

Heyelan Duyarlılık Analizinde Eğitim Seti Boyutunun Harita Doğruluğuna Etkilerinin Araştırılması

(Investigation of The Effects of Training Set Size on Map Accuracy in Landslide Susceptibility Analysis)

İsmail ÇÖLKESEN¹ , Emrehan Kutluğ ŞAHİN² 

¹Gebze Teknik Üniversitesi, Harita Mühendisliği Bölümü, Gebze, Kocaeli

²Bolu Abant İzzet Baysal Üniversitesi, İnşaat Mühendisliği Bölümü, Bolu

icolkesen@gtu.edu.tr

Geliş Tarihi (Received): 17.12.2019

Kabul Tarihi (Accepted): 29.05.2020

ÖZ

İnsanlar tarafından yapılan altyapılara zarar veren ve birçok yönden insan hayatını olumsuz bir biçimde etkileyen heyelanlar tarihte en yıkıcı doğal afetlerden birisi olarak nitelendirilmektedir. Afet yönetimindeki büyük önemi nedeniyle doğruluğu yüksek ve güncel heyelan duyarlılık haritalarının üretilmesi küresel ve yerel ölçekli çalışmalar için esas teşkil etmektedir. Literatürde, heyelan duyarlılık haritalarının üretiminde çeşitli stratejiler, algoritmalar, yöntemler ve bunların farklı kombinasyonlarının kullanımı önerilmektedir. Buna karşın heyelan duyarlılığının değerlendirilmesinde birtakım sınırlamalar ve belirsizlikler mevcuttur. Duyarlılık metodlarının tahmin edebilme kabiliyeti genellikle geçmiş heyelan olaylarına ilişkin envanter verisine dayalıdır. Envanter verisini temsil eden örnekleme boyutunun (eğitim veri seti boyutu) duyarlılık haritası üretiminde ve tahmin doğruluğu üzerine etkileri heyelan duyarlılık haritası üretiminde önemli bir araştırma konusudur ve ihtiyaç duyulan örnek boyutunun yüzdesi veya oranı henüz ortaya koyulmamıştır. Bu çalışmanın temel amacı heyelan duyarlılık analizinde eğitim veri seti boyutunun tematik harita doğruluğuna olan etkilerinin araştırılmasıdır. Bu amaca uygun olarak, literatüre uygun olarak 12 farklı boyutta eğitim verisi oluşturulmuş ve söz konusu veri setleri ile duyarlılık haritaları üretilmiştir. Oluşturulan farklı boyuttaki eğitim veri setleri kullanarak heyelan duyarlılık haritaları üretilmesinde lojistik regresyon (LR) ve rastgele orman (RO) algoritmaları kullanılmıştır. Performans analizinde genel doğruluklar, AUC (Area under the ROC Curve) ve ROC eğrileri (Receiver Operating Characteristic Curve) kullanılmış ve tahmin doğrulukları arasındaki farklılıklar Wilcoxon's işaretli sıralar testi yardımıyla değerlendirilmiştir. Çalışma sonuçları, LR ve RO algoritmalarının tahmin performanslarının eğitim veri seti boyutuna duyarlı olduğunu, RO ve LR için istatistiksel olarak en uygun örnekleme oranının 70:30 ve 30:70 olduğu belirlenmiştir. Bununla birlikte, ileri tahmin algoritması olan RO algoritmasının performansı klasik yöntem olan LR algoritmasına göre tüm durumlarda istatistiksel olarak daha üstün performans sergilediği görülmüştür.

Anahtar Kelimeler: Heyelan Duyarlılığı Haritalama, Eğitim Veri Seti Boyutu, Rastgele Orman, Lojistik Regresyon.

ABSTRACT

Landslides have been qualified as a one of the most destructive natural disasters in history that cause damage to man-made structures and negatively affect human life in many ways. Due to its great importance for hazard management, producing accurate and up-to-date landslide susceptibility maps is essential for many local- to global-scale studies. For this purpose, various strategies, algorithms, methods and their different combinations have been suggested in the literature. However, there are several limitations and uncertainties in assessing landslide susceptibility. The predictive capability of susceptibility methods mainly relies on the inventory of past landslide events. The effects of sampling size (training data size) on estimation accuracy is one of the important research topics in landslide susceptibility mapping. However, the percentage or proportion of sample size required for construction of susceptibility model has not been clearly defined and investigated by the researchers in the literature. The main purpose of this study is to investigate the effects of the training data set size on the thematic map accuracy in landslide susceptibility analysis. For this purpose, 12 different size training data sets were determined considering the sampling ratio used in landslide-related studies in the literature and susceptibility maps were produced with these sampling ratios. Logistic regression (LR) and random forest (RF) algorithms were applied to produce landslide susceptibility maps using different size training data sets. In the performance evaluations, overall accuracy, AUC (Area under the ROC Curve) and ROC curves (receiver operating characteristic curve) were used and the differences between estimation accuracy were evaluated with Wilcoxon's signed rank test. Results showed that prediction performances of LR and RO algorithms were sensitive to training dataset size and the most suitable sampling ratios were found to be 70:30 and 30:70 for RF and LR, respectively. Furthermore, the performance of the RF algorithm, known as advanced prediction algorithm, were found to be superior compared to those of traditional LR algorithm in all considered sampling ratios.

Keywords: Landslide Susceptibility Mapping, Training Data Set Size, Random Forest, Logistic Regression.

