

Türkçe Resmi Olmayan Metinlerde İroni Tespiti için Sinirsel Yöntemlerin İncelenmesi

Investigating the Neural Models for Irony Detection on Turkish Informal Texts

Yeşim Cemek*, Cenk Cidecio*, Aslı Umay Öztürk*, Recep Fırat Çekinel*, Pınar Karagöz*

*Orta Doğu Teknik Üniversitesi (ODTÜ), Bilgisayar Mühendisliği Bölümü, Ankara,
yesim.cemek@metu.edu.tr, cidecio.cenk@metu.edu.tr, ozturk.asli@metu.edu.tr, rfcekinel@ceng.metu.edu.tr,
karagoz@ceng.metu.edu.tr

Özetçe —İroni, mizah, alaycılık ya da vurgu amaçlı olarak söylenen, sözün tersini kasteden ifade olarak tanımlanmaktadır. Metinde ironi tespiti problemini, duygu tespiti probleminin özel bir durumu olarak düşünebiliriz. Ancak ironi tespiti, yapısı gereği, çoğu duygu tespiti probleminden daha zorlu bir problemdir. Yazılı ironinin modellenmesi ve tespiti başlı başına önem taşımının yanı sıra, birçok metin madenciliği işleminin doğruluğunu artırılması için de fayda sağlamaktadır. Bu çalışmada, ironi tespiti problemini Türkçe resmi olmayan metinler üzerinde incelemekteyiz. Literatürde çoğunlukla İngilizce metinler için önerilen çeşitli çözümler bulunmaktadır. Ancak Türkçe metinler üzerinde yapılan çalışma sayısı oldukça sınırlıdır. Resmi olmayan metinler üzerinde çalışmamızın nedeni, ironi ifadelerinin blog, sosyal medya gönderileri gibi metinlerde, gazete ve kitaplara kıyasla daha sık görülmesidir. Naif Bayes Sınıflandırıcı ya da Destek Vektör Makineleri gibi geleneksel denetimli öğrenme tabanlı çözümlerin performansı Türkçe metinlerde ironi tespiti için daha önce incelenmişti. Bu çalışmada, sinirsel ağ tabanlı modellerin performansını geleneksel gözetimli öğrenme yöntemleriyle karşılaştırmalı olarak incelemekteyiz.

Anahtar Kelimeler—ironi, hiciv, gözetimli öğrenme, sinirsel model.

Abstract—Irony is defined as the expression of one's meaning by using language that normally signifies the opposite, typically for humorous or emphatic effect. In the textual context, it can be considered as a specific type of opinion mining problem. However, due to the nature of the problem, it is generally more challenging to detect. Irony detection is useful to understand the semantics of a text. It also helps improving the accuracy of many other text mining tasks. In this paper, we study the irony detection problem on Turkish informal texts. In the literature there are several solutions proposed mostly for English texts. However, the number of studies on Turkish texts is very limited. The reason for using informal text is that ironic expressions are seen more frequently in informal communication media such as blogs, social media posts. Previously, performance of conventional supervised learning based solutions, such as Naive Bayes Classifier and SVM, on Turkish texts have been studied. In this work, we investigate the performance of neural network based models in comparison to conventional supervised learning models.

Keywords—irony, sarcasm, supervised learning, neural model.

I. GİRİŞ

Günümüzde sosyal medya ve web ortamında kullanıcılar tarafından üretilen metin verisi oldukça hızlı bir şekilde artmaktadır. Bu içerik üzerinde metin analizi ile bilgi çıkarımı pek çok uygulama için girdi sağlamaktadır. Örneğin ürün yorumları, ürün ve kullanıcılar hakkında çıkarımlar için önemli bir kaynaktır. Benzer şekilde sosyal medya mesajlarının analizi, kullanıcılara yapılan başka hesap veya kullanıcı takip önerilerinin ya da kullanıcı profillerinin belirlenmesinin başarısını artırmaktadır.

