Benzerlik işlemi olarak kosinüs benzerliği kullanılmakta, τ sıcaklık parametresi ise çekim-itim dengesini ayarlamaktadır (Khosla vd. 2020; Gao vd. 2021). Bu düzenleme sayesinde model, örneğin finans alanındaki nötr bir düzenleme haberini sosyal yaşam alanındaki nötr bir durum bildirimi ile aynı gömme bölgesine yerleştirmeyi öğrenebilmektedir.

Toplam amaç fonksiyonu, sınıflandırma, karşıtsal düzenleme ve alan-adversarial terimleri birleştiren aşağıdaki biçimde yazılabilir:

$$\mathcal{L}_{\text{toplaml}} = \mathcal{L}_{\text{CE}} + \lambda_{\text{sup}} \mathcal{L}_{\text{SupCon}} + \lambda_{\text{adv}} \mathcal{L}_{\text{dom}} \quad (31)$$

Burada λ_{sup} ve λ_{adv} katsayıları, temsili ne ölçüde karşıtsal olarak sıkılaştıracağımızı ve ne ölçüde alan bilgisinden arındıracağımızı denetlemektedir. Bu amaç fonksiyonu, *MahSA* bölümünde görülen nötr-sınıf kalibrasyonu ve sözlük katkısını temsil öğrenmesi düzeyinde daha genel hale getirmektedir.

3.7.7 Deneysel karşılaştırma aileleri ve değerlendirme protokolleri

Bu alt çalışmada yalnızca tek bir model değil, giderek güçlenen bir karşılaştırma ailesi kullanılmıştır. Yüzey temelli taban çizgiler, statik gömme tabanlı yapılar, çok dilli kodlayıcı ince ayarı, alan-uyarlamalı ön eğitim, alan-karşıtı düzenleme ve kutupluluk enjeksiyonu kademeli olarak karşılaştırılmıştır. Böylece elde edilen kazanımların yalnızca daha büyük bir kodlayıcı kullanmaktan mı, yoksa gerçekten temsil uzayının duyguya göre yeniden şekillenmesinden mi kaynaklandığı izlenebilmektedir (Devlin vd. 2019; Conneau vd. 2020; Bojanowski vd. 2017). Çizelge 3.16, bu bölümde karşılaştırılan temel model ailelerini göstermektedir.

Alanlar arası sağlam temsil öğrenmesi deneylerinde çok sayıda model ailesi karşılaştırılmış, ancak bu bölümde nicel başarı ölçütlerinin ayrıntılı sunumu yapılmamıştır. Bunun nedeni, başarı değerleri ve karşılaştırmalı tabloların izleyen bölümde bütün deneyler için ortak bir sunum çerçevesi altında verilmesidir. Bu aşamada yöntem açısından önemli olan nokta, farklı alanlarda aynı duygu sınıfına ait örnekleri birbirine yaklaştıran ve alan bağımlı kestirmeleri zayıflatan temsil tasarımının kurulmuş olmasıdır.

Bu çerçevede alanlar arası sağlam cümle temsilleri bölümü, tezin önceki yöntem katmanlarını temsil öğrenmesi düzeyinde birleştiren kapanış halkasıdır. SentiAzNet'ten gelen açık kutupluluk bilgisi ile *MahSA* bölümünde doğrulanan morfoloji-duyarlı tasarım burada ortak bir kuramsal soruya bağlanmaktadır: Azerbaycan dili için daha iyi duygu sınıflandırıcıları üretmek yeterli midir, yoksa duyguya göre düzenlenmiş ve alan kaymasına dayanıklı bir temsil uzayı mı inşa etmek gerekir? Bu bölümün yanıtı ikinci seçenektir ve alanlar arası sağlam temsil öğrenmesi çalışması

Çizelge 3.16 Alanlar arası sağlam temsil öğrenmesi bölümünde karşılaştırılan model aileleri

Varyant	Eklenen bileşen	Yöntemsel amacı
TF-IDF + lineer sınıflandırıcı	Yalnızca yüzeysel istatistikler	Alan kayması altında salt sözcük sıklığına dayalı taban çizgi sağlamak
fastText + lineer / MLP	Statik alt-sözcük gömmeleri	Karakter tabanlı genelleme ile bağlamdan bağımsız bir karşılaştırma sunmak
mBERT / XLM-R ince ayarı	Çok dilli bağlamsal kodlayıcı	Güçlü kodlayıcı tabanlı başlangıç noktası oluşturmak
+ DAPT/TAPT	Alan-uyarlamalı ön eğitim	Sözcük dağarcığı ve stil uyumunu artırmak
+ Adv. alan düzenlemesi	Alan-karşıtı eğitim	Alanı ezberleten ama duyguya taşınmayan ipuçlarını bastırmak
+ Kutupluluk / morfoloji enjeksiyonu	SentiAzNet + yüzey işaretleri	Duyguya özgü açık sembolik sinyalleri temsile taşımak
+ SupCon (tam model)	Karşıtsal düzenleme	Aynı duyguyu taşıyan farklı alan örneklerini gömme uzayında yaklaştırmak

tam olarak bu savı sınamak için konumlandırılmıştır (Aykaç vd. 2026a; Aykaç vd. 2026; Aykaç vd. 2025; Ramponi ve Plank 2020).

3.8 Açıklanabilirlik, Hata Analizi ve Bütünleşik Değerlendirme Protokolü

Yedi aşamalı metodolojinin yedinci ve son aşamasında, tezde geliştirilen yöntemler yalnızca nicel başarı değerleri üzerinden değil, kararların hangi dilsel işaretlerden beslendiğini görünür kılan açıklanabilirlik ve hata analizi düzenekleri üzerinden de değerlendirilmiştir. Finans ve iş dünyası alanındaki ABSA deneylerinde SHAP tabanlı özellik atıfları, dikkat örüntüleri ve sözlük eşleşmeleri birlikte incelenmiştir. *MahSA* deneylerinde ise sınıf bazlı hata çözümlemeleri, karışıklık matrisleri ve bileşen çıkarma deneyleri aracılığıyla hangi bilgi bloklarının özellikle nötr sınıf ve zor alanlarda katkı ürettiği analiz edilmiştir. Alanlar arası sağlam temsil öğrenmesi çalışmasında ise gömme uzayının yapısı, yakın komşu duygu tutarlılığı ve kümeleme uyumu gibi probe ölçümleriyle desteklenmiştir.

Bu yaklaşımın amacı, bir modelin yalnızca ne kadar başarılı olduğunu göstermek değil; başarımın hangi bilgi kanallarından beslendiğini ve hangi hata bölgelerinde sınırlı kaldığını sistematik biçimde ortaya koymaktır. Nicel değerlendirme ölçütleri ve sonuçların sunum çerçevesi izleyen bölümün

başında ayrıntılı olarak verilmektedir.

Son olarak, bu metodoloji tasarımında tekrarlanabilirlik temel bir ilke olarak benimsenmiştir. Veri bölmeleri, ön işleme kararları, özellik blokları, model aileleri, temel hiperparametreler ve karşılaştırma protokolleri ilgili alt başlıklarda açık biçimde tanımlanmış; nicel sonuçlar ise izleyen bölümde ortak bir değerlendirme çerçevesi altında raporlanmıştır. Böylece tez, yalnızca yeni bir yöntem ailesi önermekle kalmamakta, aynı zamanda bu yöntemin hangi varsayımlar altında yeniden kurulabileceğini de belgelemiş olmaktadır.

Bu bölümde Azerbaycan dili için önerilen duygu analizi çerçevesinin veri, sözlük, modelleme, temsil öğrenmesi ve açıklanabilirlik bileşenleri yöntemsel düzeyde sunulmuştur. Bir sonraki bölümde, burada tanımlanan her bir aşamaya ilişkin deneysel bulgular artan karmaşıklık düzeni içinde raporlanacak ve birlikte tartışılacaktır.

4. BULGULAR VE TARTIŞMA

Bu bölümde, tez kapsamında geliştirilen kaynak ve yöntemlere ilişkin deneysel bulgular artan karmaşıklık düzeni içinde sunulmakta ve birlikte değerlendirilmektedir (Aykaç vd. 2025, 2026b; Aykaç vd. 2026; Aykaç vd. 2026a). Sunum sırası, kaynak geliştirmeden alan-özü modellemeye, oradan çok alanlı hibrit yapıya ve alanlar arası sağlam temsil öğrenmesine uzanan bütüncül araştırma akışını yansıtacak biçimde kurulmuştur. Böylece bulgular, birbirinden bağımsız deney kümeleri olarak değil; aynı araştırma sorusunun farklı düzeylerde sınıandığı ardışık bir bulgu dizisi olarak okunabilmektedir.

Bu bölümde dört araştırma hattı birlikte ele alınmaktadır: ilk olarak *SentiAzNet* kutupluluk sözlüğünün iç ve dış geçerlik bulguları; ikinci olarak finans ve iş dünyası alanındaki aspekt-tabanlı duygu analizi deneyleri; üçüncü olarak çok alanlı *MahSA* hibrit mimarisinin sonuçları; son olarak da tez kapsamında tamamlanan ve yayıma hazırlanan alanlar arası sağlam cümle temsilleri çalışmasının bulguları. Her alt bölümde yalnızca en yüksek skorlar verilmemiş; bunun yanında hangi bilgi türlerinin hangi koşullarda yarar sağladığı ve hangi hata bölgelerinin sürmeye devam ettiği de tartışılmıştır.

4.1 Deneysel Gerçekleme Ortamı ve Bulguların Sunum Çerçevesi

Tez kapsamında geliştirilen yöntemler Python tabanlı bir deneysel ortamda gerçekleşmiştir. Derin öğrenme deneylerinde PyTorch ve HuggingFace tabanlı kütüphaneler; geleneksel makine öğrenmesi, TF-IDF ve lineer sınıflandırıcı bileşenlerinde ise scikit-learn ekosistemi kullanılmıştır. Çok dilli ve Azerbaycan dili odaklı kodlayıcı aileleri aynı deney protokolü altında karşılaştırılmış; böylece sözlük, karakter, morfoloji ve bağlamsal temsil bloklarının etkileşimi ortak bir gerçekleme zemini üzerinde incelenebilmiştir. Deneylerde tekrarlanabilirliği artırmak amacıyla veri bölmeleri, eğitim çevrimleri ve temel hiperparametreler mümkün olduğunca sabit tutulmuş; her alt çalışma için görev-özel ayrıntılar ilgili alt başlıklarda ayrıca verilmiştir.

Bu tezde kullanılan deneysel ölçütler alt göreve göre değişmektedir (Aykaç vd. 2025, 2026b; Aykaç vd. 2026; Aykaç vd. 2026a). Sözlük doğrulama katmanında doğruluk, Macro-F1, MCC ve ROC-AUC gibi ölçütler öne çıkarken; sınıflandırma deneylerinde temel karşılaştırma makro ortalamalı

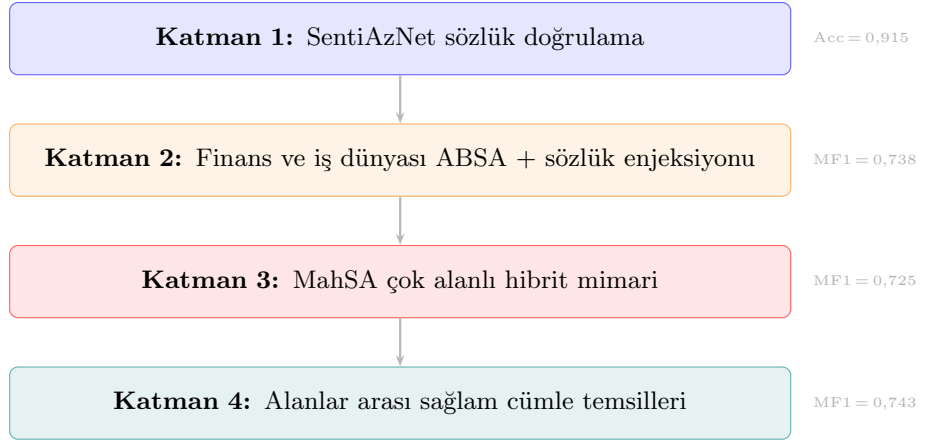
F1 üzerinden yapılmıştır. Bunun başlıca nedeni, hem finans ve iş dünyası ABSA verisinde hem de çok alanlı yorum derlemlerinde sınıf dağılımlarının tam dengeli olmaması ve özellikle nötr sınıfın daha güç bir sınıf olarak ortaya çıkmasıdır. Bu nedenle bölüm boyunca yorumlar, tek başına doğruluk değerine değil; daha çok sınıflar arası dengeli başarıyı yansıtan Macro-F1 çizgisine dayandırılmıştır.

Karşılaştırmalı kazançları daha görünür kılmak için göreceli iyileşme oranı Denklem 32 ile, alanlar arası sağlamlık deneylerinde kullanılan ortalama bırak-bir-alan-dışı-bırak başarımı ise Denklem 33 ile ifade edilmiştir. Bu iki gösterim, farklı veri büyüklükleri ve görev zorlukları arasında ortak bir yorumsal çerçeve korumaya yardımcı olmaktadır.

$$\Delta_{\text{rel}}(\%) = \frac{S_{\text{prop}} - S_{\text{base}}}{S_{\text{base}}} \times 100 \quad (32)$$

$$\overline{F1}_{\text{LODO}} = \frac{1}{|\mathcal{D}|} \sum_{d \in \mathcal{D}} F1_d \quad (33)$$

Genel bulgu akışı Şekil 4.1’de özetlenmiştir. Şekil, tezin deneysel anlatısını dört katman hâlinde göstermektedir: önce sözlük kalitesi değerlendirilmektedir; ardından bu sözlüğün alan-özgü bir sinirsel modele katkısı sınanmakta; sonrasında çok alanlı hibrit yapı üzerinden genel başarı örüntüsü ele alınmakta; en sonunda ise temsil uzayının alan kaymasına karşı nasıl daha dayanıklı hâle getirilebildiği incelenmektedir.



Şekil 4.1 Tezin deneysel kanıt zincirinin özet görünümü

4.2 SentiAzNet Kutupluluk Sözlüğüne İlişkin Bulgular

SentiAzNet hattının bulguları iki düzeyde ele alınmıştır (Aykaç vd. 2025). İlk düzey, sözlüğün kontrollü iç doğrulama başarımını; ikinci düzey ise bu kaynağın gerçek kullanıcı yorumları üzerinde ne ölçüde taşınabilir kaldığını gösteren dış geçerlik değerlendirmesini içermektedir. Bu iki katman birlikte değerlendirildiğinde, *SentiAzNet*'in yalnızca sözlük oluşturma sürecinin bir çıktısı olarak değil, sonraki modelleme aşamalarına aktarılabilen yararlı bir sembolik kaynak olarak da görülebileceği anlaşılmaktadır.