1. GİRİŞ

Heyelan duyarlılık analizinin esası bir bölgede gerçekleşen heyelan olayları ile heyelana etki edebilecek faktörlerin ilişkilendirilerek bölgenin tümü için heyelan duyarlı alanların tahmin edilmesi olarak ifade edilebilir. Bu açıdan ele alındığında heyelan duyarlılık haritalarının hazırlanması genel olarak heyelana etki eden faktörlerin belirlenmesi, yeterli sayıda örnekleme alanının tespiti, uygun bir tahmin algoritmasının seçimi ile modelleme ve doğruluk analizlerini içeren karmaşık işlem adımlarından oluşmaktadır.

Literatürde, heyelan duyarlılık haritalarının üretiminde farklı analiz tekniklerinin kullanımı, bu tekniklerin karşılaştırılması ve harita doğruluğuna etkilerinin incelenmesi önemli bir araştırma konusu olarak ön plana çıkmaktadır. Bu kapsamda ele alınan heyelan duyarlılık haritalarının doğruluğunu doğrudan veya dolaylı yoldan etkileyebilecek etmenlerin başında duyarlılık analizlerinin temel girdisi olan veri seti ve duyarlılık tahmininde kullanılacak algoritma gelmektedir. Literatürde yapılan çalışmalar incelendiğinde, heyelan duyarlılık analizlerinin gerçekleştirilmesinde temel olarak üç farklı yöntemin kullanıldığı görülmektedir. Bunlar istatistiksel, nicel ve sezgisel yöntemlerdir (Chung ve diğerleri, 1995; Suzen ve Doyuran, 2004; Kavzoglu ve diğerleri, 2015a). Literatürde özellikle bulut teknolojileri ve veri tabanı tasarımındaki gelişmelere paralel olarak özellikle coğrafi bilgi sistemleri (CBS) destekli çeşitli veri işleme teknikleri yardımıyla heyelan duyarlılık analizlerinin yapılmasına yönelik çalışmalar mevcuttur (Yalcin ve diğerleri, 2011, Kavzoglu ve diğerleri, 2014). Kullanılan tahmin algoritmaları arasında çok değişkenli istatistik tabanlı bir algoritma olan lojistik regresyon algoritması heyelan duyarlılık haritalarının oluşturulmasında yaygın olarak kullanılmaktadır (Dai ve Lee, 2002; Bai ve diğerleri, 2011, Kavzoglu ve diğerleri, 2015b). Kullanılan temel tahmin algoritmalarının yanında yapay sinir ağları, destek vektör makineleri, karar ağaçları ve toplu öğrenme algoritmaları gibi sezgisel algoritmalar ile duyarlılık haritalarının üretilmesi önemli bir araştırma konusu olmuştur (Gómez ve Kavzoglu, 2005; Yao ve diğerleri, 2008; Yilmaz, 2010; Colkesen ve diğerleri, 2016, Chen ve diğerleri 2017). Özellikle son yıllarda heyelan duyarlılık haritası üretimi işleminin bir regresyon işleminden ziyade heyelan olan ve olmayan alanların tespit edilmesi şeklinde ikili bir sınıflandırma problemi olarak ele alındığı görülmektedir. Bu kapsamda literatürde farklı sınıflandırma algoritmalarının

heyelan duyarlılığı tespitinde kullanımına yönelik araştırmaların sayısında önemli bir artış görülmektedir. Örneğin, Hong ve diğerleri (2015) Çin'in Yihuang bölgesine ait heyelan duyarlılık haritalarının üretilmesinde kernel lojistik regresyon, karar ağaçları ve destek vektör makineleri sınıflandırıcılarını kullanmış ve karar ağacı sınıflandırıcısının diğer algoritmalarından daha iyi sonuç verdiğini ifade etmiştir. Tien Bui ve diğerleri (2016) sınıflandırma yaklaşımıyla heyelan duyarlılık haritası üretiminde destek vektör makineleri, yapay sinir ağları, kernel tabanlı lojistik regresyon ve lojistik model ağaç algoritmalarının kullanımını incelemiştir. Pham ve diğerleri (2016) Viet Nam'ın Luc Yen bölgesine ilişkin heyelan duyarlılık analizinde Naive Bayes ve rotasyon orman algoritmalarından oluşan hibrit bir sınıflandırma modeli kullanımını incelemiştir. Bui ve diğerleri (2016) yağışa bağlı heyelan alanlarının modellenmesinde toplu öğrenme algoritmalarından fonksiyonel ağaçlar, AdaBoost, Bagging ve MultiBoost sınıflandırma algoritmalarını kullanmıştır.