Metin analizi ile bilgi çıkarımı problemlerinden biri de *ironi tespiti*dir. *İroni* kelimesi, Türk Dil Kurumu sözlüğünde¹ "Söylenen sözün tersini kastederek kişiyle veya olayla alay etme" şeklinde tanımlanmaktadır. Çeşitli kaynaklarda, *ironi*, mizah ya da vurgu amaçlı olarak söylenen sözün tersini kastetmek olarak da ifade edilmektedir. İroni farklı amaçlara yönelik olarak yapılabilmekle birlikte, tanımlardaki ortak vurgu sözün görünen anlamı ile kastedilen arasındaki zıtlıktır. Metinde ironi tespitini, duygu analizi probleminin oldukça özel bir alt problemi olarak düşünebiliriz. İroni tespiti, metni otomatik olarak anlamlandırmak için başlı başına öneme sahip olmakla birlikte, metinden yazar stil ve kişiliğinin belirlenmesi, ve özellikle de metin üzerindeki duygu analizi doğruluğunu iyileştirilmesi için de fayda sağlamaktadır.

Literatüre baktığımızda metinde ironi tespiti çalışmalarının, diğer metin analizi çalışmalarında olduğu gibi İngilizce metinler üzerine yoğunlaştığı görülmektedir [2], [4], [9], [13]. Problem, temel olarak bir sınıflandırma problemi olarak tanımlanmakta ve gözetimli öğrenme (supervised learning) teknikleri ile çözüm sunulmaktadır [8], [9]. Yöntem açısından, sunulan ilk çalışmalarda Destek Vektör Makineleri, Bayes Sınıflandırıcılar gibi geleneksel olarak tabir edilebilecek tekniklerin kullanıldığı görülmektedir. Yakın tarihli çalışmalarda, yapay zeka çalışmalarındaki genel eğilimle uyumlu olarak, ironi tespiti için de sinirsel modeller öne çıkmaktadır.

¹Çalışmada ilk üç araştırmacı eşit katkı sağlamıştır.

Bu çalışma TÜBİTAK tarafından 117E566 no'lu proje kapsamında desteklenmiştir.

¹www.sozluk.gov.tr

İngilizce üzerinde yapılan çalışmaların aksine, Türkçe metinler üzerinde yapılmış olan çalışmalar oldukça sınırlı sayıdadır [6], [12]. Bu çalışmalarda da, ironi tespiti için sınıflandırma problemi olarak tanımlanmış, geleneksel gözetimli öğrenme tekniklerinin, farklı özniteliklerin kullanımı durumunda tespit performansı incelenmiştir. Bildiğimiz kadarıyla sinirsel modellerin Türkçe için ironi tespiti henüz incelenmemiştir. Bu çalışmamızda, metin analizi altında duygu analizi için başarılı sonuçlar verdiği raporlanmış olan *Uzun-Kısa Süreli Bellek* (Long Short Term Memory (LSTM)) sinir ağı mimarisi ve Google AI tarafından doğal dil işleme problemlerine yönelik olarak geliştirilmiş olan sinir ağı mimarisi BERT'in (Bidirectional Encoder Representations from Transformers (Dönüştürücülerden Çift Yönlü Kodlayıcı Temsilleri)) [5] Türkçe metinlerde ironi tespiti performansını inceledik. Ironi tespitini, bu tür ifadelerin daha yoğun olarak yer aldığı sosyal medya mesajları üzerinde uyguladık. Yöntemlerimizin doğruluk performansı analizleri, kullanıcı çalışması ile etiketlenmiş dengeli bir veri kümesi üzerinde yapılarak geleneksel gözetimli öğrenme teknikleriyle karşılaştırmalı olarak sunulmaktadır.