İç doğrulama bulguları Çizelge 4.1'de verilmiştir. Negatif sınıfta gözlenen daha yüksek F1 değeri, veri dengesizliği ve olumsuz kutupların daha belirgin leksik işaretler taşımasıyla ilişkilendirilebilir. Buna karşılık pozitif sınıfta hataların görece artması, bazı olumlu sözcüklerin bağlama daha duyarlı olmasına ve özellikle kültürel kullanımlardaki dolaylı olumlu ifadelerin daha değişken görünmesine bağlanabilir. Buna rağmen genel doğruluk 0.915, Macro-F1 ise 0.892 düzeyindedir. ROC-AUC değerinin 0.98 olması da, sözlüğe dayalı özellik kümesinin ikili kutupluluk ayrımında güçlü bir ayırt edicilik sunduğuna işaret etmektedir.

Çizelge 4.1 SentiAzNet iç doğrulama sonuçları

Sınıf / ölçüt	Precision	Recall	F1	Örnek sayısı
Negatif	0.974	0.912	0.942	486
Pozitif	0.772	0.924	0.841	158
Genel ölçütler: Accuracy = 0.915, Macro-F1 = 0.892, MCC = 0.79, ROC-AUC = 0.98				

Kararlılık ve dış geçerlik sonuçları Çizelge 4.2’de sunulmuştur. Beş katlı çapraz doğrulamada Macro-F1 değerinin 0.89 ± 0.01 düzeyinde kalması, modelin farklı bölünmeler altında istikrarlı davrandığına işaret etmektedir. Buna karşılık gerçek yorumlardan oluşan dış veri üzerinde doğruluk 0.6324 düzeyine inmektedir. Bu düşüş şaşırtıcı değildir; çünkü sözlük doğrulama deneyi lemma-merkezli ve daha kontrollü bir ortamda yürütülürken, dış veri alan çeşitliliği, yazım gürültüsü, dolaylı anlam ve çok anlamlılık gibi etkenleri aynı anda içermektedir. Dolayısıyla dış geçerlik bulgusu, *SentiAzNet*’in yetersiz olduğunu değil; bu kaynağın sonraki üç sınıflı ve bağlam-duyarlı modellerde yardımcı sembolik bilgi olarak kullanılmasının daha uygun olduğunu düşündürmektedir.

Çizelge 4.2 SentiAzNet kararlılık ve dış geçerlik bulguları

Ölçüt	Değer
5-katlı çapraz doğrulama Macro-F1	0.89 ± 0.01
Dış veri doğruluğu	0.6324
Dış veri pozitif precision	0.77
Dış veri negatif recall	0.78
Dış veri macro-ortalama F1	0.63

Bu alt çalışma üç temel noktaya işaret etmektedir. İlk olarak, Azerbaycan dili için elle denetlenmiş ve diğer dillere ait zengin kaynaklardan beslenen güvenilir bir kutupluluk sözlüğü üretilebildiği görülmektedir. İkinci olarak, sözlük tabanlı sinyal tek başına belirli bir ayırt edicilik sunsa da, gerçek kullanıcı yorumlarında bağlam bilgisinin eksikliği nedeniyle tek başına nihai çözüm niteliği taşımamaktadır. Üçüncü olarak, tam da bu nedenle *SentiAzNet* tez boyunca *nihai sınıflandırıcı* olmaktan çok, *yardımcı sembolik ön bilgi* işlevi üstlenmektedir. Finans ve iş dünyası ABSA ile *MahSA* katmanlarında gözlenen kazanımlar da bu yorumla uyumludur.

Bu örüntü, sözlük tabanlı duygu analizi literatüründeki daha genel gözlemlerle de uyumludur. Kutupluluk sözlükleri çoğu durumda yüksek kesinlikli ve yorumlanabilir bir başlangıç sağlar; ancak çok anlamlılık, alan bağımlılığı ve örtük duygu gibi bağlama gömülü olgular karşısında tek başına sınırlı kalır (Taboada vd. 2011; Shaik vd. 2023; Mutinda vd. 2023). Dolayısıyla bu alt bölümden çıkan ana sonuç, Azerbaycan dili için bir sözlük geliştirmenin yeterli olmadığı; fakat iyi tasarlanmış bir sözlüğün sonraki sinirsel ve hibrit modeller için değerli bir sembolik dayanak sağlayabildiğidir.

4.3 Finans ve İş Dünyası ABSA Bulguları ve Tartışma

İkinci deneysel katman, *SentiAzNet* bilgisinin finans ve iş dünyası alanında bağlamsal bir modele nasıl aktarılabilceğini incelemektedir (Aykaç vd. 2026b). Bu alt problemde amaç, tüm cümleye tek bir duygu etiketi vermek değil; verilen bir aspekt için duygu kutbunu tahmin etmektir. Sözlük kapsamı seyrek olsa da, kutupluluk ipuçlarının uygun biçimde kullanılması performansa yararlı bir katkı sunabilmektedir. Özellikle dikkat önyargısı yaklaşımı, sözlük sinyalinin bağlamsal kodlayıcı içinde daha etkili değerlendirilmesine imkân tanımaktadır.

Ana sonuçlar Çizelge 4.3’da özetlenmiştir. Yalın XLM-R modeli 0.702 Macro-F1 düzeyindeyken, sözlük özellik birleştirme yaklaşımı bu değeri 0.718’e, dikkat önyargısı yaklaşımı ise 0.727’ye taşımaktadır. Finans alanına özgü küçük ölçekli kutup düzeltmeleri içeren *SentiAzNet-Fin* ile birlikte dikkat önyargısı kullanan P3 modeli 0.769 doğruluk ve 0.738 Macro-F1 değerine ulaşmaktadır. Denklem 32 kullanıldığında, P3 modelinin yalın XLM-R’a göre görece Macro-F1 kazancı yaklaşık %5.13 olmaktadır. Bu sonuç, sözlük sinyalinin kapsamı sınırlı olsa da kararın kritik noktalarında yararlı bir yönlendirici işlev görebildiğini düşündürmektedir.

Çizelge 4.3 Finans ve iş dünyası ABSA ana sonuçları

Model	Accuracy	Macro-P	Macro-R	Macro-F1	W-F1
B1: Rule-based (SentiAzNet only)	0.412	0.358	0.341	0.349	0.398
B2: SVM + TF-IDF	0.621	0.583	0.571	0.577	0.618
B3: BiLSTM + FastText	0.658	0.612	0.604	0.608	0.652
B4: mBERT (fine-tuned)	0.714	0.681	0.673	0.677	0.710
B5: XLM-R (vanilla)	0.738	0.706	0.698	0.702	0.734
P1: XLM-R + SentiAzNet (Feature)	0.751	0.722	0.714	0.718	0.748
P2: XLM-R + SentiAzNet (Attention Bias)	0.759	0.731	0.724	0.727	0.756
P3: XLM-R + SentiAzNet-Fin (Attention Bias)	0.769	0.742	0.735	0.738	0.770

Bileşen katkıları ve alternatif eğitim düzenleri Çizelge 4.4’da yer almaktadır. Burada özellikle “yalnız önyargı” varyantının sınırlı bir artış üretmesi dikkat çekicidir; çünkü bu durum, mimarideki değişikliğin tek başına değil, sözlük bilgisinin kendisinin esas katkırı taşıdığını düşündürmektedir. Odak kaybı (focal loss) ve ağırlıklı çapraz entropi gibi sınıf dengesizliği duyarlı kayıplar da güçlü sonuçlar vermektedir; ancak en belirgin iyileşme yine alan-uyarlanmış sözlük ile dikkat önyargısının birlikte kullanıldığı yapıdan elde edilmektedir.

Çizelge 4.4 Finans ve iş dünyası ABSA ablasyon ve alternatif eğitim sonuçları

Yapılandırma	Macro-F1	$\Delta F1$ (puan)	Accuracy
XLM-R (base)	0.702	–	0.738
+ SentiAzNet (Feature concat)	0.718	+1.6	0.751
+ SentiAzNet (Dikkat önyargısı)	0.727	+2.5	0.759
+ Domain adaptation (SentiAzNet-Fin)	0.738	+3.6	0.769
– SentiAzNet (yalnız önyargı)	0.708	+0.6	0.742
+ SentiAzNet + Focal Loss	0.735	+3.3	0.762
+ SentiAzNet + Weighted CE Loss	0.733	+3.1	0.760

Hata dağılımı incelendiğinde, nötr sınıfın en güç bölgelerden biri olduğu görülmektedir. En iyi modelde dahi satır-normalize karışıklık matrisinde nötr örneklerin yalnızca %61.3’ü doğru sınıflandırılmakta, %25.0’ı negatif sınıfa kaymaktadır. Bu örüntü, finans ve iş dünyası metinlerinde olumsuz çağrışım taşıyan ancak açık bir şikâyet de içermeyen kısa ifadelerin yaygınlığıyla uyumlu görünmektedir. Çalışmada incelenen hata türleri Çizelge 4.5’de gösterilmiştir. Sarkazm ve ironi,

yüzeyde olumlu sözcükler içerdiği için hem sözlük sinyalini hem de yüzey istatistiklerini yanıltabilmektedir. Nötr-negatif belirsizliği ise alanın doğası gereği daha yapısal bir zorluk olarak öne çıkmaktadır.

Çizelge 4.5 Finans ve iş dünyası ABSA’da gözlenen başlıca hata tipleri

Hata türü	Pay
Sarkazm / ironi	%28.3
Nötr ↔ negatif belirsizliği	%24.1
Karma duygu	%18.7
Örtük duygu	%14.2
Kod değiştirimi	%8.9
Alan jargonu	%5.8

Bu alt bölümün tez açısından temel önemi, *SentiAzNet* bilgisinin bağlamsal modele doğrudan ve sert biçimde değil, seçici ve yumuşak bir sinyal olarak verilmesinin daha uygun görüldüğünü ortaya koymasıdır. Dikkat önyargısı yaklaşımının özellik birleştirme yaklaşımına göre daha iyi sonuç vermesi, sözlük bilgisinin sabit bir etiket gibi değil, dikkat dağılımını yönlendiren yardımcı bir ipucu olarak kullanıldığında daha yararlı olabildiğini düşündürmektedir. Bununla birlikte sarkazm, nötr-negatif sınır vakaları ve kod değiştirimi bu alt problem açısından açık güçlük alanları olarak varlığını sürdürmektedir.

Bu sonuçlar, finansal ve hedefe bağlı duygu analizi literatürüyle iki açıdan ilişkilendirilebilir. İlk olarak, finansal metinlerde genel dil kutupluluğunun her zaman geçerli olmadığı; kimi teknik terimlerin alan bağlamında yeniden yorumlanması gerektiği uzun süredir bilinmektedir (Loughran ve McDonald 2011; Araci 2019). *SentiAzNet-Fin* ile elde edilen ek artış, bu gözlemin Azerbaycan dili için de geçerli olduğunu göstermektedir. İkinci olarak, güncel ABSA derlemeleri hedef-koşullu bağlamsal kodlayıcıların önemini vurgulasa da, düşük kaynaklı dillerde hafif sözlük enjeksiyonunun sağlayabileceği yarar daha sınırlı incelenmiştir (Hua vd. 2024; Šmíd ve Král 2025). Buradaki deneyler, seyrek kapsamlı bir sözlüğün bile doğru yerde kullanıldığında anlamlı iyileşme üretebildiğini göstermesi bakımından dikkat çekicidir.

Öte yandan, iyileşmenin mutlak olarak sınırlı kalması da önemlidir. Bu durum, sözlük bilgisinin

bağlamın yerini alamadığını; ancak kararın kritik düğüm noktalarında bağlamsal modele tamamlayıcı destek verdiğini göstermektedir. Başka bir ifadeyle ABSA katmanı, tezde savunulan “sembolik ve sinirsel bilgi birlikteliği” fikrinin ilk kontrollü doğrulama alanı olarak değerlendirilebilir.

4.4 MahSA Çok Alanlı Hibrit Mimarisine İlişkin Bulgular

Tezin ana deneysel omurgasını *MahSA* oluşturmaktadır (Aykaç vd. 2026). Bu alt çalışmada problem artık yalnızca tek alanlı ya da ikili bir senaryo değildir; beş farklı alana dağılmış, üç sınıflı ve nötr sınıfı belirgin biçimde içeren daha kapsamlı bir duygu analizi düzeni söz konusudur. Bu bağlamda *MahSA*’nın ayırt edici yönü, *SentiAzNet* sinyalinin karakter n-gramları, morfoloji-duyarlı yüzey ipuçları, wPPMI ile bağlama yayılan kutupluluk bilgisi ve bağlamsal transformer temsilleriyle bir araya getirmesidir.

Alan bazlı ana sonuçlar Çizelge 4.6’de verilmiştir. *MahSA*, tüm alanların ortalamasında 0.725 Macro-F1 ile XLM-R ince ayar modelinin 0.688 ve Az-RoBERTa ince ayar modelinin 0.697 düzeylerinin üzerine çıkmaktadır. En belirgin mutlak artış kamu hizmetleri alanında gözlenmektedir; burada XLM-R ince ayar modeli 0.653’te kalırken *MahSA* 0.692’ye ulaşmaktadır. Bu örüntü, küçük ve nötr-ağırlıklı alanlarda sembolik önbilginin ve yüzey ipuçlarının daha fazla değer taşıyabildiğini düşündürmektedir.