Kontrollü sınıflandırma, heyelan duyarlılık tahminlerini de içerisine alan birçok çalışmada en sık kullanılan veri analiz tekniğidir. Literatürde sınıflandırma veya tahmin doğruluğunun kullanılan eğitim veri setinin bir fonksiyonu olduğu, dolayısıyla eğitim veri setinin boyutunun ve kalitesinin model doğruluğu ile doğrudan ilişkili olduğu ifade edilmektedir (Foody, 1999; Tsai ve Philpot, 2002; Foody ve Mathur, 2004). Bu nedenle kullanılacak algoritma farklı biçimlerde eğitim veri setine ihtiyaç duyabilmektedir (Foody, 1999). Diğer bir ifadeyle, oluşturulan eğitim veri seti bir algoritma için yüksek doğrulukta sonuçlar ortaya çıkarabilirken, diğer bir algoritmanın doğruluğunda azalmaya neden olabilmektedir. Literatürde eğitim veri setinin tespitine yönelik çalışmalar mevcuttur. Örneğin, Mather (2004), eğitim veri setinin, her bir sınıf için (örneğin heyelan olan ve olmayan alanlar) en az kullanılan faktör sayısının 10 veya 30 katı (10p ve 30p kuralı) kadar örnek nokta veya piksel içermesi gerektiğini ifade etmiştir. Congalton ve Green (2008), eğitim veri seti boyutunun belirlenmesinde çoklu terimli olasılık (multinomial distribution) teorisine dayanan istatistiksel bir yaklaşım önermişlerdir. Duyarlılık analizinde kullanılacak algoritmaların tahmin kabiliyeti genellikle geçmiş heyelan olaylarına ilişkin envanter verisine dayalı olduğundan, heyelanla ilişkili çalışmalarda en fazla tartışılan konulardan birisi de eğitim ve doğrulama veri setinin kalitesidir. Bu noktada envanter verisini temsil eden örnekleme boyutunun (eğitim veri seti boyutu) duyarlılık haritası üretimine ve tahmin

doğruluğu üzerine etkileri heyelan duyarlılık haritası üretiminde göz önüne alınması gereken önemli bir husustur. Duyarlılık analizinde kullanılan algoritmalar, veri seti yardımıyla tahmin modeli oluşturulmasında farklı yaklaşımlar veya kabuller sonrasında seçilecek eğitim – test veri seti boyutu model doğruluğu üzerinde etkili olacaktır. Örneğin lojistik regresyon ve Naive Bayes algoritmaları gibi istatistiksel temellere dayanan parametrik bir algoritma kullanımında eğitim verisinin oluşturulmasında olasılık fonksiyonları kullanılmakta ve bir noktadaki duyarlılık seviyesi tahmin edilebilmektedir. Bu nedenle istatistiksel hesaplamalar yardımıyla tahminler gerçekleştirebilen parametrik algoritmalar için fazla sayıda örneğe ihtiyaç duyulmaktadır. Söz konusu istatistiki esaslara dayalı tahmin algoritmalarının aksine son yıllarda daha az örnek kullanımıyla daha doğru tahminler yapabilen yapay sinir ağları, karar ağaçları ve destek vektör makineleri gibi makine öğrenme algoritmaları birçok çalışmada başarıyla kullanılmaktadır (Kavzoglu ve Colkesen, 2012; Şahin, 2017)).

Heyelan duyarlılık haritası üretiminde en uygun parametre setinin belirlenmesi (özellik seçimi), duyarlılık haritası üretimi ve sonuç haritalarının doğruluklarının değerlendirilmesi gibi birçok işlem adımında eğitim ve test verisine ihtiyaç duyulmaktadır. Son yıllardaki literatür incelendiğinde eğitim ve test verisinin belirlenmesinde farklı yöntemlerin kullanıldığı görülmüştür. Örneğin, Chen ve diğerleri (2014) coğrafi bilgi sistemlerine dayanan istatistik bilgi değeri metodu ile Çin'in Chencang bölgesi için heyelan duyarlılık haritası üretmişlerdir. Çalışmada kullanılan eğitim ve test verisinin tespitinde 120 heyelan olan alanı içeren envanter haritasının %70'i (84 olan heyelan) eğitim verisi ve %30'luk kısmı (36 heyelan olan alan) test verisi olacak şekilde değerlendirilmiştir. Yang ve diğerleri (2015) çalışmasında heyelan duyarlılık haritası üretimi ve doğrulama işlemleri için heyelan envanter verisinin %80'lik kısmını (4.550 heyelan olan alan) eğitim verisi olarak kullanırken %20'lik kısmını (1.138 heyelan alanı) test verisi olarak kullanılmıştır. Bir diğer çalışmada, Pradhan ve Kim (2014) ise heyelan duyarlılık haritasının model üretim ve doğruluk safhalarında kullanılmak üzere 748 adet heyelan alanının %90'ını (673 heyelan alanı) eğitim verisi ve %10'unu (75 heyelan alanı) ise test verisi olarak değerlendirmeye almıştır. Youssef ve diğerleri (2015) çalışmalarında heyelan duyarlılık haritaları üretiminde üç farklı metot (iki değişkenli istatistiksel yaklaşım, frekans oranı ve ağırlıklı kanıt) kullanmışlardır. Söz konusu çalışmadaki