II. LİTERATÜR ÖZETİ

Literatürde, metinlerde ironi tespiti üzerine görece az sayıda çalışma bulunmaktadır. Birçok çalışma, tweetler gibi kısa yazıların aksine düzenli metinler üzerine yapılmıştır. Bunların yanında, daha az sayıda, tweetler ve sosyal medya mesajları üzerinde çalışmalar da bulunmaktadır. Cümle uzunluğu üzerinde sınır olması, düz yazılara oranla çok daha fazla gramer hatası bulundurması, kullanıcıların platforma özel ifadeleri sıklıkla kullanıyor olmaları, bu çalışmaları daha farklı kılmaktadır.

Joshi ve diğerleri [9] çalışmalarında, geleneksel yöntemlerdeki özniteliklerin yanı sıra kelime temsili (word-embedding) bazlı öznitelikler, cümledeki kelime çiftlerinin yakınlık değerleri, kelimelerin cümledeki yerlerine göre uzaklık değerlerini de kullanmışlardır. Bunlara ek olarak, bağlamı değerlendirmeye katan çalışmalar da mevcuttur. Örneğin Bamman ve Smith [3], yazar, okur ve bu aktörlerin paylaştıkları iletişim bağlamı hakkında bilgileri kullanmışlardır. Wallace ve diğerleri [13] ise, Reddit'teki ² iki zıt görüşün Subreddit'ine ait metinleri kullanarak, ima edilen duygu ve metinlerin ait oldukları Subreddit'i de öznitelikler olarak kullanmışlardır.

Ironi tespiti problemi, ironik metinleri sözel ironi ve durumsal ironi gibi alt sınıflara ayırmak için detaylandırılabilir. Literatürde buna örnek olan bir çalışmada Van Hee ve diğerleri [8], ironi tespiti için hiyerarşik bir sınıflandırma yaklaşımı kullanmıştır. Ironik metinler öncelikle *kutupsal çakışma kaynaklı ironi* (cümlede kutupsallık mevcutsa) ve *diğer ironi* olmak üzere ikiye ayrılmaktadır. Daha sonra *diğer ironi* altbaşlığın-dakiler, *durumsal ironi* ve *diğer sözlü ironi* olmak üzere tekrar iki sınıfa ayrılır. Sözdizimsel (syntactic), sözcüksel (lexical), duygusal (sentiment lexicon) ve anlamsal (semantic) olmak üzere gruplanan öznitelikler, ayrı ayrı incelendiği gibi bir arada da analiz edilmiştir. Çalışma sonucunda en iyi sonucu tüm özniteliklerin birlikte incelendiği modelin verdiği belirtilmiştir. Pamungkas ve Patti [11] ise, metinlerde kullanılan *emoji*leri ve ifadeleri de değerlendirmeye katmışlardır. Diğer özniteliklere bu emoji ve ifadelerin duygusal puanlarını da ekledikleri

zaman ironi tespiti yapan modellerinin başarısının yükseldiği belirtilmiştir.

Barbieri ve diğerleri [4] probleme Hiciv, Eğitim, Mizah, İroni, Politika ve Gazete makalesi olmak üzere altı farklı sınıf kullanarak, ikili sınıflandırma problemi olarak yaklaşmışlardır. Model oluşturmak için karar ağacı (decision tree) yöntemini kullanmışlardır. Ironi tespitinde kelime örüntülerinin öznitelik olarak kullanımı yerine, kelime sıklıkları ve bunlar arasındaki uzaklık, zarf ve sıfatların yoğunluğu, yaygın ve nadir eşanlamlı sözcüklerin kullanımı gibi özelliklere odaklanmışlardır. Ahmed ve diğerleri [1] ise tweet'lerden oluşan veri kümesindeki sınıfları, yakınlık penceresi (vicinity window) adı verilen bir metrik kullanarak yönlü değersiz çizgeler (directed unweighted graph) olarak temsil etmişlerdir. Örneklem ile veri kümesindeki sınıflar arasındaki çizge benzerliklerini hesaplamak için çeşitli çizge karşılaştırma yöntemleri kullanarak, elde ettikleri değerleri özniteliklere eklemişlerdir. Zhang ve diğerleri [14], ironi tespiti için transfer öğrenimi kullanarak ironi etiketli metin üzerinde gözetimli öğrenme ile dış kaynaklardan elde edilen duygu bilgisini (sentiment information) birleştirmişlerdir. Sunulan sonuçlar, duygu bilgisinin kullanılmasının modellerin örtük ve açık bağlam uyumsuzluğunu tespit etme yeteneğini artırdığını göstermektedir. Ghanem ve diğerleri [7] ise ironi tespiti problemi için, Fransızca, İngilizce ve Arapça dillerinde transfer öğrenme tabanlı bir yaklaşımın başarısını incelemişlerdir.