Çizelge 4.6 MahSA ve temel karşılaştırma modelleri için alan bazlı Macro-F1 sonuçları

Model	Tek./Dij.	Finans	Sosyal	Perakende	Kamu	Tüm alanlar
SentiAzNet-only	0.592	0.601	0.548	0.612	0.532	0.577
Char TF-IDF	0.648	0.652	0.612	0.668	0.595	0.635
CLS embeddings	0.662	0.668	0.628	0.682	0.615	0.651
mBERT ince ayar	0.678	0.685	0.652	0.695	0.630	0.668
Az-RoBERTa ince ayar	0.708	0.715	0.678	0.722	0.662	0.697
XLM-R ince ayar	0.698	0.705	0.672	0.712	0.653	0.688
MahSA (XLM-R)	0.735	0.742	0.708	0.748	0.692	0.725

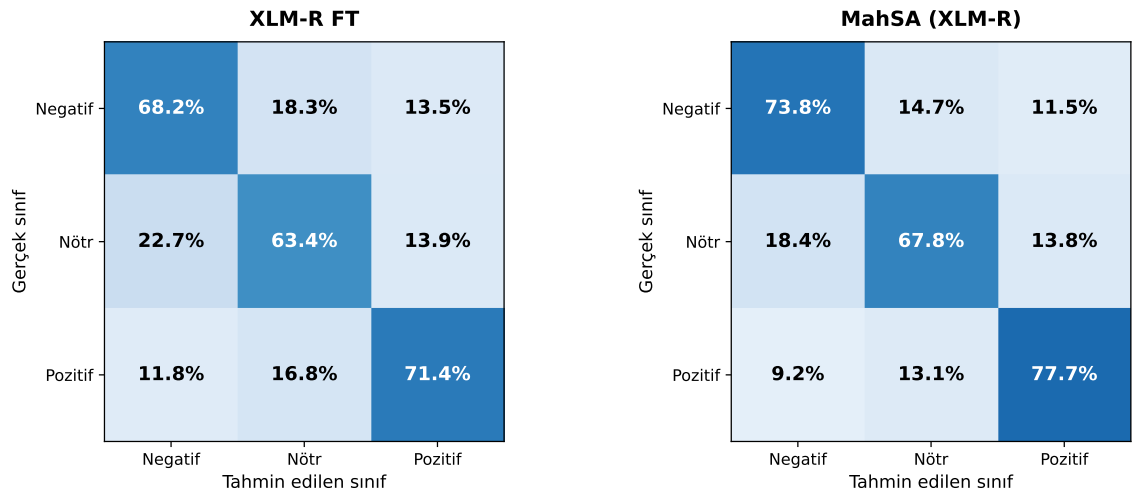
MahSA’nın güçlü yönünü daha iyi görebilmek için seçilmiş güçlü modellerin sınıf bazlı ve genel

sonuçları Çizelge 4.7’de özetlenmiştir. Burada *MahSA* negatif sınıfta 0.748, nötr sınıfta 0.678 ve pozitif sınıfta 0.749 F1 elde etmektedir. Özellikle nötr sınıfta XLM-R ince ayar modeline göre +4.4 puanlık artış dikkat çekicidir. Bu durum, hibrit yapının yalnızca kutuplu sınıfları daha iyi yakalamasıyla değil, aynı zamanda karar sınırını nötr örnekler için daha dengeli kurabilmesiyle de ilişkilendirilebilir.

Çizelge 4.7 Seçilmiş güçlü modeller için sınıf bazlı ve genel sonuçlar

Model	Negatif F1	Nötr F1	Pozitif F1	Genel Macro-F1
Az-RoBERTa ince ayar	0.724	0.642	0.725	0.697
XLM-R ince ayar	0.718	0.634	0.712	0.688
Qwen 2.5 + QLoRA	0.742	0.668	0.744	0.718
MahSA (XLM-R)	0.748	0.678	0.749	0.725

Şekil 4.2 satır-normalize karışıklık matrislerini göstermektedir. XLM-R ince ayar modelinde gerçek nötr örneklerin %22.7’si negatif sınıfa kayarken, *MahSA*’da bu oran %18.4’e düşmektedir. Benzer biçimde gerçek negatif örneklerin nötr olarak tahmin edilme oranı %18.3’ten %14.7’ye inmektedir. Dolayısıyla *MahSA*’nın en belirgin kazanımı, doğrudan negatif-pozitif karışıklığını azaltmasından çok; *nötr ile kutuplu sınıflar arasındaki sınır bölgesini* daha dengeli yönetebilmesinden kaynaklanmaktadır.



Şekil 4.2 XLM-R ince ayar ve MahSA için tez kapsamında elde edilen satır-normalize karışıklık matrisleri

MahSA’nın bileşen katkıları Çizelge 4.8’de sunulmuştur. En belirgin düşüşün sözlük ve wPPMI

bileşenleri çıkarıldığında ortaya çıkması, sembolik bilginin model için tamamlayıcı bir rol üstlendiğine işaret etmektedir. Karakter n-gramlarının çıkarılması da anlamlı bir kayba yol açmaktadır; bu bulgu, gürültülü kullanıcı dilinde alt-sözcük ve yazım varyasyonlarının kritik önem taşıdığını düşündürmektedir. Buna karşılık yalnızca Lex+Morph yapısının 0.637 Macro-F1 ile belirgin biçimde daha düşük kalması, sembolik bilgi ile bağlamsal temsillerin birbirinin alternatifi olmaktan çok, birbirini tamamlayan kanallar olarak değerlendirilmesi gerektiğini göstermektedir.

Çizelge 4.8 MahSA için ablasyon sonuçları

Model varyantı	Macro-F1	Nötr F1
MahSA (full)	0.725	0.678
$L(x)$ ve $\psi(x)$ olmadan	0.703	0.643
$M(x)$ olmadan	0.718	0.665
$C(x)$ olmadan	0.712	0.658
Lex+Morph only	0.637	0.598
MahSA + ordinal loss	0.722	0.672
MahSA + hierarchical loss	0.719	0.668

Alan kayması altındaki davranış Çizelge 4.9’da verilmektedir. Her hedef alanda *MahSA*, XLM-R ince ayar modelinin üzerinde sonuç üretmektedir. Kamu alanında artış 0.605’ten 0.652’ye, yani +4.7 puana ulaşmaktadır. Bu bulgu, *MahSA*’nın yalnızca alan içi uyum sağlayan bir model olmadığını; alan değişiminde de sözlük, karakter ve morfoloji bilgisi sayesinde daha dayanıklı kalabildiğini düşündürmektedir.

Çizelge 4.9 MahSA için LODO (bir alanı dışarıda bırakma) sonuçları

Model	Tek./Dij.	Finans	Sosyal	Perakende	Kamu
XLM-R ince ayar	0.658	0.668	0.642	0.672	0.605
MahSA	0.695	0.702	0.678	0.708	0.652

Toplam olarak *MahSA* alt çalışması üç ana noktaya işaret etmektedir. İlk olarak, sözlük destekli hibrit tasarım çok alanlı ve üç sınıflı senaryoda etkili bir ana çözüm sunmaktadır. İkinci olarak, bu katkı özellikle nötr sınıfta ve küçük hedef alanlarda daha görünür hâle gelmektedir. Üçüncü olarak, alan kayması altında dahi sembolik ve yüzeysel ipuçları, salt bağlamsal kodlayıcıların üzerinde istikrarlı bir ek yarar sağlayabilmektedir.

Bu örüntü, morfolojik açıdan zengin dillerde alt-sözcük tabanlı çok dilli modellerin tek başına her zaman yeterli olmadığını vurgulayan çalışmalarla uyumludur (Rust vd. 2021; Dehkharghani vd. 2017; Erşahin vd. 2019; Altınel Girgin vd. 2024). Özellikle $L(x)$ ve $\psi(x)$ bloklarının çıkarıldığında hem genel Macro-F1’in hem de nötr sınıf F1’inin belirgin biçimde düşmesi, kutupluluk bilgisinin yalnızca örtük bağlamsal temsiller içinde bırakılmaması gerektiğini; açık sembolik ve dağıtımsal sinyallerin hâlâ ayrı bir katkı taşıdığını göstermektedir. Karakter n-gramlarının çıkarılmasıyla gelen düşüş de, gürültülü kullanıcı metinlerinde yazım varyasyonu ve alt-sözcük parçalanmasının yöntemsel olarak merkezi olduğunu doğrulamaktadır.

Bir başka dikkat çekici nokta, büyük parametrelili bir ince ayar düzeni olan Qwen 2.5 + QLoRA’nın güçlü bir taban çizgisi oluşturmaya rağmen *MahSA*’nın üzerinde kalamamasıdır. Bu durum, belirli dil ve görev koşullarında daha büyük modelin tek başına otomatik üstünlük sağlamadığını; dile özgü sembolik ve morfoloji-duyarlı bilginin dikkatli füzyonunun daha dengeli bir sonuç üretebildiğini düşündürmektedir (Bai vd. 2023; Dettmers vd. 2023). Bu yönüyle *MahSA*, “daha büyük model” yaklaşımı ile “daha uygun bilgi bileşimi” yaklaşımı arasındaki farkı deneysel olarak görünür kılmaktadır.

4.5 Alanlar Arası Sağlam Cümle Temsillerine İlişkin Bulgular

Bu alt bölüm, tez kapsamında tamamlanan ve yayıma hazırlanan temsil öğrenmesi çalışmasının sonuçlarını içermektedir (Aykaç vd. 2026a). Buradaki amaç, önceki alt çalışmalarda doğrulanan sözlük ve morfoloji bilgisinin yalnızca son katman sınıflandırıcısına değil, doğrudan cümle temsillerinin geometrisine de katkı sağlayıp sağlayamayacağını sınamaktır. Bu nedenle değerlendirme yalnızca sınıflandırma başarımı ile sınırlı tutulmamış; en yakın komşu duygu tutarlılığı ve kümeleme uyumu gibi ek göstergeler de raporlanmıştır.

Alan bazlı alan-içi sonuçlar Çizelge 4.10’de sunulmuştur. Yalın XLM-R ince ayar yaklaşımı ortalamada 69.8 Macro-F1 düzeyindeyken, DAPT, kutupluluk enjeksiyonu, morfoloji ipuçları ve denetimli karşıtsal düzenleştirme içeren tam model bu değeri 74.3’e taşımaktadır. En düşük başlangıç performansına sahip kamu alanında 66.5’ten 71.3’e yükseliş görülmesi özellikle dikkat çekicidir. Bu artış, temsil öğrenmesinin yalnızca yüksek kaynaklı ya da görece kolay alanlarda değil, daha kırılğan alanlarda da yarar üretebildiğini göstermektedir.

Çizelge 4.10 Alanlar arası sağlam temsil öğrenmesi çalışmasında alan-içi Macro-F1 sonuçları

Model	Tek./Dij.	Finans	Sosyal	Perakende	Kamu
TF-IDF + Linear	60.2	57.4	54.1	56.8	53.5
fastText + Linear	62.8	60.1	57.3	59.4	55.9
mBERT (ft)	70.5	67.2	65.4	67.8	63.1
XLM-R (ft)	73.4	70.8	68.2	70.1	66.5
Önerilen: +DAPT	74.9	72.1	70.5	72.3	68.4
Önerilen: +Enjeksiyon	75.8	73.4	71.1	73.0	69.2
Önerilen: Tam model	77.5	75.2	72.8	74.6	71.3

Alan kayması altındaki LODO sonuçları Çizelge 4.11’da yer almaktadır. Burada iyileşmenin daha da belirgin hâle geldiği görülmektedir. Ortalama Macro-F1 değeri 59.7’den 68.6’ya yükselmekte; yani mutlak kazanç 8.9 puana ulaşmaktadır. Bu fark, temsil uzayına kutupluluk ve morfoloji bilgisinin özellikle alan değişimi durumunda daha fazla yardımcı olduğunu düşündürmektedir. Başka bir deyişle, önerilen yöntem yalnızca sınıflandırıcıyı güçlendirmemekte; farklı alanlarda aynı duyguyu taşıyan örnekleri birbirine daha yakın yerleştiren daha düzenli bir gömme uzayı da üretmektedir.

Çizelge 4.11 Alanlar arası sağlam temsil öğrenmesi çalışmasında LODO sonuçları

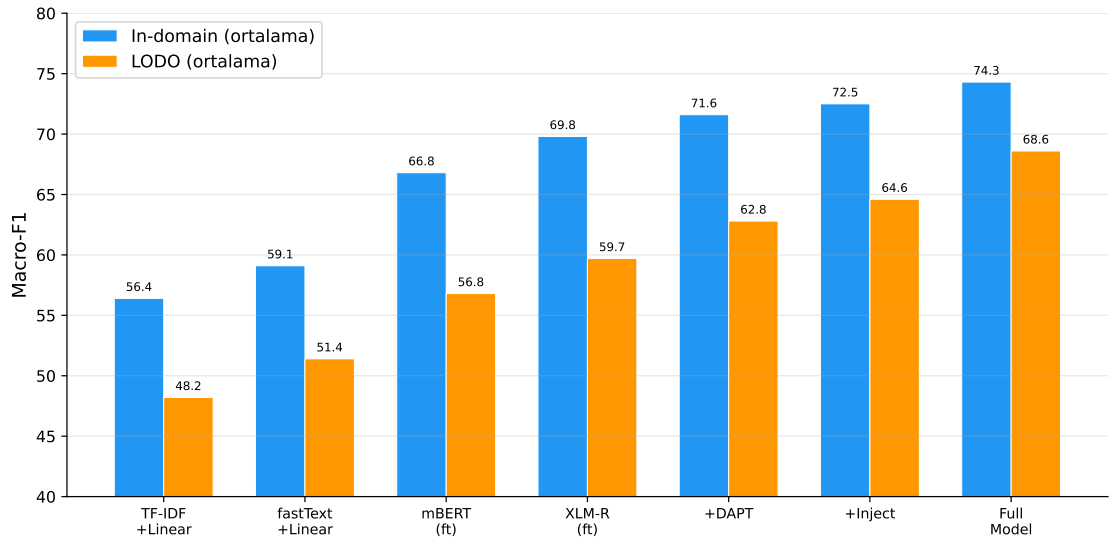
Eğitim: 4 alan → test: 1 alan	Tek./Dij.	Finans	Sosyal	Perakende	Kamu
XLM-R (ft)	65.1	61.3	57.4	59.2	55.6
XLM-R + DAPT	68.4	64.5	60.1	62.3	58.9
XLM-R + Adv (DANN)	69.2	65.1	61.5	63.4	59.8
Önerilen: Enjeksiyon	70.1	66.2	62.3	64.1	60.5
Önerilen: DAPT + Enjeksiyon	71.8	68.4	64.2	66.5	62.1
Önerilen: Tam model (+SupCon)	73.5	70.1	66.8	68.2	64.4

Temsil uzayının yapısına ilişkin bulgular Çizelge 4.12’da özetlenmiştir. Burada yalnızca Macro-F1 değil, NN-consistency@10 ve NMI gibi gömme-probe ölçütleri de raporlanmaktadır. XLM-R ince ayar tabanından tam modele geçildiğinde NN-consistency@10 değeri 0.58’den 0.72’ye, NMI değeri ise 0.32’den 0.46’ya yükselmektedir. Bu örüntü, performans artışının yalnızca daha iyi bir son sınıflandırma sınırından kaynaklanmadığını; daha düzenli yapılandırılmış bir duygu temsili uzayının da öğrenildiğini düşündürmektedir.