bu üç metotla heyelan duyarlılık haritası üretimi ve her bir metodun performans karşılaştırmalarında kullanmak üzere heyelan envanter verisinin (106 heyelan alanı) %75'ini eğitim ve %25'ini ise test verisi olacak şekilde değerlendirmeye alınmıştır. Brenning (2005) heyelan duyarlılık haritası üretimi için lojistik regresyon, destek vektör makineleri ve önyüklemeli sınıflandırma ağaçları kullanmıştır. Çalışmada gerçekleştirilen tüm analizlerde (harita üretimi, performans karşılaştırma ve doğrulama vb.) heyelan envanter verisi kapsamındaki olan ve olmayan heyelan alanlarının %50'si eğitim ve %50'si test verisi olacak şekilde kullanılmıştır. Duyarlılık haritası üretimindeki işlem adımları içerisinde kilit öneme sahip olan eğitim ve test veri seti tespiti ve farklı veri setinin duyarlılık haritası üzerindeki etkileri son yıllarda birçok çalışma için temel araştırma konusu olmuştur. Heckmann ve diğerleri (2014) yoğun yağış sonrası oluşan çamur akıntısı sonrası oluşabilecek duyarlılık haritası üretimi için model üretiminde lojistik regresyon metodu kullanılmıştır. Duyarlılık haritası üretiminde $n=50$ 'den başlayarak $n=5000$ 'e kadar 1000 adet örneklemden oluşan model üretilmiştir. Çalışma sonrasında denenen örnekler içerisinde $n=300-350$ arası değerlerde en iyi model üretimi tespit edilirken bu sayının artmasının doğruluk üzerinde etkisinin olmadığı görülmüştür. Petschko ve diğerleri (2014) çalışmasında Güney Avusturya'nın heyelan duyarlılık haritası modelleri üretiminde esnek bir istatistiksel metot olan genelleştirilmiş katkı modeli kullanılmıştır. Çalışmada litoloji sınıflarına dayanarak üretilen 16 model bulunmaktadır. Bu modellerin performanslarının tespitinde mekânsal ve rastgele alt kümeleme ile tekrarlanan k -katlı çapraz doğrulama tekniği kullanılarak değerlendirilmeye alınmıştır. Heyelan olan ve olmayan alanların kapsadığı envanter verisi üzerinden çapraz doğrulama tekniği ile 9 kez tekrarlanarak (5-katlı ve 20 tekrar) 12.652 örneklemden 50 örnekleme kadar eğitim veri seti üretilmiştir. Doğrulama sonuçları heyelan olan ve olmayan her bir alandaki 200 örneklem verisi (toplamda 400) bu çalışma için tavsiye edilen veri seti büyüklüğü olmuştur. Literatürden de görüleceği üzere, heyelan duyarlılık haritasının üretilmesine ilişkin yapılan çalışmalarda eğitim veri seti boyutunun harita doğruluğuna olan etkilerinin önemli bir araştırma konusu olduğu ve çalışmalarda envanter verisinin tamamının veya büyük bir kısmının model oluşumunda kullanıldığı görülmektedir.

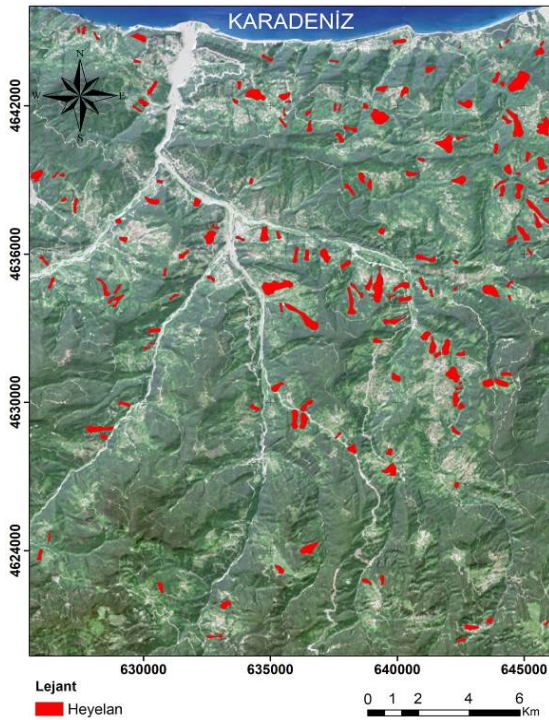
Bu çalışma kapsamında, heyelan duyarlılık haritası üretiminde farklı boyutlarda eğitim ve test

veri setinin kullanımının harita doğruluğuna etkileri araştırılarak en uygun veri seti boyutu incelenacaktır. Çalışmada, seçilecek bir pilot bölge için öncelikle heyelan envanter verisi ve heyelana etki eden faktörler belirlenerek modellemeye esas veri seti oluşturulacaktır.

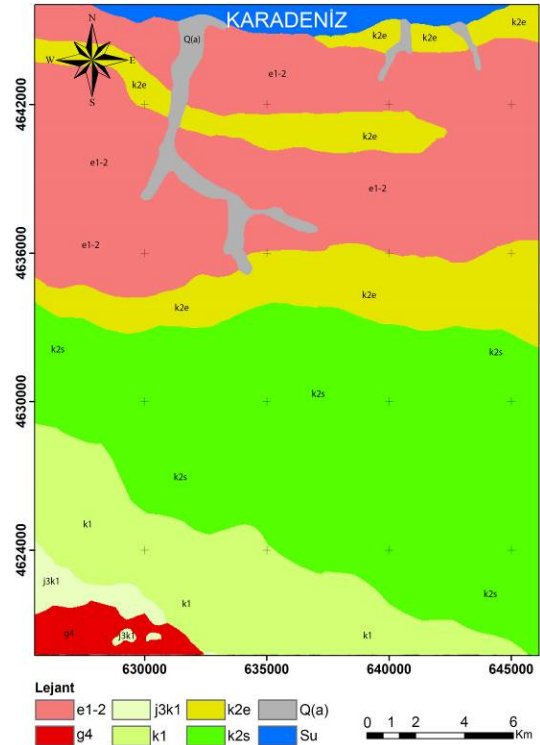
Çalışmanın amacına uygun olarak, envanter verisinin değişik yüzde ve oranlarda eğitim verisi olarak kullanılacağı farklı boyutlarda veri setleri oluşturulacaktır. Bunun yanında literatürde farklı modelleme problemlerinin çözümünde eğitim veri seti boyutunun belirlenmesi amacıyla yaygın olarak kullanılan diğer teknikler (örneğin 10p, 30p kuralı ve çok terimli olasılık teorisi) değerlendirmeye alınarak heyelan duyarlılık analizinde uygulanabilirliği incelenmektedir. Oluşturulan farklı boyutta veri setleri kullanılarak parametrik lojistik regresyon algoritması ve parametrik olmayan rastgele orman algoritması kullanılarak pilot bölgeye ilişkin ayrı ayrı heyelan duyarlılık haritaları üretilecektir. Doğruluk analizleri neticesinde, farklı boyutta eğitim veri seti kullanımının harita doğruluğu üzerindeki etkileri ortaya konulacaktır.