Ironi tespiti için Türkçe metinler üzerinde yapılan çalışmalar oldukça sınırlıdır. Taslioglu ve Karagoz [12] tarafından yapılan çalışmada Türkçe sosyal ağ mesajları üzerinde Destek Vektör Makineleri (SVM), En yakın K Komşu (K-NN), Naif Bayes Sınıflandırıcı ve Rastgele Orman (Random Forest) karar ağacı yöntemleri kullanılarak farklı öznitelik kümelerinin ironi tespitine etkisi incelenmiştir. Dulger'in [6] sunduğu çalışmada ise düzenli Türkçe metinler üzerinde ironi tespitine odaklanılmıştır. Metin içeriklerinden elde edilen duygu polarite puanlarının yanı sıra tırnak işareti, büyük harf kullanımı gibi desen tabanlı özniteliklerin de yer aldığı 13 öznitelik üzerinde inceleme yapılmıştır.

Sunduğumuz çalışmada [12]'deki yaklaşıma benzer şekilde sosyal medya mesajları analiz edilmektedir. Önce çalışmalardan farklı olarak, İngilizce metinler için bazı çalışmalarda kullanılmış ama daha önce Türkçe metinler üzerinde uygulanmamış olan sinirsel modellerin performansının incelenmesini odaklanmaktayız.

III. YÖNTEMLER

A. Sinirsel Model Tabanlı Yöntemler

Sinirsel model tabanlı bir yöntem olan Uzun-Kısa Süreli Hafıza ağı (LSTM) bir tekrarlayan sinir ağı (RNN) mimarisidir. Standart yapay sinir ağları mimarilerinden farklı olarak geri bildirim bağlantıları içerir ve tekrarlayan sinir ağlarına uzun süreli bellek sağlar. Bu özelliğiyle zaman serisi (time series) verileri üzerinde çalışmak için idealdir. İki Yönlü Uzun-Kısa Süreli Hafıza ağları (Bi-LSTM) ise verilen girdileri geçmişten geleceğe ve gelecekte geçmişe olacak şekilde iki kere çalıştırır. Bu şekilde Uzun-Kısa Süreli Hafıza ağından farklı olarak sadece geçmişten gelen bilgiyi kullanmaz, gelecekte gelen bilgileri koruyup geçmişteki ile birleştirir ve herhangi bir zamanda bu bilgilere ulaşmayı mümkün kılar. Kurulan

²Sosyal haber ve paylaşım sitesi (www.reddit.com)

mimariye, temsil (embedding) katmanı eklenerek metinlerin vektör temsilleri elde edilebilir ve model eğitme sürecinde kullanılabilir.