Çizelge 4.12 Alanlar arası sağlam temsil öğrenmesi çalışmasında ablasyon ve gömme-probe sonuçları

Varyant	Macro-F1	NN-consistency@10	NMI
XLM-R ince ayar	69.8	0.58	0.32
+ DAPT	71.6	0.61	0.35
+ Polarity injection	72.5	0.65	0.39
+ Morphology cues	73.1	0.67	0.41
+ SupCon	74.3	0.72	0.46

Şekil 4.3 ortalama alan-içi ve LODO başarılarını birlikte göstermektedir. Şekilden de izlendiği üzere, mutlak artış alan kayması altında daha belirgindir. Bu örüntü, tezde savunulan temel görüşle uyumludur: Azerbaycan dili gibi morfolojik açıdan zengin ve alanlar arasında hızlı dağılım değişimi gösteren bir dilde, kutupluluk ve morfoloji bilgisini doğrudan temsil öğrenmesine dâhil etmek, özellikle hedef alan eğitim sırasında görülmediğinde daha değerli hâle gelmektedir.



Şekil 4.3 Alanlar arası sağlam temsil öğrenmesi çalışmasında model ailesine göre alan-içi ve LODO ortalama Macro-F1 sonuçları

Bu alt çalışma, tez açısından iki temel sonuca işaret etmektedir. İlk olarak, alan kayması altında daha iyi genelleme için yalnızca daha büyük modeller değil, daha iyi yapılandırılmış temsil uzayları da gerekmektedir. İkinci olarak, *SentiAzNet* ve morfoloji-duyarlı ipuçları yalnızca klasik hibrit modellerde değil, temsil öğrenmesi düzeyinde de yarar üretebilmektedir. Böylece tezin son katmanı, “daha iyi sınıflandırıcı” fikrini “daha sağlam duygu temsili” fikrine doğru genişletmektedir.

Bu bulgular, alan uyarlamalı ön eğitim ve temsil dayanıklılığı literatürüyle de anlamlı biçimde örtüşmektedir. DAPT türü ön eğitim düzenlerinin hedef alan performansını artırabildiği daha önce gösterilmiştir; ancak burada gözlenen ek kazanım, kutupluluk enjeksiyonu ve morfoloji ipuçlarıyla birleştirildiğinde daha belirgin hale gelmektedir (Gururangan vd. 2020; Fang vd. 2024; Nwaiwu 2025). Aynı şekilde, yalnızca karşıtsal düzenlileştirmenin değil; duyguya duyarlı koşullandırılmış temsil uzayının etkili olması, alanlar arası sağlamlığın salt “alanı unutturma” problemi olmadığını; hangi duygusal işaretlerin korunacağı sorusunu da içerdığını göstermektedir.

Bu nedenle temsil öğrenmesi katmanından çıkan esas sonuç, alan kayması probleminin yalnızca son karar yüzeyinde değil, gömme uzayının iç geometrisinde de ele alınması gerektiğidir. NN-consistency@10 ve NMI değerlerindeki artışlar, performans kazancının yalnızca son sınıflandırma katmanından gelmediğini; daha tutarlı ve duyguya duyarlı bir komşuluk yapısının öğrenildiğini göstermesi bakımından ayrıca önem taşımaktadır.

4.6 Bütünleşik Tartışma

Dört deneysel katman birlikte değerlendirildiğinde, tezin katkı yapısının daha net görülebildiği söylenebilir. İlk katmanda, Azerbaycan dili için güvenilir bir kutupluluk sözlüğü üretilbildiği görülmektedir. İkinci katmanda, bu sözlüğün bağlamsal modellere yumuşak biçimde enjekte edildiğinde alan-özü bir görevde kayda değer bir katkı sağlayabildiği anlaşılmaktadır. Üçüncü katmanda, sözlük sinyali karakter n-gramları, morfoloji ipuçları ve bağlamsal gömmelerle birleştğinde çok alanlı bir problemde daha güçlü bir yapı ortaya çıkmaktadır. Son katmanda ise aynı bilgi, alanlar arası daha sağlam cümle temsilleri öğrenmek için kullanılmaktadır. Başka bir deyişle, bulgular tekil ve bağlamdan kopuk başarı artışlarından çok, aynı temel fikrin farklı deney düzeylerinde sınandığı tutarlı bir örüntü sunmaktadır. Bu tutarlılık, tezin dört ayrı makalenin toplamından ibaret olmadığını; kaynak geliştirme, modelleme ve temsil öğrenmesini birbirine bağlayan tek bir araştırma programı ortaya koyduğunu göstermektedir.

Bu bütünleşik bulgu seti, üç temel savı destekler niteliktedir (Aykaç vd. 2025, 2026b; Aykaç vd. 2026; Aykaç vd. 2026a). İlk olarak, Azerbaycan dili için duygu analizinde yüksek nitelikli sembolik kaynakların hâlâ önemli olduğu görülmektedir. İkinci olarak, bu kaynaklar en yüksek yararı salt sözlük-tabanlı kullanımda değil, sinirsel ve hibrit modellerin içinde yardımcı sinyal olarak

sunmaktadır. Üçüncü olarak, morfoloji-duyarlı ve alan uyarlamalı temsil öğrenmesi, özellikle nötr sınıf ve küçük ya da hedef alanlar söz konusu olduğunda belirgin katkılar sağlayabilmektedir. Bu bulgular, tez başlığında belirtilen “Azerbaycan dili için bir duygu analizi çerçevesi geliştirme” amacının, yalnızca bir yazılım uygulaması değil; kaynak, yöntem, modelleme ve yorumlanabilirlik katmanlarını birlikte içeren bilimsel bir çerçeve olarak ele alındığını göstermektedir.

Bulgular literatürle birlikte okunduğunda daha geniş bir yöntemsel sonuç da ortaya çıkmaktadır. Son yıllarda çok dilli büyük modellerin düşük kaynaklı dillere güçlü bir başlangıç performansı sunduğu açıktır; ancak bu tezdeki deneyler, bu modellerin dil-özel ve alan-özel bilgi kaynaklarıyla desteklendiğinde daha dengeli sonuç verdiğini göstermektedir. Bu tablo, bir yandan çok dilli transformer literatürünün sağladığı kazanımları doğrulamakta, diğer yandan da morfolojik zenginlik ve alan kayması koşullarında hibrit tasarımların hâlâ anlamlı bir araştırma yönü olduğunu göstermektedir (Conneau vd. 2020; Rust vd. 2021; Dehkharghani vd. 2017; Mutinda vd. 2023). Özellikle *MahSA* ve alanlar arası sağlam temsil öğrenmesi katmanlarında görülen kazançlar, “daha büyük model” ile “daha uygun bilgi bileşimi” arasındaki farkı görünür kılmaktadır.

Ablasyon sonuçları da bu genel tabloyu desteklemektedir. Ne ABSA katmanında ne de *MahSA* katmanında iyileşme tek bir bileşenin dramatik etkisine indirgenebilmektedir. Tersine, sözlük, karakter, morfoloji ve bağlamsal temsil bloklarının her biri belirli hata bölgelerini azaltmakta; asıl güçlü sonuç ise bunların birlikte kullanıldığı düzenlerde ortaya çıkmaktadır. Bu durum, hibrit mimarinin bileşenlerinin rastgele yığılmadığını; farklı bilgi eksikliklerini karşılıklı olarak telafi ettiğini göstermektedir. Özellikle nötr sınıfta görülen düzenli kazanç, bu sinerjinin en açık göstergelerinden biridir. Çünkü nötr sınıf çoğu zaman belirgin kutupluluk işaretlerinden çok, bağlamın, örtüklüğün ve sınır durumlarının doğru yorumlanmasına bağlıdır.

Hata analizi tarafında ise tüm katmanlarda tekrarlayan bir örüntü görülmektedir: sarkazm, ironi, örtük duygu, karma değerlendirme, kod değiştirimi ve alan jargonunun hızlı değişimi. Bu hata tiplerinin bölüm boyunca yeniden ortaya çıkması önemlidir; çünkü sorunların belirli bir modelde rastgele oluşmadığını, veri ve dil yapısının doğasından kaynaklanan kalıcı güçlükler olduğunu düşündürmektedir. Özellikle finans ve iş dünyası ile kamu hizmetleri alanlarında nötr sınıfın kırılğan kalması, tezin en güçlü modellerinde bile bağlamın ve pragmatik ipuçlarının tümüyle çözülemediğini göstermektedir. Bu yönüyle bulgular, başarı artışlarını görünür kılmak kadar,

henüz tam olarak çözülmemiş zorluk alanlarını da sistematik biçimde haritalamaktadır.

Pratik açıdan bakıldığında bölümden çıkan ana sonuç şudur: Azerbaycan dili için güvenilir bir duygu analizi sistemi kurmak isteyen bir araştırmacı ya da uygulayıcı, yalnızca tek bir güçlü bağlamsal model seçerek optimum sonuca ulaşmayı beklememelidir. Dilin morfolojik yapısı, veri gürültüsü, alan kayması ve nötr sınıfın belirsizliği nedeniyle sembolik ön bilgi, yüzey ipuçları ve alan-duyarlı düzenlemeler hâlâ kayda değer bir işlev görmektedir. Bu nedenle bu bölümde sunulan sonuçlar, yalnızca önerilen yöntemlerin performansını raporlamamakta; aynı zamanda Azerbaycan dili için hangi tasarım ilkelerinin daha umut verici olduğunu da ortaya koymaktadır.

Bununla birlikte açıkta kalan güçlükler de vardır. Sarkazm ve ironi, kod değiştirimi, örtük duygu, karma değerlendirmeler ve alan jargonunun hızlı değişimi bütün alt çalışmalarda kalıcı hata kaynakları olarak görünmektedir. Özellikle nötr sınıf, finans ve iş dünyası metinlerinde ve kamu hizmetleri alanında daha kırılgan bir sınıf olmaya devam etmektedir. Bu nedenle tezden çıkan genel tablo, “sorun bütünüyle çözülmüştür” biçiminde değil; “hangi bilgi türlerinin hangi hata bölgelerinde daha sistematik yarar sağladığı artık daha açık görülebilmektedir” biçiminde okunmalıdır. Bu nokta, bir sonraki bölümde sunulan genel değerlendirme ve gelecek çalışma önerileri için de zemin oluşturmaktadır.

5. GENEL DEĞERLENDİRME

Bu bölümde, dördüncü bölümde sunulan deneysel bulgular tek tek sonuç kümeleri düzeyinde değil, tezin tamamına kuş bakışı bakan bir sentez düzeyinde ele alınmaktadır. Amaç, tartışma bölümündeki ayrıntılı yorumları tekrarlamak değil; problemden yöntem, yöntemden bulguya ve bulgudan anlama uzanan araştırma hattını tek bir tutarlı anlatı içinde yeniden kurmaktır. Bu doğrultuda önce tezin genel değerlendirmesi yapılmakta, ardından araştırma sorularına verilen yanıtlar sistematik biçimde sunulmakta, katkılar olgunlaşmış halleriyle yeniden kristalize edilmekte, çalışmanın sınırlılıkları çerçevelenmekte ve gelecekte izlenebilecek araştırma doğrultuları tartışılmaktadır.

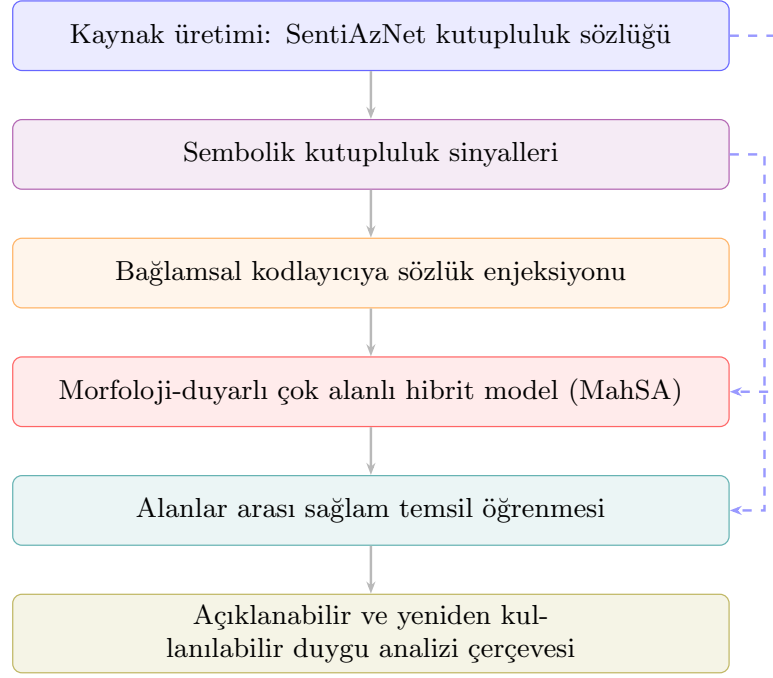
5.1 Tezin Genel Değerlendirmesi

Bu tez, Azerbaycan dili için duygu analizinin tek başına bir sınıflandırma problemi olarak ele alınamayacağını gösteren bir araştırma hattı üzerine kurulmuştur. Çalışmanın çıkış noktası, düşük kaynaklı ve morfolojik açıdan zengin bir dilde duygu analizinin yalnızca veri azlığı nedeniyle değil; aynı zamanda kutupluluk sözlüklerinin sınırlılığı, alanlar arası kutup kaymaları, nötr sınıfın yapısal zorluğu ve model kararlarının açıklanabilirliğine yönelik gereksinimler nedeniyle de güçleştiği gözlemdir. Tez boyunca geliştirilen yöntemler ve elde edilen bulgular, bu zorlukların tek bir teknik hamle ile değil, kaynak geliştirme, modelleme, temsil öğrenmesi ve açıklanabilirlik boyutlarının birlikte ele alınmasıyla daha tutarlı biçimde yönetilebildiğini ortaya koymuştur.

Bu bakımdan tez, birbirinden kopuk alt çalışmalar toplamı olarak değil, kademeli biçimde genişleyen bir araştırma anlatısı olarak okunmalıdır. İlk aşamada Azerbaycan dili için güvenilir bir kutupluluk kaynağı oluşturulmuş; ikinci aşamada bu sembolik kaynağın alan-özgü bir bağlamsal modele nasıl aktarılabilceği sınanmış; üçüncü aşamada çok alanlı ve nötr sınıfı merkeze alan hibrit bir duygu analizi mimarisi geliştirilmiş; sonrasında ise dikkat, SHAP ve hata çözümlemeleri ile kararların hangi bilgi kanallarından beslendiği görünür kılınmıştır. Temsil öğrenmesi boyutunda ulaşılan son aşama ise araştırma hattını tekil sınıflandırma başarımının ötesine taşıyarak alan kaymasına karşı daha dayanıklı ve daha taşınabilir duygu temsilleri üretmenin de bu çerçevenin doğal bir parçası olduğunu göstermiştir.