2. ÇALIŞMA ALANI VE VERİ SETİ

Uygulama kapsamında, Türkiye Pafta bölümlenmesinde 1/100.000'lik pafta indeksinde E33 paftasına denk gelen alan çalışma alanı olarak belirlenmiştir. Söz konusu çalışma alanı Türkiye'nin Orta Karadeniz bölgesinin Sinop ili sınırları içinde yer almaktadır (Şekil 1a). Çalışma alanının Kuzey-Batısında Türkeli ilçesi, Kuzey Doğusunda Sinop Merkezi ve Güneyinde ise Boyabat ilçesi bulunmaktadır. Yaklaşık 2.330 km²'lik bir alanı kapsayan çalışma bölgesi coğrafik olarak 34°00' ve 34°29' doğu meridyenleri ve 41°59' ve 41°29' kuzey paralelleri arasında kalmaktadır. Bölgenin coğrafi özellikleri, üzerinde bulunduğu jeolojik yapısı ve bölgenin mevcut meteorolojik koşulları göz önüne alındığında, çalışma alanının heyelan tehlikesine duyarlı bir bölge içerisinde yer aldığı ifade edilebilir. Özellikle 1998 yılında gerçekleşen ve Batı Karadeniz bölgesinde etkili olan aşırı yağışlar ve su baskınları, çalışma alanını da içerisine alan birçok ilde can ve mal kayıplarına neden olan heyelan vakalarının yaşanmasına sebep olmuştur. Yaşanan afetler binlerce konutu etkileyerek hafif ve ağır hasarlara sebep olurken birçok insanın da ölümüne neden olmuştur.



(a)



(b)

Şekil 1. (a) Çalışma alanı ve konumu; (b) Çalışma alanına ait litoloji haritası.

Çalışmada kullanılacak verilerin temininde istenilen doğruluk ve güvenilirlik kapsamında veriyi üreten kurum ve kuruluşlarla irtibata geçilmiştir. Yapılan çalışmalar sonucunda litoloji (Şekil 1b) ve heyelan envanter verisi (Şekil 1a) Maden Tetkik ve Arama Genel Müdürlüğü'nden, sayısal yükseklik verisi Harita Genel Müdürlüğü'nden (HGM), Orman ve Su İşleri Bakanlığı'ndan, akarsu ve yol verisi elde edilmiştir. Eş yükseklik verisi halindeki veriler 1/25.000 ölçeklidir. Vektörel ortamdaki eş yükseklik verileri CBS ortamında 30m piksel boyutunda raster sayısal yükseklik modeli verisinin üretiminde kullanılmıştır. Bu nedenle farklı piksel boyutlarına sahip diğer veri türlerinde (örn. litoloji ve arazi örtüsü/kullanımı vb.) 30m piksel boyutuna indirgenmiş ya da yükseltilmiştir. Bu işlem verilerin sahip olduğu veri çözünürlüğü dikkate alındığında herhangi bir bilgi kaybı yaşanmadığı ayrıca belirtilmelidir.

Çalışma alanındaki heyelana etki eden faktörlerin tespitinde son yıllarda yapılan çalışmalarda heyelan duyarlılık haritaları üretiminde sıklıkla kullanılan faktörler dikkate alınmıştır. Heyelan duyarlılık analizlerinde kullanılan en önemli veri kaynağı Sayısal Yükseklik Modeli (SYM) ve söz konusu model yardımıyla oluşturulan bakı, plan, profil ve eğrilik haritası, drenaj yoğunluğu, eğim, eğim uzunluğu, akarsu güç indeksi, sediment taşıma indeksi ve topoğrafik nemlilik indeksini de içerisine alan verilerdir. SYM yardımıyla söz konusu veri setlerinin oluşturulmasında çeşitli raster analizlerinden faydalanılmıştır. Tablo 1'de çalışma kapsamında değerlendirmeye alınan veri setleri ve SYM yardımıyla oluşturulan faktörlerin ismi, özelliği, elde edildiği yer ve veri yapısına ilişkin bilgiler verilmiştir. Kurumlardan elde edilen veya SYM'den üretilen tüm faktörler ortak bir konumsal referans sistemine taşınmıştır.