B. BERT

Dil modelleri kelime dizilerinde geçen kelimelerin olasılıksal dağılımını hesaplar. Geleneksel dil modellerinde, bir kelime dizisinde önceki kelimeleri kullanarak takip eden kelime olasılıksal olarak hesaplanmaktadır. Bir dil modeli olan BERT [5] (Bidirectional Encoder Representations from Transformers (Dönüştürücülerden Çift Yönlü Kodlayıcı Temsilleri) ise çıkarımlarında hem önceki hem de ileride karşılaşacağı parçaları (token) kullanmaktadır. Bu işlem BERT’te maskeleye (masked language model) ile yapılmaktadır. Her kelime dizisindeki elemanların % 15’i maskelenir ve maskelenen kelime bağlam kullanılarak tahmin edilmeye çalışılır. Bunun yanı sıra, iki cümle arasında ardışıklık ilişkisi olup olmadığına göre de eğitilmiştir. Bunun nedeni, soru cevaplama gibi bazı doğal dil işleme görevlerinde önceki cümlelerden de çıkarımlar yapmanın gerekli olmasıdır. Dönüştürücü (Transformer) mimarisine sahip olan BERT farklı dillerde büyük metin derlemeleri (text corpus) ile önceden eğitilmiştir (pre-trained). Önceden eğitilmiş BERT modeli ince ayar (fine tuning) işlemiyle farklı doğal dil işleme görevlerinde kullanılabilir. Çalışmamızda PyTorch kütüphanesinde çok dilli (multilingual-based) BERT modelini ironi tespiti için kullandık. Bunun için 11 katmanlı hazır BERT modeli üzerine sınıflandırma katmanı ekleyerek ³ oluşturulan 12 katmanlı modelle ikili sınıflandırma (binary classification) uyguladık.

C. Geleneksel Gözetimli Öğrenme Yöntemleri

Geleneksel sınıflandırma yöntemlerinden olan Naif Bayes (NB) yöntemleri, gözetimli öğrenme (supervised learning) teknikleri altında yer almaktadır. Çalışma prensibi, bütün özelliklerin birbirinden koşullu bağımsızlık taşıdığını (conditional independence) varsayan *saf*, *naif* (*naive*) bir varsayıma dayanmaktadır. Bayes teoremi, yöntemin temelini oluşturur. Özelliklerin birbirine benzerliğinin (likelihood) hesaplanmasında kullanılan formüle göre farklı türlere ayrılır: Gauss NB, Multinomial NB, Bernoulli NB gibi.

Öte yandan Destek Vektör Makinesi yöntemi (Support Vector Machine (SVM)) ise yine başka bir gözetimli öğrenme metodudur. Bu yöntemin çalışma prensibini, temel olarak, bir düzlem üzerinde konumlanmış noktaları etiketlerine göre olabildiğince eşit bir şekilde ikiye ayıran ve iki tarafa da uzaklığını maksimize eden ideal hiperdüzlemi bulmaya çalışmak olarak tanımlayabiliriz.

IV. DENEYLER

A. Veri Kümesi

Veri kümesi ağırlıklı olarak sosyal medya platformu Twitter’den ve başka sosyal ağlarından derlenerek hazırlandı. Oluşturulan veriler araştırmacıların erişimine açıktır ⁴. Veri kümesinde, 110’u ironik, kalan 110’u ise ironik olmayan cümleler

olarak, toplam 220 cümle bulunmaktadır. Veriler 3 kullanıcının katıldığı bir kullanıcı anketinde oy çoğunluğu ile etiketlenmiştir. Veri, geleneksel yöntemlere beslenmeden önce NLTK (Natural Language Toolkit) ⁵ kullanılarak önce cümlelere, daha sonra alt parçalarına (token) ayrılmış olup, devamında her cümleye ait olan parça listesinden öznitelik çıkarımı yapılmıştır. Bütün veride, parçalara ayırma işleminden sonra toplam 2421 parça oluşmuştur. Bu da, 220 cümlede, cümle başına 11.005 parça düştüğünü göstermektedir. Ancak bu bütün parçaların kelime olmadığını, bazı parçaların *emoji/emoticonlar* ve noktalama işaretleri olduklarını not düşmeliyiz.