Şekil 5.1, tez boyunca geliştirilen kaynak, modelleme, temsil öğrenmesi ve açıklanabilirlik katman-

larının birbirini besleyen bütüncül yapısını özetlemektedir. Şekil, tezin tek bir model başarısına indirgenemeyeceğini; *SentiAzNet*, alan-özümlü sözlük enjeksiyonu, *MahSA* ve alanlar arası sağlam temsil öğrenmesinin aynı araştırma programının birbirini tamamlayan halkaları olduğunu göstermektedir.



Şekil 5.1 Tezin kavramsal ve deneysel katkılarının birbirini besleyen bütüncül akışı

Bu bütüncül yapı içinde ilk temel kazanım, Azerbaycan dili için yeniden kullanılabilir ve denetlenmiş bir kutupluluk kaynağının üretilmiş olmasıdır. *SentiAzNet*, tezde geliştirilen sonraki bütün yöntem katmanları için yalnızca bir yardımcı sözlük değil, aynı zamanda araştırma hattının sembolik bilgi temelini oluşturan kurucu bileşen olarak işlev görmüştür. İç ve dış doğrulama bulguları birlikte değerlendirildiğinde, bu kaynağın tek başına nihai çözüm olmaktan çok, sonraki modelleri tamamlayan yüksek hassasiyetli bir ön bilgi kaynağı olarak en fazla değer ürettiği anlaşılmaktadır.

İkinci temel kazanım, bu sembolik ön bilginin çağdaş bağlamsal modellere aktarılmasının yöntemsel olarak mümkün ve yararlı olduğunun gösterilmesidir. Finans ve iş dünyası alanındaki aspekt-tabanlı duygu analizi deneyleri, sözlük bilgisinin modele doğrudan ve sert biçimde değil, dikkat dağılımını yönlendiren veya temsil katmanını tamamlayan yumuşak bir sinyal olarak verildiğinde daha etkili sonuçlar üretebildiğini ortaya koymuştur. Böylece tez, sembolik bilgi ile sinirsel modelleme arasında karşıtlık kurmak yerine, bunların uygun tasarım kararlarıyla birbirini tamamlayabileceğini göstermiştir.

Tezin ana yöntemsel omurgasını oluşturan *MahSA* ise bu çizgiyi daha da ileri taşımıştır. Çok alanlı ve üç sınıflı bir senaryoda, sözlük tabanlı kutupluluk sinyalleri, karakter n-gramları, morfoloji-duyarlı yüzey ipuçları, wPPMI ile bağlama yayılan duygu sinyalleri ve bağlamsal transformer temsilleri tek bir nöro-sembolik çerçeve içinde birleştirilmiştir. Elde edilen bulgular, özellikle nötr sınıfın modellenmesinde ve veri miktarı görece sınırlı alanlarda, bu çok kaynaklı tasarımın tek başına bağlamsal taban modellere göre daha dengeli ve daha kararlı davranabildiğini göstermiştir. Bu sonuç, Azerbaycan dili gibi eklemeli ve düşük kaynaklı dillerde temsil çeşitliliğinin yalnızca ek bir ayrıntı değil, performansın yapısal belirleyicilerinden biri olduğunu düşündürmektedir.

Temsil öğrenmesi boyutunda elde edilen bulgular ise tezin katkısını yalnızca “daha iyi bir sınıflandırıcı üretme” düzeyinde bırakmamaktadır. Alanlar arası sağlam cümle temsilleri üzerine kurulan son aşama, duygu analizinde asıl meselenin yalnızca karar katmanında doğru etiketi üretmek olmadığı; farklı alanlarda benzer duygusal işlev gören örnekleri ortak bir temsil uzayında daha tutarlı biçimde düzenleyebilmek olduğunu göstermiştir. Bu bakımdan tez, sınıflandırma başarımı ile temsil geometrisi arasında doğrudan bir ilişki kurmakta ve alan kaymasına karşı dayanıklılığı ayrı bir yan konu değil, tezin merkezî araştırma hedeflerinden biri olarak konumlandırmaktadır.

Genel bir sentez düzeyinde bakıldığında, bu tezden çıkan ana sonuç şudur: düşük kaynaklı ve morfolojik açıdan zengin bir dilde duygu analizinin başarımı, tek başına daha büyük ve daha genel modeller kullanılarak açıklanamaz. Dil-özel sembolik kaynaklar, morfolojiye duyarlı yüzey işaretleri, alan farklarını dikkate alan tasarımlar ve açıklanabilirlik düzenekleri bir araya geldiğinde, hem daha güvenilir hem de daha yorumlanabilir bir çerçeve kurulabilmektedir. Bu yönüyle tez, büyük dil modelleri çağında dahi dile özgü dilbilimsel bilginin ve dikkatli kaynak geliştiriminin önemini koruduğunu ampirik olarak desteklemektedir.

5.2 Araştırma Sorularına Verilen Yanıtlar

Giriş bölümünde ortaya konulan araştırma soruları, dördüncü bölümde sunulan bulgular ışığında sistematik biçimde yanıtlanabilmektedir. Bu bölümde amaç, söz konusu soruları yeniden kurmak ve her birinin ne ölçüde karşılandığını tezin genel argümanına bağlayarak netleştirmektir.

AS1. Azerbaycan dili için, diğer dillere ait zengin kaynaklar, geri çeviri, anlamsal doğrulama

ve uzman etiketlemesi birlikte kullanılarak güvenilir ve yeniden kullanılabilir bir kutupluluk sözlüğü oluşturulabilir mi? Bu soruya olumlu yanıt verilmiştir. *SentiAzNet* hattı, çoklu kaynak birleştirmesi, geri çeviri ve insan doğrulaması bir araya getirildiğinde Azerbaycan dili için güvenilir bir kutupluluk kaynağının üretilbildiğini göstermiştir. İç doğrulamadaki güçlü sonuçlar ve dış veri üzerindeki taşınabilirlik, bu kaynağın yalnızca laboratuvar koşullarına özgü olmadığını, ancak en yüksek değerini yardımcı sembolik ön bilgi olarak ürettiğini ortaya koymuştur. Böylece tezin ilk araştırma sorusu, tez boyunca kullanılan bütün yöntem katmanlarının dayandığı sembolik tabanın gerçekten kurulabildiğini göstermesi bakımından karşılanmıştır.

AS2. Oluşturulan kutupluluk sözlüğünden türetilen sinyaller, bağlamsal transformer tabanlı modellere eklendiğinde Azerbaycan dili için duygu analizi başarımına anlamlı katkı sağlayabilir mi? Bu soruya da olumlu yanıt verilmiştir; ancak katkının en yüksek olduğu koşullar belirli sınırlar içinde tanımlanmıştır. Finans ve iş dünyası alanındaki ABSA deneyleri, sözlük bilgisinin bağlamsal modele doğrudan özellik olarak eklenmesiyle de yarar sağladığını, ancak dikkat önyargısı gibi daha seçici bir entegrasyonla daha tutarlı kazanımlar üretebildiğini göstermiştir. Bu bulgu, sembolik bilginin salt varlığının değil, modele nasıl enjekte edildiğinin de araştırma problemi olduğunu ortaya koymuştur. Dolayısıyla ikinci soru, “sözlük yararlı mı?” sorusundan öte, “sözlük hangi tasarım altında yararlı?” sorusuna da yanıt üretmiştir.

AS3. Kutupluluk sözlüğü bilgisi, karakter düzeyi ipuçları, bağlamsal gömme temsilleri ve morfoloji-duyarlı yüzey özellikleri birlikte kullanıldığında, çok alanlı ve üç sınıflı duygu analizi probleminde tek başına bağlamsal modellere göre daha güçlü sonuçlar elde edilebilir mi? Bu sorunun yanıtı genel olarak evettir. *MahSA* mimarisi, farklı bilgi kaynaklarını tek bir füzyon yapısında birleştirerek güçlü bağlamsal taban modellere kıyasla daha dengeli bir başarı örüntüsü üretmiştir. Bu sonuç, tek bir temsil ailesine yaslanmak yerine, farklı bilgi kanallarını kontrollü biçimde bir araya getiren hibrit tasarımların Azerbaycan dili için daha uygun bir yöntemsel yönelim sunduğunu göstermektedir.

AS4. Önerilen hibrit yaklaşım, özellikle nötr sınıfın modellenmesinde ve veri miktarı görece sınırlı alanlarda, mevcut güçlü taban modellere göre daha kararlı bir performans sergileyebilir mi? Bulgular bu soruya da olumlu yanıt vermektedir. Özellikle nötr sınıfın belirsiz ve dağınık yapısı ile kamu hizmetleri gibi daha küçük ve zor alanlarda elde edilen sonuçlar, hibrit tasarımın

yalnızca genel ortalamayı artırmadığını; aynı zamanda klasik taban modellerin daha kırılgan olduğu bölgelerde daha dengeli davrandığını göstermiştir. Bu durum, tezin nötr sınıf ve alan kaymasını ikincil değil, merkezî problem alanları olarak konumlandıran yaklaşımını desteklemektedir.

AS5. Alan-özgü kutupluluk farklılıkları ve alan kayması dikkate alındığında, duyguya duyarlı ve alanlar arası daha sağlam cümle temsilleri öğrenmek mümkün müdür? Bu soruya verilen yanıt, mevcut bulgular temelinde olumlu ve umut verici niteliktedir. Alanlar arası sağlam temsil öğrenmesi hattı, alan-içi başarı ile bırak-bir-alan-dışı-bırak düzenindeki başarıyı birlikte değerlendirerek, duyguya özgü bilgi ile alan-bağımlı kestirmeler arasındaki gerilimin temsil düzeyinde yönetilebileceğini göstermiştir. Böylece tezin temsil öğrenmesi boyutu, genel savın yalnızca sınıflandırıcı katmanında değil, gömme uzayının kuruluşunda da karşılık bulduğunu ortaya koymuştur.

AS6. SHAP, dikkat örüntüleri, sözlük eşleşmeleri ve hata çözümlemeleri birlikte kullanıldığında, model kararlarının açıklanabilirliği artırılabilir ve başarımın hangi dilsel işaretlerden beslendiği sistematik biçimde ortaya konabilir mi? Bu sorunun yanıtı kısmen fakat anlamlı ölçüde evettir. Tezde kullanılan açıklanabilirlik ve hata analizi düzenekleri, modellerin hangi sözcük, hangi morfolojik işaret veya hangi sözlük sinyalden beslendiğini daha görünür hale getirmiştir. Bununla birlikte ironi, sarkazm, derin pragmatik örüntüler ve kod değiştirimi gibi bazı zorlu hata bölgeleri, açıklanabilirliğin teknik araçlarla bütünüyle tamamlanmadığını da göstermiştir. Dolayısıyla altıncı soru, açıklanabilirliğin ek bir raporlama tercihi değil, tezin kurucu değerlendirme katmanlarından biri olduğunu doğrulamıştır.

Bu araştırma sorularına verilen yanıtlar bir arada okunduğunda, tezin temel savı farklı düzeylerde sınanmış ve her düzeyde belirli ölçülerde desteklenmiştir. Kaynak geliştirme, modelleme, temsil öğrenmesi ve açıklanabilirlik hatlarının birlikte ilerlemesi, çalışmayı bağımsız deney kümeleri toplamı olmaktan çıkarıp tek bir araştırma yayına dönüştürmektedir.

5.3 Tezin Bilimsel ve Uygulamalı Katkılarının Sentezi

Giriş bölümünde katkılar birer araştırma vaadi olarak sunulmuştu. Bu aşamada ise aynı katkılar, tez boyunca elde edilen bulgular ışığında olgunlaşmış ve kanıtlanmış halleriyle yeniden ifade

edilebilir. Başka bir deyişle bu bölümde vaat dilinden kanıt diline geçilmektedir.

Kaynak ve veri katkıları. Tezin en görünür katkılarından biri, Azerbaycan dili için denetlenmiş bir kutupluluk sözlüğünün ve büyük ölçekli çok alanlı bir yorum derleminin geliştirilmiş olmasıdır. *SentiAzNet*, yalnızca bir kaynak üretimi değil; daha sonra sinirsel modellere aktarılacak sembolik bilginin güvenilir biçimde inşa edilebildiğinin kanıtıdır. Benzer biçimde beş alana yayılan elle etiketlenmiş yorum derlemi, Azerbaycan dili için nötr sınıf davranışı, alan kayması ve çok alanlı değerlendirme gibi sorunların sistematik biçimde çalışılabilmesini mümkün kılmıştır. Bu iki çıktı birlikte düşünüldüğünde, tez yalnızca bir model önermemiş; aynı zamanda bu modelin çalışabileceği araştırma altyapısını da kurmuştur.

Metodolojik katkılar. Finans ve iş dünyası ABSA çalışması ile *MahSA* mimarisi, tezin yöntemsel çekirdeğini oluşturmaktadır. ABSA hattı, sözlük bilgisinin çağdaş bağlamsal modellere hangi tasarım ilkeleriyle aktarılabilceğini göstermiş; *MahSA* ise bu ilkeleri daha geniş bir problem ortamında çok kaynaklı ve morfoloji-duyarlı bir hibrit mimari altında birleştirmiştir. Böylece tez, “büyük model kullanmak” ile “dilbilimsel bilgi kullanmak” arasındaki sahte karşıtlığı aşarak, düşük kaynaklı diller için daha uygun bir nöro-sembolik tasarım çizgisi önermiştir.

Ampirik katkılar. Tezde yürütülen kapsamlı karşılaştırmalar, ablasyon deneyleri ve alan bazlı çözümlemeler, hangi bileşenin hangi koşulda değer ürettiğini görünür kılmıştır. Özellikle nötr sınıf, alanlar arası performans farkları ve veri miktarı sınırlı alanlar üzerinde yürütülen çözümlemeler, katkıların yalnızca tek bir ortalama başarı skoru düzeyinde değil, performansın dağılımı ve kararlılığı düzeyinde de değerlendirildiğini göstermektedir. Bu yönüyle çalışma, Azerbaycan dili için yalnızca “daha iyi sonuç” değil, “daha anlaşılabilir sonuç deseni” üretmiştir.