Heyelan envanter haritasının kapsadığı topoğrafik verinin heyelanın meydana gelmesinden sonra bölgede gelişen değişimleri de içereceği göz önünde bulundurulmalıdır. Bu durumda, envanter veri seti içerisinde heyelan hadisesinin gerçekleştiği alanlarla birlikte heyelan oluşumu üzerinde etkisi olmayan doğal koşullara veya topoğrafik özelliklere sahip olan alanların varlığı söz konusu olabilmektedir. Bu durumun nedenle, analiz işlemleri sırasında kararı etkileyebilecek ve bozucu etkilere sebep olabilecek alanlar (yollar, göl ve göletler, su kenarları, heyelan akması vb.) envanter verisinden çıkarılarak veri setinde güncelleştirme işlemi yapılmıştır. 1/25.000 ölçeğinde üretilmiş heyelan envanter haritasında çalışma bölgesine

ait 456 adet heyelan alanı bulunmaktadır. Vektör poligon türündeki heyelan verisi 30m çözünürlükte raster verisine dönüştürülmüştür.

3. MATERYAL VE YÖNTEM

Çalışma kapsamında heyelan duyarlılık haritalarının üretilmesinde takip edilen işlem adımları, genel olarak verilerin temini ve ön işlemler, çalışma alanlarına ilişkin heyelana etki eden faktörlerin belirlenmesi, envanter haritasının oluşturulması, heyelana etki eden faktörler içerisinden istatistiksel yaklaşımla optimum faktör sayısının tespiti, modellerin oluşumunda kullanılacak farklı boyutlarda eğitim ve doğrulama veri setlerinin oluşturulması, LR ve RO algoritmaları kullanılarak duyarlılık haritalarının üretilmesi ve doğruluk/duyarlılık analizlerini içerisine alan 6 temel işlem adımından oluşmaktadır. Söz konusu işlem adımları, gerçekleştirilen işlemler ve elde edilen sonuçlar bu bölümde ayrıntılı olarak ele alınmıştır.

a. Özellik/Faktör Seçimi

Heyelan duyarlılık analizlerini de içerisine alan birçok tahmin probleminde, temel veri setini oluşturan faktörler arasında korelasyonlu, gürültü içeren ve tahmin doğruluğunu olumsuz olarak etkileyebilecek özelliklerin belirlenmesi ve temel veri seti içerisinden çıkarılması söz konusudur (Webb ve Copsey, 2011). Sınıflandırma veya regresyon analizinde girdi veri setindeki özellik sayısının artmasının teorik olarak oluşturulacak tahmin modelinin doğruluğunu arttıracığı öngörülmektedir. Ancak, veri seti içindeki bazı özellikler tahmin doğruluğunu olumsuz şekilde etkileyebilecek düzeyde gürültüye sahiptir. Bu sebeple yüksek gürültü oranına sahip özelliklerin belirlenerek veri seti içerisinden çıkarılması tahmin doğruluğunun iyileştirilmesi noktasında etkili olabilmektedir. Diğer taraftan, veri seti boyutunun azaltılması modelleme sırasında ihtiyaç duyulan hesap uzayının boyutu gözetildiğinde donanımsal alt yapının (bellek, ağ bant genişliği ve depolama alanı) niceliği üzerinde tasarrufa neden olarak hız, zaman ve iş gücünde tasarrufa neden olacaktır. Bu çalışmada değerlendirmeye alınacak faktör veri setleri içerisinden en uygun faktör seti kümesinin belirlenmesinde Pearson's korelasyon analizi kullanılmıştır. Faktör veri setleri korelasyon analizi yardımıyla değerlendirilmiş ve faktörlere ait önem dereceleri belirlenerek sonuçlar analiz edilmiştir.

Tablo 1. Çalışmada kullanılan faktörlerin kaynak ve türleri.

Faktör Özelliği	Faktör Adı	Kaynağı	Türü
Yüzey jeolojisi	Heyelan envanter verisi	Literatür	Poligon
Yüzey Topoğrafya bilgisi	Litoloji		
	Yükseklik	Sayısallaştırılmış çizgi	
	Plan Eğriliği		
	Profil Eğriliği		
	Bakı	SYM	Raster
	Eğim		
Yüzey-su ilişkisi	Akarsudan olan Uzaklık		
	Topoğrafik Nemlilik İndeksi	SYM	Raster
	Akarsu Güç İndeksi		
	Sediment Taşıma İndeksi		
Yüzey Vegetasyon Bilgisi	NDVI		
Yüzey Kullanım Bilgisi	Arazi kullanımı/örtüsü	Uydu görüntüsü	Raster

İki değişken için Pearson's korelasyon katsayısı hesabında değişkenlerin kovaryanslarının standart sapmalarına oranları hesaplanmaktadır (Biesiada ve Duch, 2007). X değerlerine sahip x özelliği ve y değerlerine sahip Y sınıfları rastgele değişkenler olarak tanımlandığı düşünüldüğünde korelasyon katsayısı Eşitlik (1) yardımıyla hesaplanır.

$$\rho(X,Y) = \frac{E(XY) - E(X)E(Y)}{\sqrt{\sigma^2(X)\sigma^2(Y)}} = \frac{\sum_i (x_i - \bar{x}) \sum_i (y_i - \bar{y})}{\sqrt{\sum_i (x_i - \bar{x})^2 \sum_i (y_i - \bar{y})^2}} \quad (1)$$