B. LSTM ve Bi-LSTM Mimarileri ile Analiz

LSTM kullanarak gerçekleştirdiğimiz deneyler 5 dönem (epoch) boyunca 5 kere çalıştırılarak elde edilmiştir. Kullandığımız ağ mimarisi bir temsil (embedding) katmanı, 100 birim içeren Uzun-Kısa Süreli Bellek katmanı ve bir yoğun (Dense) katmandan oluşmaktadır. Yitim (Loss) fonksiyonu olarak İkili Çapraz Entropi (Binary Cross Entropy) ve iyileştirici (Optimizasyon) olarak Adam İyileştirici kullanılmıştır. Modelimiz her bir cümle için kelime temsili (word embedding) değerleri kullanılarak eğitilmiştir. LSTM ve Bi-LSTM Ağları kullanarak tekrarladığımız deneylerin sonuçları Tablo I’de görülebilir. Deneylerde kullandığımız mimaride ağırlıkların her deneyde rastgele başlatılmasından dolayı farklı sonuçlar elde ettik. Bu nedenle 5 çalıştırma için elde ettiğimiz sonuçların ortalama ve standart sapma değerlerini sunmaktayız. Sonuçlardan da görüleceği üzere Bi-LSTM ve standart LSTM ortalama olarak aynı doğruluk değerine ulaşmışlardır, ancak kullanılan diğer yöntemlerin aşağısında kalmıştır.

C. BERT Mimarisi Analizi

Çalışmamızda sosyal medyadan alınıp etiketlenen veriler kullanılmıştır. Yapılan deneylerde üst değişkenler (hyperparameter) sabittir. Bu değerler [5] çalışmasına göre ve yaptığımız ön değerlendirmeler sonucunda belirlenmiştir. Dönem (epoch) sayısı 5, öğrenme hızı (learning rate) 0.00004, toptan boyutu (batch size) 16 olarak kullanılmıştır. Bunun yanı sıra ilk katmanlardaki parametre değerleri eğitilmeyip sadece son katmanlardaki parametrelerin güncellenmesi (weight freeze) yöntemiyle de sistemin başarısı ölçülmüştür. Yapılan bütün deneylerde 10 katlı çapraz doğrulama yöntemi kullanılmıştır. Tablo II’deki sonuçlarda, temel model bütün katmanlardaki parametrelerin eğitim sırasında güncellendiği modeli ifade etmektedir. Bu modelde 10 katlı çapraz doğrulama sonucunda %83.5 başarı elde edilmiştir. Üç katman modelinde ise BERT’in sadece son üç katmanındaki parametreler eğitilmiş olup önceki katmanlardaki değerler Bert’in ön eğitilmiş (pre-trained) parametre değerleridir. Aynı şekilde dört katman modelinde sadece son dört katman eğitilmiştir. Yapılan deneylere göre Bert’in en başarılı olduğu sistem sadece son altı katmanın eğitim için kullanıldığı modeldir. Bu sistemde modelin genel başarı oranı %87.50 olarak ölçülmüştür. Çok katmanlı sinirsel modellerde ilk katmanlarda çoğunlukla daha genel özellikler çıkarılmaktadır. Çalışmamızda kullandığımız verinin nitelik açısından kısıtlı olması göz önünde bulundurularak ilk katmanlarda BERT’in kendi değerleri kullanılmıştır. Bu sayede bütün katmanların yeniden eğitildiği sisteme göre

³<https://github.com/ThilinaRajapakse/pytorch-transformers-classification> adresindeki BERT modeli kullanılmıştır. 10 katlı çapraz doğrulama ve ağırlık dondurma (weight freeze) için kodlarda gerekli değişiklikler yapılmıştır.

⁴<https://github.com/teghub/Turkish-Irony-Dataset>

⁵<https://www.nltk.org/> NLTK içinde 50’den fazla ulusal derlem bulunduran, kolay bir Python arayüzüne sahip bir doğal dil işleme kütüphanesidir.