Temsil öğrenmesi katkısı. Alanlar arası sağlam cümle temsilleri üzerine kurulan son aşama, tezin katkısını sınıflandırma katmanının ötesine taşımaktadır. Bu katkı, farklı alanlardan gelen ama benzer duygu işlevi taşıyan örnekleri daha tutarlı bir gömme uzayında düzenlemeye dönük bir araştırma yönü açmaktadır. Böylece tez, duygu analizini yalnızca etiket üretme problemi olarak değil, duyguya duyarlı temsil inşası problemi olarak da yeniden düşünmeye davet etmektedir.

Açıklanabilirlik ve değerlendirme katkısı. Tezde kullanılan SHAP çözümlemeleri, dikkat örün-

tüleri, sözlük eşleşmeleri, ablasyon düzenekleri ve hata analizi hatları; modellerin yalnızca ne kadar başarılı olduğunu değil, hangi bilgi kanallarından beslendiğini de görünür hale getirmiştir. Bu durum, açıklanabilirliği sonradan eklenmiş bir raporlama katmanı olmaktan çıkarıp yöntemsel tasarımın ve deneysel değerlendirme mantığının kurucu unsurlarından biri haline getirmektedir.

Uygulamaya dönük katkı. Tez kapsamında geliştirilen yöntem ailesi, modüler yapısı sayesinde farklı uygulama senaryolarına aktarılacak bir gerçekleştirim zemini de üretmiştir. Bu zemin, sosyal medya izlemesi, kullanıcı yorumu analizi, finans ve iş dünyası metinlerinde hedefe bağlı duygu çözümleme ve kamu hizmetlerine ilişkin geri bildirimlerin değerlendirilmesi gibi uygulamalara araştırma temelli bir başlangıç noktası sağlamaktadır.

Büyük resim açısından bakıldığında, bu tezden çıkan en önemli derslerden biri şudur: düşük kaynaklı ve morfolojik açıdan zengin dillerde ilerleme, yalnızca daha büyük ve daha genel modeller kullanmakla sağlanmamaktadır. Tezin bulguları, dile özgü sembolik bilginin, yüzeyde görülen morfolojik işaretlerin, alan farklarını dikkate alan modelleme tercihleri ile birlikte hâlâ kritik bir rol oynadığını göstermektedir. Bu yönüyle çalışma, geniş model paradigmasının önemini reddetmemekte; ancak özellikle düşük kaynak koşullarında ölçeklenmiş genel temsiller ile dil-özel bilginin birlikte düşünülmesi gerektiğini savunmaktadır. Bu tezden çıkan ilkeler, yalnızca Azerbaycan dili için değil, benzer tipolojik özellikler taşıyan ve veri açısından sınırlı başka diller için de anlamlı bir yöntemsel çerçeve sunmaktadır.

5.4 Sınırlılıklar

Bu çalışmanın temel sınırlılıkları, tezin katkılarını geçersiz kılan kusurlar olarak değil, bulguların hangi koşullarda geçerli olduğunu belirleyen sınırlar olarak okunmalıdır. Çalışma tek bir dil üzerinde gerçekleştirilmiş, ağırlıklı olarak metin tabanlı veriyle sınırlandırılmış ve deneysel değerlendirme büyük ölçüde yazılı kullanıcı yorumları üzerine kurulmuştur. Bu nedenle önerilen çerçevenin diğer eklemeli dillere, farklı platform türlerine ya da çok kipli ortamlara doğrudan genellenebilirliği bu tez kapsamında deneysel olarak sınanmamıştır. Bununla birlikte yöntem ailesinin modüler yapısı ve farklı bilgi bloklarının açık biçimde tanımlanmış olması, bu tür genellemelerin sonraki çalışmalar için yapısal olarak mümkün olduğunu göstermektedir.

İkinci temel sınırlılık, duygu analizinin en dirençli hata bölgelerinde ortaya çıkmaktadır. İroni, sarkazm, örtük değerlendirme, karma duygu, alan-özgü jargon ve kod değiştirimi gibi örüntüler, hem sözlük tabanlı hem de bağlamsal modeller için güçlük üretmeye devam etmektedir. Benzer biçimde tezde benimsenen morfoloji-duyarlı yaklaşım, tam bir biçimbilim çözümleyiciden ziyade yüzey işaretleri ve kurallı sezgiler üzerinden kurulmuştur. Bu tercih gürültülü kullanıcı verisi açısından pragmatik bir avantaj sağlamakla birlikte, biçimbilimsel yapının tüm derinliğini kapsadığı iddiasında değildir. Ayrıca açıklanabilirlik katmanı teknik açıdan güçlü bulgular sunmasına rağmen, insan-merkezli kullanıcı değerlendirmeleriyle desteklenmemiştir. Dolayısıyla açıklamaların teknik tutarlılığı ile son kullanıcı yararı arasındaki ilişki, ileride ayrıca sınanması gereken bir boyut olarak kalmaktadır.

Bu sınırlılıklar, tezin temel argümanını zayıflatmaktan çok, onun geçerlilik alanını netleştirmektedir. Başka bir ifadeyle çalışma, problemin tamamını çözdüğü iddiasında değildir; fakat düşük kaynaklı ve morfolojik açıdan zengin bir dilde hangi bilgi türlerinin hangi koşullarda sistematik yarar sağladığını artık daha açık biçimde ortaya koymaktadır.

5.5 Gelecekte Yapılabilecek Çalışmalar ve Öneriler

Bu tez, tamamlanmış bir çalışma olmanın yanı sıra, devam ettirilebilir bir araştırma hattı da açmaktadır. Bu nedenle gelecekte yapılabilecek çalışmalar yalnızca mevcut sınırlılıkların tersi olarak değil, tezin ortaya koyduğu yöntemsel ilkeleri yeni bağlamlara taşıyan bir araştırma gündemi olarak düşünülmelidir.

Kısa vadeli uzantılar. İlk aşamada, mevcut çerçevenin doğrudan devamı niteliğinde çalışmalar öne çıkmaktadır. Sarkazm, ironi, örtük değerlendirme ve ince taneli nötr örüntülerin ayrıca etiketlendiği veri kümeleri oluşturulması; *SentiAzNet*'in yarı denetimli veya aktif öğrenme ile dinamik biçimde güncellenmesi; açıklanabilirlik katmanının insan-merkezli değerlendirmelerle sınanması; ve nötr sınıf için daha ayrıntılı kalibrasyon mekanizmalarının geliştirilmesi bu bağlamda ilk akla gelen yönelimlerdir. Bunlar, tezin mevcut mimarisini koruyarak onun en belirgin zayıflık bölgelerini daha sistematik biçimde hedef alabilir.

Orta vadeli araştırma yönleri. İkinci düzeyde, tezin temel fikirlerini yeni bağlamlara taşıyan

alışmalar ne ıkmaktadır. Bu kapsamda Azerbaycan dili iin geliřtirilen erevenin Trke bařta olmak zere tipolojik olarak yakın diğerk Trk dillerine aktarılması; apraz platform, apraz zaman ve apraz alan aktarım senaryolarında sınanması; aspekt ıkarımı, hedef belirleme, tutum analizi ya da argmantasyon gibi komřu grevlere uyarlanması; ve metin dıřı sinyallerin de hesaba katıldıđı ok kipli duygu analizi hatlarına geniřletilmesi anlamlı arařtırma dođrultuları sunmaktadır. Bu ynelimler, tezin dil ve grev sınırlarını kontroll bimde geniřletmeye imkn verecektir.

Uzun vadeli vizyon. Daha uzun vadede ise bu tezden ıkan asıl vizyon, dřk kaynaklı diller iin morfoloji-duyarlı ve sembolik bilgiyle zenginleřtirilmiř dil teknolojilerinin daha genel n eđitim ve byk model paradigmalarına nasıl entegre edilebileceđi sorusudur. Tez boyunca elde edilen bulgular, dile zg bilginin byk modeller karřısında deđersizleřmediđini; tersine, uygun bimde temsil edildiđinde onları tamamlayıcı ve glendirici bir rol oynayabildiđini gstermektedir. Bu nedenle gelecekte morfoloji-duyarlı sinyallerin, kutupluluk nbilgisinin ve alan-duyarlı temsil ilkelerinin parametre-verimli ince ayar veya n eđitim dzenekleri iine daha erken ařamada dahil edilmesi, dřk kaynaklı diller iin mevcut modelleme paradigmasını dnřtrebilecek bir arařtırma ufku sunmaktadır.

Sonc olarak bu tez, Azerbaycan dili iin duygu analizinin yalnızca daha yksek bir bařarı skoru retme meselesi olmadıđını; kaynak geliřtirme, morfoloji-duyarlı modelleme, alanlar arası sađlamlık ve aıklanabilirlik boyutlarını birlikte gerektiren btncl bir arařtırma alanı olduđunu gstermiřtir. 124.251 yorumluk ok alanlı derlem zerinde geliřtirilen ve *SentiAzNet* ile bařlayan bu arařtırma hattı, *MahSA* ile glenmiř, alanlar arası sađlam temsil đrenmesi ile geniřlemiř ve bylece dřk kaynaklı, morfolojik aıdan zengin diller iin srdrlebilir bir yntemsel yol haritası sunmuřtur. Bu ynyle tez, yalnızca tamamlanmıř bir doktora alışması deđil; benzer diller ve benzer problemler iin ileriye tařınabilecek btnkl bir arařtırma programının da bařlangı noktası olarak deđerlendirilebilir.

6. SONUÇ

Bu tez, Azerbaycan dili için duygu analizinin tek başına bir sınıflandırma problemi olarak ele alınmasının yeterli olmadığını; bu alanın güvenilir dil kaynağı üretimi, morfoloji-duyarlı modelleme, alanlar arası sağlamlık ve açıklanabilir karar mekanizmalarını birlikte içeren bütüncül bir araştırma çerçevesi gerektirdiğini deneysel olarak göstermiştir. Bu amaçla önerilen yedi aşamalı metodolojik omurga içinde kutupluluk sözlüğü geliştirme, alan-özümlü aspekt-tabanlı modelleme, çok alanlı hibrit duygu analizi ve alan kaymasına dayanıklı temsil öğrenmesi tek bir araştırma hattında birleştirilmiş; elde edilen bulgular, sembolik ve sinirsel bilgi kaynaklarının birlikte kullanılmasının Azerbaycan dili için daha güvenilir, daha açıklanabilir ve daha genellenebilir bir duygu analizi yaklaşımı sunduğunu doğrulamıştır.

Tezin ilk temel bulgusu, Azerbaycan dili için yeniden kullanılabilir bir kutupluluk sözlüğü üretilebildiğinin gösterilmesidir. Çok kaynaklı sözlük hattı üzerinden 9.614 İngilizce lemma ile başlayan süreç sonunda 2.482 Azerbaycan dili girdisinden oluşan *SentiAzNet* geliştirilmiş; sözlük doğrulama hattı iç değerlendirmede 0.915 doğruluk ve 0.892 Macro-F1 üretmiştir. Bu sonuç, düşük kaynaklı bir dil için yüksek kesinlikli sembolik önbilginin üretilebildiğini ve sonraki modelleme katmanları için sağlam bir başlangıç zemini sunduğunu göstermiştir.

İkinci temel bulgu, sözlük bilgisinin bağlamsal modellere uygun tasarımıyla aktarıldığında anlamlı kazanç sağlayabildiğidir. Finans ve iş dünyası alanındaki aspekt-tabanlı deneylerde yalın XLM-R modeli 0.702 Macro-F1 düzeyinde kalırken, sözlük bilgisi ve alan-uyarlamalı dikkat önyargısı içeren yapı 0.738 Macro-F1 değerine ulaşmıştır. Bu bulgu, sözlük bilgisinin tek başına nihai karar mekanizması olarak değil, bağlamsal kararları yönlendiren seçici bir yardımcı sinyal olarak kullanıldığında daha etkili olduğunu ortaya koymuştur.

Üçüncü temel bulgu, çok alanlı ve üç sınıflı duygu analizi probleminde morfoloji-duyarlı nöro-sembolik hibrit tasarımın güçlü bir üstünlük sağlayabildiğidir. *MahSA*, beş alanın ortalamasında 0.725 Macro-F1 elde ederek XLM-R ince ayar modelinin 0.688 ve Az-RoBERTa ince ayar modelinin 0.697 düzeylerini aşmıştır. Kazanımın özellikle nötr sınıf ve kamu hizmetleri gibi görece zor alanlarda daha belirgin hale gelmesi, sembolik ön bilgi, karakter temsili ve morfoloji-duyarlı yüzey işaretlerinin veri sınırlı koşullarda kritik değer taşıdığını göstermiştir.

Dördüncü temel bulgu, duygu analizinde başarımlarının artışının yalnızca son sınıflandırma katmanından değil, temsil uzayının niteliğinden de beslendiğinin gösterilmesidir. Alanlar arası sağlam temsil öğrenmesi deneyleri, duyguya duyarlı ve alan kaymasına karşı daha dayanıklı cümle temsillerinin hem alan-ıç i hem de bırak-bir-alan-dışı-bırak düzenlerinde daha dengeli sonuçlar üretebildiğini ortaya koymuştur. Böylece tez, Azerbaycan dili için duygu analizini yalnızca daha iyi sınıflandırma skoru üretme hedefinden çıkarıp, daha sağlam ve daha taşınabilir temsil öğrenme problemine de genişletmiştir.

Bu tez, literatürdeki dört temel boşluğa doğrudan yanıt vermiştir: Azerbaycan dili için güvenilir bir kutupluluk sözlüğü geliştirilmiş, sözlük bilgisinin alan-özgü bağlamsal modellere nasıl entegre edileceği gösterilmiş, büyük ölçekli ve çok alanlı bir duygu analizi veri zemini kurulmuş ve morfoloji-duyarlı hibrit modelleme ile alanlar arası sağlam temsil öğrenmesi aynı araştırma çerçevesi içinde birleştirilmiştir. Buna ek olarak açıklanabilirlik, bu çalışmada ikincil bir raporlama katmanı olarak değil, model kararlarının hangi dilsel işaretlerden beslendiğini görünür kılan kurucu bir bileşen olarak ele alınmıştır. Bu yönüyle tez, Azerbaycan dili için duygu analizini parçalı deneylerden oluşan bir uygulama alanı olmaktan çıkarıp, kaynak, yöntem, deneysel doğrulama ve yorumlama boyutlarını bütünleştiren sistematik bir araştırma programına dönüştürmektedir.