Pearson's korelasyon katsayısı [-1,1] arasında değerler almaktadır. Örnek vermek gerekirse bulunan istatistik değer -1 veya 1'e eşit olması koşulunda iki değişken arasında mükemmel bir doğrusal ilişki olduğu, fakat İstatistik değer 0'a eşit ise de ikisi arasında ilişki olmadığını göstermektedir. Hesaplanan değer $0 < \rho(X,Y) < 1$ veya $0 > \rho(X,Y) > -1$ aralığında değerler aldığında, özellikler arasında pozitif veya negatif ilişki olduğu söylenebilir $\rho(X,Y)$ büyük değerler aldıkça ilişkinin gücü artmaktadır.

b. Eğitim ve Doğrulama Veri Seti Üretimi

Heyelan envanter haritalarında bölge içerisinde gerçekleşen heyelanlar konumları ve büyüklükleri esas alınarak genellikle vektör-tabanlı poligonlar şeklinde belirlenmektedir. Bununla birlikte heyelan duyarlılık haritası üretiminde esas veri seti olarak kullanılan heyelan duyarlılık faktörleri genellikle raster veri yapısı formatındadır. Diğer bir ifadeyle heyelan duyarlılık haritası üretiminde vektör tabanlı konumsal analizlerden ziyade raster tabanlı

(piksel tabanlı ya da grid tabanlı) analizler yardımıyla heyelan duyarlılıkları tahmin edilmektedir (Carrara, 1983; van Westen, 1993).

Çalışmada eğitim veri seti boyutundaki değişimlerin heyelan duyarlılık haritası doğruluğuna etkilerini incelemek amacıyla, envanter verisinin %10, %20, %30, %40, %50, %60, %70, %80 ve %90 oranlarında rastgele örneklenerek 9 farklı eğitim veri seti oluşturularak değerlendirilmeye alınmıştır. Test veya doğrulama veri seti olarak, envanter veri seti içerisinde eğitim veri seti için seçilen örnekler haricinde kalan pikseller kullanılmıştır. Bununla birlikte, eğitim seti büyüklüğü tespiti için kabul edilen çeşitli tekniklerden (ör. 10p, 30p kuralları ve çok terimli olasılık teorisi) yararlanılmıştır. Sonuç olarak 12 farklı boyutta eğitim ve doğrulama veri seti oluşturulmuştur.

Mather (2004), eğitim veri setinde, her bir sınıf için (örneğin heyelan olan ve olmayan alanlar) en az kullanılan faktör sayısının 10 veya 30 katı (10p ve 30p kuralı) kadar örnek nokta veya piksel içermesi gerektiğini ifade etmiştir. Örneğin, heyelana etki eden 10 faktör dikkate alınarak gerçekleştirilecek heyelan duyarlılık analizinde, tahmin modeli eğitiminde 10p kuralına göre heyelan olan alanları temsil eden minimum 100 örneğe ihtiyaç duyulurken, 30p kuralına göre minimum örnek sayısı 300'dür. Congalton ve Green (2008), eğitim veri seti boyutunun belirlenmesinde çoklu terimli olasılık teorisine dayanan istatistiksel bir yaklaşım önermişlerdir. Bu yaklaşıma göre her bir sınıf için gerekli olan eğitim örneği sayısı (n),

$$n = B \Pi_i (1 - \Pi_i) / b_i^2 \quad (2)$$

eşitliği yardımıyla hesaplanmaktadır. Bu eşitlikte istenen doğruluğu oransal olarak ifade eder (örneğin istenen güven aralığı %95, istenen hassasiyet $\alpha = 0,05$), Π_i her bir sınıfı temsil eden örneklerin tüm örnekler içindeki yüzdesini, B serbestlik derecesi 1 olan ki-kare olasılık tablosundan hesaplanan tablo değeri olup üst yüzdelik değeri şeklinde hesaplanır. Burada α güven aralığını ve k sınıf sayısını göstermektedir (Congalton ve Green, 2008). Örneğin heyelan olan ve heyelan olmayan alanların tespiti için gerçekleştirilecek iki sınıflı ($k=2$) bir duyarlılık analizinde, istenen güven aralığı %95, istenen hassasiyet %5 ve heyelan olan alanların çalışma alanının %30'unu kapsadığı ($\Pi_i = \%30$) göz önüne alınsın. B tablo değeri, serbestlik derecesi 1 olan ki-kare olasılık tablosundan ve üst yüzdelik değeri $1 - \alpha/k$ ($1 - 0,05/2 = 0,0975$) için hesaplanır. Bu durumda, B ki-kare tablo değeri $\chi^2_{(1,0,975)} = 5,024$ 'tür. Eşitlik (2) yardımıyla gerekli olan örnek sayısı $n = 5,024(0,30)(1 - 0,30) / (0,05)^2 \cong 420$ şeklinde elde edilir. Dolayısıyla göz önüne alınan örnek problem için eğitim veri setinde heyelan olan alanları temsil edecek minimum örnek sayısı 420 olmalıdır.

c. Rastgele Orman Algoritması

Rastgele orman (RO) algoritması belirli bir sınıflandırma ve regresyon problemi çözümünde birçok karar ağacı modelinin bir arada kullanılmasını esas alan toplu öğrenme algoritmalarından birisidir (Breiman, 2001). RO model oluşumunda birden çok karar ağacını kullanması nedeniyle literatürde karar ormanı olarak da bilinmektedir. Algoritma, sınıf etiketi belli olmayan bir örneğe ilişkin tahmin işleminin gerçekleştirilmesinde ormanı oluşturan karar ağaçlarının her birinin yaptığı tahminlerin birleştirilmesi ve ilgili örnek için nihai kararın verilmesi prensibini esas almaktadır. Bu amaç doğrultusunda öncelikli olarak kontrollü sınıflandırma işleminin gerçekleştirmek üzere hazırlanan orijinal eğitim veri seti içerisinde kullanıcı tarafından belirlenen karar ağacı sayısı kadar rastgele altkümeler oluşturulur. Bu işlemdeki temel düşünce ormandaki her bir karar ağacı modelinin farklı eğitim veri setleri ile oluşturulması ve ormandaki çeşitliliğin sağlanmasıdır (Kuncheva ve Whitaker, 2003). Rastgele oluşturulan altkümelerin her birinin yaklaşık olarak %67'si karar ağacı modeli oluşumu için kullanılırken geriye kalan %33'lük kısmı oluşturulan ağaç modelinin performans analizini gerçekleştirmek için kullanılmaktadır.