Yöntem	Doğruluk (Accuracy)	Hassasiyet (Precision)	Duyarlılık (Recall)	F1 Skoru
LSTM	50.00 $\pm$ 0.0	50.00 $\pm$ 0.0	100.00 $\pm$ 0.0	66.67 $\pm$ 0.0
Bi-LSTM	50.00 $\pm$ 8.2	48.11 $\pm$ 14.4	38.00 $\pm$ 19.2	41.80 $\pm$ 17.8

TABLO I: LSTM ve Bi-LSTM modellerinin Türkçe sosyal medya verilerinde 10 katlı çapraz doğrulama başarısı (% olarak)

Yöntem	Doğruluk	Hassasiyet	Duyarlılık	MCC [10]
Temel	83.50	86.09	80.00	0.681
Üç katman	83.00	85.86	80.00	0.675
Dört katman	84.00	87.45	80.00	0.691
Beş katman	86.00	87.85	84.00	0.729
Altı katman	87.50	90.71	84.00	0.760
Yedi katman	85.00	86.46	83.00	0.701

TABLO II: BERT modelinin Türkçe sosyal medya verilerinde 10 katlı çapraz doğrulama başarısı

Yöntem	Doğruluk	Hassasiyet	Duyarlılık	F1 Skoru
Gauss NB	%57.50	40.00	20.00	26.67
Multinom. NB	55.00	50.00	27.50	35.5
Destek Vektör Mak. (SVM)	60.50	66.00	30.00	41.25
BERT	87.50	90.71	84.00	87.23
LSTM	50.00	50.00	100.00	66.67

TABLO III: Geleneksel ve sinirsel modellerin Türkçe sosyal medya verilerinde 10 katlı çapraz doğrulama başarısı (% olarak)

daha iyi bir performans elde edilmiştir. Tablo II’de, doğruluk, hassasiyet, duyarlılık metriklerinin yanı sıra MCC metriği ile de sonuç sunmaktayız. Matthews korelasyon katsayısı (MCC) [10] ikili sınıflandırma işlemlerinde modelin başarısını ölçmek için kullanılan bir başka indikatördür. 1 değeri optimum sonucu göstermekteyken değer düştükçe modelin sınıflandırma başarısı azalmaktadır. MCC’nin 0 olduğu durum rastgele sınıflandırmaya eşdeğerdir. BERT modelinde son altı katmanın yeniden eğitildiği sistemde MCC değeri 0.760 olarak ölçülmüştür. Bu sonuç sistemimizin genel başarısı olan %87.50 ile uyumludur.

D. Geleneksel Yöntemlerle Analiz

Deneylerimize referans noktası oluşturması açısından, geleneksel sınıflandırma yöntemlerinin ironi tespit performansını da inceledik. Bunlardan en yaygın kullanılan yöntemlerden ikisi olan Naif Bayes sınıflandırıcı varyasyonları ve SVM ile ironi tespiti sonuçları, sinirsel yöntemlerle karşılaştırmalı olarak Tablo III’te sunulmuştur. Bu çalışmalarımızda sınıflandırmaya yardımcı olması açısından çeşitli öznitelikler kullandık. Bunlardan bazılarını, emojiler, emoticonlar (:) veya :(gibi sosyal medya ve mesajlaşma jargonunda bulunan duygu belirten semboller), ünlem sözcükleri (eyvah, aman, aferin), kelime vektörü bazı özellikler olarak sıralayabiliriz. Bu listedeki özniteliklerden yararlanılarak toplamda 8 farklı özellik kullanılmıştır. Tablo III’te görüleceği üzere, kullandığımız özniteliklere daha hassasiyet gösterebilen yöntemler kullanıldıkça, doğruluk, hassasiyet ve duyarlılık değerleri yükselmiş, ancak yine de BERT mimarisi kadar yüksek değerler elde edilememiştir.

V. SONUÇ

Çalışmamızda Türkçe sosyal ağ metinlerinde ironi tespiti problemi için sinirsel yöntemlerin performansını inceledik. Sinirsel model olarak LSTM ve Bi-LSTM mimarileri ve BERT mimarisi kullanarak model oluşturduk. Karşılaştırma amaçlı

olarak, ironi tespiti için literatürde sıkça kullanılmış olan öznitelikler ile, SVM ve Naif Bayes algoritmalarını kullanarak model geliştirdik. Deney sonuçları, ironi tespiti için BERT mimarisinin LSTM, Bi-LSTM ve geleneksel tekniklerden daha yüksek performans verdiğini göstermektedir. Veri miktarının sınırlı olması nedeniyle LSTM ve Bi-LSTM mimarilerinin performansının da sınırlı kaldığı değerlendirilmektedir. Gelecek çalışmalarda veri miktarının artırılması ve sinir ağı mimarisinin farklılaştırılması ile performansın iyileşme potansiyeli bulunmaktadır.