Pratik açıdan bakıldığında, tez kapsamında geliştirilen çerçeve Azerbaycan dilinde sosyal medya izleme, müşteri geri bildirimi çözümleme, finans ve iş dünyası yorumlarının hedefe bağlı analizi ve kamu hizmetlerine ilişkin değerlendirme sistemleri için daha güvenilir bir temel sunmaktadır. Özellikle veri miktarının sınırlı, dilsel gürültünün yüksek ve nötr sınıfın problemli olduğu koşullarda, yalnızca büyük çok dilli modellere dayanmak yerine, dile ve alana özgü sembolik kaynakları model tasarımına dahil etmenin daha dengeli ve daha izlenebilir sistemler kurulmasına imkân verdiği gösterilmiştir.

Bununla birlikte, bu tezin sınırları da açıktır. İroni ve sarkazm içeren örnekler, kod değiştirimi, hızla değişen alan-özgü jargon ve tam morfolojik çözümleme gerektiren daha karmaşık yapılar hâlâ önemli güçlük alanları olarak kalmaktadır. Ayrıca çalışma tek bir dil üzerinde ve metin tabanlı senaryolarla yürütülmüş; çok kipli duygu analizi ile diğer Türk dillerine doğrudan deneysel genelleme bu tez kapsamında sınırlanmamıştır. Ancak bu sınırlılıklar, tezin temel katkılarını geçersiz kılmamakta; aksine önerilen çerçevenin hangi yönlerden geliştirilebileceğini açık biçimde

göstermektedir.

Gelecekte, *SentiAzNet*'in dinamik olarak güncellenen bir kaynağa dönüştürülmesi, *MahSA* ilkelerinin diğer Türk dillerine aktarılması, çok kipli duygu analizi senaryolarının çalışılması ve morfoloji-duyarlı bilginin büyük dil modellerinin alan-uyarlamalı ön eğitim süreçlerine daha sıkı biçimde entegre edilmesi özellikle değerli araştırma yönleri olarak görünmektedir. Bu yönelimler, bu tezde önerilen yaklaşımın tek bir uygulama ile sınırlı olmadığını, sürdürülebilir ve genişleyebilir bir araştırma hattı açtığını göstermektedir.

Sonuç olarak bu tez, Azerbaycan dili için duygu analizinin daha büyük bir model seçme problemi değil, doğru dil kaynağını, doğru dilsel sezgileri ve doğru temsil düzeylerini aynı çerçevede buluşturma problemi olduğunu ortaya koymuştur. Sunulan yöntemsel yol haritası, düşük kaynaklı ve morfolojik açıdan zengin dillerde duygu analizinin nasıl çalışılması gerektiğine ilişkin somut, sınanmış ve genişletilebilir bir temel bırakmaktadır.

KAYNAKLAR

- Abdullah, T. ve Ahmet, A. (2022). Deep learning in sentiment analysis: Recent architectures. *ACM Computing Surveys*, 55(8):1–37.
- Ali, H. M. U., Farooq, Q., Imran, A., ve El Hindi, K. (2025). A systematic literature review on sentiment analysis techniques, challenges, and future trends. *Knowledge and Information Systems*, 67(5):3967–4034.
- Aliyu, Y., Sarlan, A., Danyaro, K. U., Rahman, A. S. B., ve Abdullahi, M. (2024). Sentiment analysis in low-resource settings: a comprehensive review of approaches, languages, and data sources. *IEEE Access*, 12:66883–66909.
- Alizada, T., Suleymanov, U., ve Rustamov, Z. (2024). Contextualized word embeddings in azerbaijani language. In *2024 IEEE 18th International Conference on Application of Information and Communication Technologies (AICT)*, pages 1–6. IEEE.
- Alnajjar, K., Hämäläinen, M., ve Rueter, J. (2023). Sentiment analysis using aligned word embeddings for uralic languages. *arXiv preprint arXiv:2305.15380*.
- Altinel Girgin, A. B., Gümüşçekiççi, G., ve Birdemir, N. C. (2024). Turkish sentiment analysis: A comprehensive review. *Sigma Journal of Engineering and Natural Sciences*, 42(4):1292–1314.
- An, Z., Wan, H., Xiong, T., ve Wang, B. (2025). Lepb-sa4re: A lexicon-enhanced and prompt-tuning bert model for evolving requirements elicitation from app reviews. *Applied Sciences*, 15(5):2282.
- Araci, D. T. (2019). FinBERT: Financial sentiment analysis with pre-trained language models. *arXiv preprint arXiv:1908.10063*.
- Aydın, C. R., Güngör, T., ve Erkan, A. (2020). Generating word and document embeddings for sentiment analysis. *arXiv preprint*.
- Aykaç, Y. E., Özkan, M., ve Samet, R. (2026a). Domain-robust, polarity-injected sentence embeddings for azerbaijani sentiment analysis. Under review.
- Aykaç, Y. E., Özkan, M., ve Samet, R. (2026b). Lexicon-augmented explainable aspect-based sentiment analysis for azerbaijani finance and business text. In *Proceedings of the 3rd International Conference on Information Technologies and Their Applications (ITTA 2026)*, Baku, Azerbaijan. Accepted, to appear.
- Aykaç, Y. E., Samet, R., Pashayev, A., ve Sabziev, E. (2025). Sentiaznet: A polarity lexicon for azerbaijani. In *2025 10th International Conference on Computer Science and Engineering (UBMK)*, pages 1684–1689. IEEE.
- Aykaç, Y. E., Ozkan, M., ve Samet, R. (2026). Morpheme-aware hybrid sentiment analysis for azerbaijani: A lexicon–embedding fusion with domain adaptation. *IEEE Access*.
- Baccianella, S., Esuli, A., Sebastiani, F., vd. (2010). Sentiwordnet 3.0: an enhanced lexical resource for sentiment analysis and opinion mining. In *Lrec*, volume 10, pages 2200–2204. Valletta.
- Badr, H., Wanas, N., ve Fayek, M. (2024). Unsupervised domain adaptation via weighted sequential discriminative feature learning for sentiment analysis. *Applied Sciences*, 14(1):406.
- Bai, J., Bai, S., Chu, Y., Cui, Z., Dang, K., Deng, X., Fan, Y., Ge, W., Han, Y., Huang, F., vd.

- (2023). Qwen technical report. *arXiv preprint arXiv:2309.16609*.
- Baldwin, T., De Marneffe, M.-C., Han, B., Kim, Y.-B., Ritter, A., ve Xu, W. (2015). Shared tasks of the 2015 workshop on noisy user-generated text: Twitter lexical normalization and named entity recognition. In *Proceedings of the workshop on noisy user-generated text*, pages 126–135.
- Barbieri, F., Camacho-Collados, J., Anke, L. E., ve Neves, L. (2020). Tweeteval: Unified benchmark and comparative evaluation for tweet classification. In *Findings of the association for computational linguistics: EMNLP 2020*, pages 1644–1650.
- Barnes, J., Klinger, R., ve Im Walde, S. S. (2018). Bilingual sentiment embeddings: Joint projection of sentiment across languages. In *Proceedings of the 56th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 2483–2493.
- Birjali, M., Kasri, M., ve Beni-Hssane, A. (2021). A comprehensive survey on sentiment analysis: Approaches, challenges and trends. *Knowledge-Based Systems*, 226:107134.
- Bojanowski, P., Grave, E., Joulin, A., ve Mikolov, T. (2017). Enriching word vectors with subword information. *Transactions of the association for computational linguistics*, 5:135–146.
- Çakıcı, Ş., Karaduman, D., Çırlan, M. A., ve Hürriyetoğlu, A. (2025). A cross-validation study of turkish sentiment analysis datasets and tools. *Language Resources and Evaluation*, pages 1–39.
- Chrysostomou, G. ve Aletras, N. (2021). Improving the faithfulness of attention-based explanations with task-specific information for text classification. In *Proceedings of the 59th Annual Meeting of the Association for Computational Linguistics and the 11th International Joint Conference on Natural Language Processing (Volume 1: Long Papers)*, pages 477–488, Online. Association for Computational Linguistics.
- Çöltekin, Ç. (2020). A corpus of turkish offensive language on social media. In *Proceedings of the twelfth language resources and evaluation conference*, pages 6174–6184.
- Conneau, A., Khandelwal, K., Goyal, N., Chaudhary, V., Wenzek, G., Guzmán, F., Grave, E., Ott, M., Zettlemoyer, L., ve Stoyanov, V. (2020). Unsupervised cross-lingual representation learning at scale. In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics (ACL)*, pages 8440–8451. Association for Computational Linguistics.
- Dehkharghani, R., Yanikoglu, B., Saygin, Y., ve Oflazer, K. (2017). Sentiment analysis in turkish at different granularity levels. *Natural Language Engineering*, 23(4):535–559.
- Demirtas, E. ve Pechenizkiy, M. (2013). Cross-lingual polarity detection with machine translation. In *Proceedings of the Second International Workshop on Issues of Sentiment Discovery and Opinion Mining*, pages 1–8.
- Dettmers, T., Pagnoni, A., Holtzman, A., ve Zettlemoyer, L. (2023). Qlora: Efficient finetuning of quantized llms. *Advances in neural information processing systems*, 36:10088–10115.
- Devlin, J., Chang, M., Lee, K., ve Toutanova, K. (2019). BERT: Pre-training of deep bidirectional transformers for language understanding. In *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (NAACL-HLT)*, pages 4171–4186. Association for Computational Linguistics.
- Dhall, S., Kumar, S., ve Kumar, S. (2025). A review on sentiment analysis in low-resource languages focusing on fake news and sarcasm detection as major challenges. *SN Computer Science*, 6(6):693.

- Diwali, A., Saeedi, K., Dashtipour, K., Gogate, M., Cambria, E., ve Hussain, A. (2024). Sentiment analysis meets explainable artificial intelligence: A survey on explainable sentiment analysis. *IEEE Transactions on Affective Computing*, 15(3):837–846.
- D’aniello, G., Gaeta, M., ve La Rocca, I. (2022). Knowmis-absa: an overview and a reference model for applications of sentiment analysis and aspect-based sentiment analysis. *Artificial Intelligence Review*, 55(7):5543–5574.
- Erşahin, B., Aktaş, Ö., Kiliç, D., ve Erşahin, M. (2019). A hybrid sentiment analysis method for turkish. *Turkish Journal of Electrical Engineering and Computer Sciences*, 27(3):1780–1793.
- Fang, F., Dutta, K., ve Datta, A. (2014). Domain adaptation for sentiment classification in light of multiple sources. *INFORMS Journal on Computing*, 26(3):586–598.
- Fang, Y., Yap, P.-T., Lin, W., Zhu, H., ve Liu, M. (2024). Source-free unsupervised domain adaptation: A survey. *Neural Networks*, 174:106230.
- Felbo, B., Mislove, A., Søggaard, A., Rahwan, I., ve Lehmann, S. (2017). Using millions of emoji occurrences to learn any-domain representations for detecting sentiment, emotion and sarcasm. In *Proceedings of the 2017 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, pages 1615–1625.
- Feng, F., Yang, Y., Cer, D., Arivazhagan, N., ve Wang, W. (2022). Language-agnostic bert sentence embedding. In *Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 878–891.
- Fleiss, J. L. (1971). Measuring nominal scale agreement among many raters. *Psychological Bulletin*, 76(5):378–382.
- Ganin, Y. ve Lempitsky, V. (2015). Unsupervised domain adaptation by backpropagation. In *International conference on machine learning*, pages 1180–1189. PMLR.
- Ganin, Y., Ustinova, E., Ajakan, H., Germain, P., Larochelle, H., Laviolette, F., Marchand, M., ve Lempitsky, V. (2016). Domain-adversarial training of neural networks. *Journal of Machine Learning Research*, 17(59):1–35.
- Gao, T., Yao, X., ve Chen, D. (2021). Simcse: Simple contrastive learning of sentence embeddings. In *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, pages 6894–6910.
- Guliyev, N., Rustamov, S., ve Rustamov, Z. (2024). Analysis of public sentiment in azerbaijani news and social media. In *Proceedings of the 2024 International Conference on Computing, Machine Learning and Data Science (CMLDS ’24)*, pages 22:1–22:6. Association for Computing Machinery.
- Gururangan, S., Marasović, A., Swayamdipta, S., Lo, K., Beltagy, I., Downey, D., ve Smith, N. A. (2020). Don’t stop pretraining: Adapt language models to domains and tasks. In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics (ACL)*, pages 8342–8360.
- Haddow, B., Bawden, R., Miceli-Barone, A. V., Helcl, J., ve Birch, A. (2022). Survey of low-resource machine translation. *Computational Linguistics*, 48(3):673–732.
- Hamilton, W. L., Clark, K., Leskovec, J., ve Jurafsky, D. (2016). Inducing domain-specific sentiment lexicons from unlabeled corpora. In *Proceedings of the 2016 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, pages 595–605.

- Hu, E. J., Shen, Y., Wallis, P., Allen-Zhu, Z., Li, Y., Wang, S., Wang, L., Chen, W., vd. (2022). Lora: Low-rank adaptation of large language models. *ICLR*, 1(2):3.
- Hua, Y. C., Denny, P., Wicker, J., ve Taskova, K. (2024). A systematic review of aspect-based sentiment analysis: domains, methods, and trends. *Artificial Intelligence Review*, 57:296.
- Hutto, C. ve Gilbert, E. (2014). Vader: A parsimonious rule-based model for sentiment analysis of social media text. In *Proceedings of the international AAAI conference on web and social media*, volume 8, pages 216–225.
- Isbarov, J., Huseynova, K., Mammadov, E., Hajili, M., ve Ataman, D. (2024). Open foundation models for Azerbaijani language. In *Proceedings of the First Workshop on Natural Language Processing for Turkic Languages (SIGTURK 2024)*, pages 18–28, Bangkok, Thailand and Online. Association for Computational Linguistics.
- Jain, S. ve Wallace, B. C. (2019). Attention is not explanation. In *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (Volume 1: Long and Short Papers)*, pages 3543–3556. Association for Computational Linguistics.
- Johnson, M., Schuster, M., Le, Q., Krikun, M., Wu, Y., Chen, Z., Thorat, N., Viégas, F., Wattenberg, M., Corrado, G., vd. (2017). Google’s multilingual neural machine translation system: Enabling zero-shot translation. *Transactions of the Association for Computational Linguistics*, 5:339–351.
- Joshi, P., Santy, S., Budhiraja, A., Bali, K., ve Choudhury, M. (2020). The state and fate of linguistic diversity and inclusion in the nlp world. In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics (ACL)*, pages 6282–6293.
- Joulin, A., Grave, E., Bojanowski, P., Douze, M., Jégou, H., ve Mikolov, T. (2016). Fasttext. zip: Compressing text classification models. *arXiv preprint arXiv:1612.03651*.
- Joulin, A., Grave, E., Bojanowski, P., ve Mikolov, T. (2017a). Bag of tricks for efficient text classification. In *Proceedings of the 15th Conference of the European Chapter of the Association for Computational Linguistics (EACL)*, volume 2, pages 427–431.
- Joulin, A., Grave, E., Bojanowski, P., ve Mikolov, T. (2017b). Bag of tricks for efficient text classification. In *Proceedings of the 15th conference of the European chapter of the association for computational linguistics: volume 2, short papers*, pages 427–431.
- Keung, P., Lu, Y., Szarvas, G., ve Smith, N. A. (2020). The multilingual amazon reviews corpus. *arXiv preprint arXiv:2010.02573*.
- Khan, L., Amjad, A., Ashraf, N., ve Chang, H.-T. (2022). Multi-class sentiment analysis of urdu text using multilingual bert. *Scientific Reports*, 12(1):5436.
- Khosla, P., Teterwak, P., Wang, C., Sarna, A., Tian, Y., Isola, P., Maschinot, A., Liu, C., ve Krishnan, D. (2020). Supervised contrastive learning. In *Advances in Neural Information Processing Systems (NeurIPS)*.
- Köksal, A. ve Özgür, A. (2021). Twitter dataset and evaluation of transformers for turkish sentiment analysis. In *29th Signal Processing and Communications Applications Conference (SIU) 2021*, pages 1–4, Istanbul, Turkey. IEEE.
- Kreuzer, C., Sparrer, C., ve Dorfleitner, G. (2026). Beyond pure hype: News sentiment and its role in the BTC and ETH futures market. *Review of Derivatives Research*, 29(1):1–36.
- Kuriyozov, E., Matlatipov, S., Alonso, M. A., ve Gómez-Rodríguez, C. (2022). Construction

- and evaluation of sentiment datasets for low-resource languages: The case of Uzbek. In *Human Language Technology. Challenges for Computer Science and Linguistics*, volume 13212 of *Lecture Notes in Computer Science*, pages 232–243. Springer, Cham. Language and Technology Conference (LTC 2019).
- Lamin, N. Z. ve Aziz, A. A. (2025). Cross-lingual sentiment analysis in low-resource languages: A recent review on tasks, methods and challenges. *International Journal of Advanced Computer Science & Applications*, 16(11).
- Li, M., Gao, Q., ve Yu, T. (2023). Kappa statistic considerations in evaluating inter-rater reliability between two raters: which, when and context matters. *BMC Cancer*, 23:799.
- Ligthart, A., Catal, C., ve Tekinerdogan, B. (2021). Systematic reviews in sentiment analysis: a tertiary study. *Artificial intelligence review*, 54(7):4997–5053.
- Liu, B. (2012). *Sentiment Analysis and Opinion Mining*. Synthesis Lectures on Human Language Technologies. Morgan & Claypool Publishers.
- Liu, P., Qiu, X., ve Huang, X.-J. (2017). Adversarial multi-task learning for text classification. In *Proceedings of the 55th annual meeting of the association for computational linguistics (volume 1: long papers)*, pages 1–10.
- Liu, Y., Ott, M., Goyal, N., Du, J., Joshi, M., Chen, D., Levy, O., Lewis, M., Zettlemoyer, L., ve Stoyanov, V. (2019). RoBERTa: A robustly optimized BERT pretraining approach. *arXiv preprint arXiv:1907.11692*.
- Loughran, T. ve McDonald, B. (2011). When is a liability not a liability? textual analysis, dictionaries, and 10-ks. *The Journal of Finance*, 66(1):35–65.
- Lu, Q., Sun, X., Long, Y., Gao, Z., Feng, J., ve Sun, T. (2023). Sentiment analysis: Comprehensive reviews, recent advances, and open challenges. *IEEE Transactions on Neural Networks and Learning Systems*.
- Lundberg, S. M. ve Lee, S.-I. (2017). A unified approach to interpreting model predictions. In *Advances in Neural Information Processing Systems*, volume 30, pages 4765–4774.
- Maas, A. L., Daly, R. E., Pham, P. T., Huang, D., Ng, A. Y., ve Potts, C. (2011). Learning word vectors for sentiment analysis. In *Proceedings of the 49th Annual Meeting of the Association for Computational Linguistics: Human Language Technologies (ACL-HLT)*, pages 142–150.
- Mabokela, K. R., Celik, T., ve Raborife, M. (2022). Multilingual sentiment analysis for under-resourced languages: a systematic review of the landscape. *IEEE Access*, 11:15996–16020.
- Mabokela, K. R., Primus, M., ve Celik, T. (2024). Explainable pre-trained language models for sentiment analysis in low-resourced languages. *Big Data and Cognitive Computing*, 8(11):160.
- McNemar, Q. (1947). Note on the sampling error of the difference between correlated proportions or percentages. *Psychometrika*, 12(2):153–157.
- Mohammad, S. M. ve Turney, P. D. (2013). Crowdsourcing a word–emotion association lexicon. *Computational Intelligence*, 29(3):436–465.
- Mughal, N., Mujtaba, G., vd. (2024). Comparative analysis of deep natural networks and large language models for aspect-based sentiment analysis. *IEEE Access*.
- Mutinda, J., Mwangi, W., ve Okeyo, G. (2023). Sentiment analysis of text reviews using lexicon-enhanced bert embedding (lebert) model with convolutional neural network. *Applied Sciences*, 13(3):1445.

- Mutlu, M. M. ve Özgür, A. (2022). A dataset and bert-based models for targeted sentiment analysis on turkish texts. In *Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics: Student Research Workshop*, pages 467–472.
- Najafi, A. ve Varol, O. (2024). Turkishbertweet: Fast and reliable large language model for social media analysis. *Expert Systems with Applications*, 255:124737.
- Nasiopoulos, D. K., Roumeliotis, K. I., Sakas, D. P., Toudas, K., ve Reklitis, P. (2025). Financial sentiment analysis and classification: A comparative study of fine-tuned deep learning models. *International Journal of Financial Studies*, 13(2):75.
- Nath, D. ve Dwivedi, S. K. (2024). Aspect-based sentiment analysis: approaches, applications, challenges and trends. *Knowledge and Information Systems*, 66(12):7261–7303.
- Nwaiwu, S. (2025). Parameter-efficient fine-tuning for low-resource text classification: a comparative study of lora, ia 3, and reft. *Frontiers in Big Data*, 8:1677331.
- Pang, B., Lee, L., ve Vaithyanathan, S. (2002). Thumbs up? sentiment classification using machine learning techniques. In *Proceedings of the 2002 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, pages 79–86.
- Perikos, I. ve Diamantopoulos, A. (2024). Explainable aspect-based sentiment analysis using transformer models. *Big Data and Cognitive Computing*, 8(11):141.
- Pires, T., Schlinger, E., ve Garrette, D. (2019). How multilingual is multilingual bert? *arXiv preprint arXiv:1906.01502*.
- Pokhriyal, H. ve Jain, G. (2025a). Sarcasm detection with induced sentimental cues using heuristic search based on unconstrained optimisation learning quantifying callousness on social media. *Neurocomputing*, 645:130499.
- Pokhriyal, H. ve Jain, G. (2025b). Supposititious sarcasm detection and sentiment analysis coping hindi language in social networks harnessing zipf–mandelbrot probabilistic optimisation and perplexity entropy learning. *ACM Transactions on Asian and Low-Resource Language Information Processing*, 24(2):17:1–17:28.
- Pontiki, M., Galanis, D., Papageorgiou, H., Androutsopoulos, I., Manandhar, S., Al-Smadi, M., Al-Ayyoub, M., Zhao, Y., Qin, B., De Clercq, O., vd. (2016). Semeval-2016 task 5: Aspect based sentiment analysis. In *Proceedings of the 10th international workshop on semantic evaluation (SemEval-2016)*, pages 19–30.
- Ramponi, A. ve Plank, B. (2020). Neural unsupervised domain adaptation in nlp—a survey. In *Proceedings of the 28th International Conference on Computational Linguistics (COLING)*, pages 6872–6882.
- Reimers, N. ve Gurevych, I. (2019). Sentence-bert: Sentence embeddings using siamese bert-networks. In *Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP)*, pages 3982–3992.
- Ribeiro, M. T., Singh, S., ve Guestrin, C. (2016). “why should i trust you?” explaining the predictions of any classifier. In *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, pages 1135–1144. Association for Computing Machinery.
- Rust, P., Pfeiffer, J., Vulić, I., Ruder, S., ve Gurevych, I. (2021). How good is your tokenizer? on the monolingual performance of multilingual language models. In *Proceedings of the 59th Annual Meeting of the Association for Computational Linguistics and the 11th International*

Joint Conference on Natural Language Processing (Volume 1: Long Papers), pages 3118–3135.

Schweter, S. (2020). BERTurk – BERT models for Turkish.

Senel, L. K., Ebing, B., Baghirova, K., Schuetze, H., ve Glavaš, G. (2024). Kardeş-NLU: Transfer to low-resource languages with the help of a high-resource cousin – a benchmark and evaluation for Turkic languages. In *Proceedings of the 18th Conference of the European Chapter of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 1672–1688, St. Julian's, Malta. Association for Computational Linguistics.

Shaik, T., Tao, X., Dann, C., Xie, H., Li, Y., ve Galligan, L. (2023). Sentiment analysis and opinion mining on educational data: A survey. *Natural Language Processing Journal*, 2:100003.

Sheetal, R. ve Aithal, P. K. (2025). Enhancing financial sentiment analysis: Integrating the loughran-mcdonald dictionary with BERT for advanced market predictive insights. *Procedia Computer Science*, 258:2244–2257.

Šmíd, J. ve Král, P. (2025). Cross-lingual aspect-based sentiment analysis: A survey on tasks, approaches, and challenges. *Information Fusion*, 120:103073.

Šmíd, J., Přibáň, P., ve Král, P. (2025). LACA: Improving cross-lingual aspect-based sentiment analysis with LLM data augmentation. In *Proceedings of the 63rd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 839–853, Vienna, Austria. Association for Computational Linguistics.

Suhartono, D., Wongso, W., ve Tri Handoyo, A. (2024). Id sarcasm: Benchmarking and evaluating language models for indonesian sarcasm detection. *IEEE Access*, 12:87323–87332.

Taboada, M., Brooke, J., Tofiloski, M., Voll, K., ve Stede, M. (2011). Lexicon-based methods for sentiment analysis. *Computational linguistics*, 37(2):267–307.

Veitsman, Y. ve Hartmann, M. (2025). Recent advancements and challenges of Turkic Central Asian language processing. In *Proceedings of the First Workshop on Language Models for Low-Resource Languages*, pages 309–324, Abu Dhabi, United Arab Emirates. Association for Computational Linguistics.

Wang, B., Lapata, M., ve Titov, I. (2021). Meta-learning for domain generalization in semantic parsing. In *Proceedings of the 2021 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, pages 366–379.

Wankhade, M., Rao, A. C. S., ve Kulkarni, C. (2022). A survey on sentiment analysis methods, applications, and challenges. *Artificial Intelligence Review*, 55(7):5731–5780.

Wu, C., Lin, Z., Liu, Y., He, C., Yu, D., Yamada, I., Zhou, J., ve Miyao, Y. (2025a). M-ABSA: A multilingual dataset for aspect-based sentiment analysis. In *Proceedings of the 2025 Conference on Empirical Methods in Natural Language Processing*, pages 2525–2541, Miami, Florida, USA. Association for Computational Linguistics.

Wu, J., Qiu, P., ve Jia, X. (2025b). Constructing a domain-specific sentiment lexicon for agricultural product reviews using bert and so-pmi. *PloS one*, 20(6):e0326602.

Xu, G., Meng, Y., Qiu, X., Yu, Z., ve Wu, X. (2019). Sentiment analysis of comment texts based on bilstm. *Ieee Access*, 7:51522–51532.

Xu, Y. ve Ibrahim, N. F. (2022). Cross-domain aspect-based sentiment analysis for enhancing customer experience in electronic commerce. *advances in artificial intelligence and machine*

- learning. 2024; 4 (3): 151. In *Conference on Computer Engineering and Applications (CVIDL & ICCEA)*. IEEE, volume 1, page 2.
- Yang, L., Li, Y., Wang, J., ve Sherratt, R. S. (2020). Sentiment analysis for e-commerce product reviews in chinese based on sentiment lexicon and deep learning. *IEEE access*, 8:23522–23530.
- Yeshpanov, R. ve Varol, H. A. (2024). KazSAnDRA: Kazakh sentiment analysis dataset of reviews and attitudes. In *Proceedings of the 2024 Joint International Conference on Computational Linguistics, Language Resources and Evaluation (LREC-COLING 2024)*, pages 9657–9667, Torino, Italia. ELRA and ICCL.
- Yildirim, S. (2024). Fine-tuning transformer-based encoder for turkish language understanding tasks. *arXiv preprint arXiv:2401.17396*.
- Yue, T., Mao, R., Wang, H., ve Hu, Z. (2023). Knowlenet: Knowledge fusion network for multimodal sarcasm detection. *Information Fusion*, 100:101921.
- Zeybek, S., Koç, E., ve Seçer, A. (2021). Ms-tr: A morphologically enriched sentiment treebank and recursive deep models for compositional semantics in turkish. *Cogent Engineering*, 8(1):1893621.
- Zhang, L., Wang, S., ve Liu, B. (2018). Deep learning for sentiment analysis: A survey. *Wiley interdisciplinary reviews: data mining and knowledge discovery*, 8(4):e1253.
- Zhang, W., Deng, Y., Liu, B., Pan, S., ve Bing, L. (2024). Sentiment analysis in the era of large language models: A reality check. In *Findings of the Association for Computational Linguistics: NAACL 2024*, pages 3881–3906, Mexico City, Mexico. Association for Computational Linguistics.
- Zhao, Z., Ma, Z., Lin, Z., Xie, J., Li, Y., ve Shen, Y. (2024). Source-free domain adaptation for aspect-based sentiment analysis. In *Proceedings of the 2024 Joint International Conference on Computational Linguistics, Language Resources and Evaluation (LREC-COLING 2024)*, pages 15076–15086.
- Zhou, J., Tian, J., Wang, R., Wu, Y., Xiao, W., ve He, L. (2020). Sentix: A sentiment-aware pre-trained model for cross-domain sentiment analysis. In *Proceedings of the 28th international conference on computational linguistics*, pages 568–579.
- Zhu, Z., Chen, H., Ye, X., Lyu, Q., Tan, C., Marasović, A., ve Wiegreffe, S. (2024). Explanation in the era of large language models. In *Proceedings of the 2024 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (Volume 5: Tutorial Abstracts)*, pages 19–25.