KAYNAKLAR

- [1] U. Ahmed, L. Zafar, F. Qayyum, and M. A. Islam, Irony Detector at SemEval-2018 Task 3: Irony Detection in English Tweets using Word Graph, *Proceedings of The 12th International Workshop on Semantic Evaluation*, pp. 581–586, Jun. 2018.
- [2] U. B. Baloglu, B. Alatas and H. Bingol, Assessment of Supervised Learning Algorithms for Irony Detection in Online Social Media, 1st International Informatics and Software Engineering Conference (UBMYK), Ankara, Turkey, pp. 1-5, 2019.
- [3] D. Bamman, N. A. Smith, Contextualized Sarcasm Detection on Twitter, *ICWSM*, 2015.
- [4] F. Barbieri, H. Saggion, and F. Ronzano, Modelling Sarcasm in Twitter, a Novel Approach, *Proceedings of the 5th Workshop on Computational Approaches to Subjectivity, Sentiment and Social Media Analysis*, pp. 50–58, June 2014.
- [5] J. Devlin, M. Chang, K. Lee, and K. Toutanova, BERT: Pre-training of deep bidirectional transformers for language understanding, *NAACL*, pp. 4171–4186, 2019.
- [6] O. Dulger, Turkce metinlerde ironi tespiti (Irony Classification in Turkish Text), *Ulusal Yazılım Mühendisliği Sempozyumu (UYMS)*, 2018.
- [7] B. Ghanem, J. Karoui, F. Benamara, P. Rosso, V. Moriceau, Irony Detection in a Multilingual Context, arXiv:2002.02427, 2020.
- [8] C. V. Hee, E. Lefever and V. Hoste, Exploring the Automatic Recognition of Irony in English Tweets, *Proceedings of COLING 2016, the 26th International Conference on Computational Linguistics*, 2016.
- [9] A. Joshi, V. Tripathi, K. Patel, P. Bhattacharyya, and Mark Carman, Are word embedding-based features useful for sarcasm detection?, *In Proc. of Conference on Empirical Methods in Natural Language Processing (EMNLP)*, 2016
- [10] B. Matthews, Comparison of the predicted and observed secondary structure of T4 phage lysozyme, *Biochimica et Biophysica Acta (BBA) - Protein Structure*, vol. 405, no. 2, pp. 442–451, 1975.
- [11] E. W. Pamungkas and V. Patti, "#NonDicevoSulSerio at SemEval-2018 Task 3: Exploiting Emojis and Affective Content for Irony Detection in English Tweets", *Proceedings of The 12th International Workshop on Semantic Evaluation*, 2018.
- [12] H. Taslioglu and P. Karagoz, Irony detection on microposts with limited set of features, *in Proceedings of the Symposium on Applied Computing (SAC '17)*, Association for Computing Machinery, New York, NY, USA, 1076–1081, 2017.
- [13] B. C. Wallace, D. K. Choe, and E. Charniak, Sparse, Contextually Informed Models for Irony Detection: Exploiting User Communities, Entities and Sentiment, *Proceedings of the 53rd Annual Meeting of the Association for Computational Linguistics and the 7th International Joint Conference on Natural Language Processing (Volume 1: Long Papers)*, 2015.
- [14] S. Zhang, X. Zhang, J. Chan, and P. Rosso, Irony detection via sentiment-based transfer learning, *Information Processing & Management*, vol. 56, no. 5, pp. 1633–1644, Sep. 2019.