Sınıf etiketi belli olmayan bir örnek (piksel) için ormandaki her bir karar ağacı tarafından tahminler yapılır ve tahminlerin birleştirilmesi neticesinde en fazla oyu alan sınıf etiketi ilgili örneğe atanır (Pal, 2005; Kavzoglu ve Colkesen, 2013). RO tahmin modelinde, karar ağaçlarının oluşturulması ve ağaç dallanmasında kullanılacak öznelitliklerin tespitinde Gini indeksi kullanılır. RO algoritmasının uygulanması sürecinde kullanıcılar tarafından belirlenmesi gereken iki parametre mevcuttur. Bunlar, ormandaki ağaç sayısı (m) ve karar ağacı modellerinin oluşturulması için her bir düğüm noktasında değerlendirmeye alınacak örneklerin sayısı (n) olarak bilinmektedir.

ç. Lojistik Regresyon Algoritması

Lojistik regresyon (LR) algoritması matematiksel olarak kolay anlaşılabilir yapısı ve hızlı işlem süresi nedeniyle heyelan duyarlılık analizi araştırmaları gibi birçok regresyon probleminin çözümünde sıklıkla kullanılan istatistiksel tabanlı bir algoritmadır (Yılmaz, 2009). LR algoritmasının temelinde, bağımsız değişkenler (heyelan olayına etki eden faktörler) ile bağımlı değişken (heyelan olayının olması) arasında çok değişkenli bir regresyon ilişkisi kurmayı sağlamaktır (Lee, 2005). Çoklu doğrusal regresyon modeli aşağıdaki eşitlik ile gösterilebilir,

$$Y = b_0 + b_1x_1 + b_2x_2 + \dots + b_nx_n \quad (3)$$

Doğrusal regresyon modelindeki Y , 0 ve 1 aralığındaki yanıt değişkeni, b_0 bağımlı değişken sabit sayısı, $b_1, \dots, b_n$ regresyon katsayılarını ve $x_1, \dots, x_n$ ise bağımsız değişkenlerin 1'den n 'e kadarki düzey değerlerini göstermektedir.

4. UYGULAMA

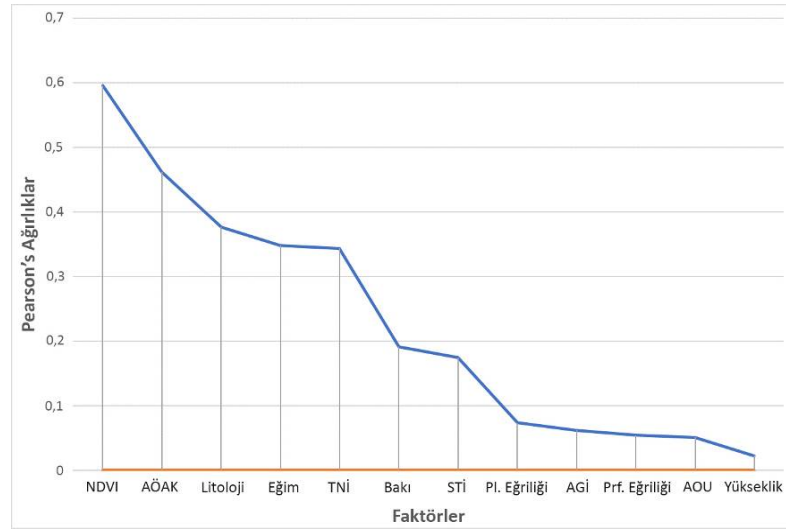
Çalışma kapsamında değerlendirmeye alınacak faktör veri setleri içerisinde en uygun faktör seti kümesinin belirlenmesinde Pearson's korelasyon analizi kullanılmıştır. Bu kapsamda her faktör için elde edilen ağırlık aşağıda Şekil 2'de verilmiştir. Şekilde verilen faktör ağırlıkları incelendiğinde en yüksek ağırlıklı faktörün NDVI verisi olduğu görülmektedir. Diğer etkin faktörler incelendiğinde sırasıyla Arazi Örtüsü ve Arazi Kullanımı (AÖAK), litoloji, eğim ve Topoğrafik Nemlilik İndeksi (TNI) faktörleridir. Diğer taraftan, plan eğriliği, Akarsu Güç İndeksi (AGİ), profil eğriliği, Akarsudan Olan Uzaklık (AOA) ve yükseklik faktörleri ise 0,1 ağırlık katsayısının altında kalmıştır ve Pearson's analizine göre en az etkin faktörler olarak tespit edilmiştir. Elde edilen sonuçlar analiz edildiğinde farklı sayıda

eğitim ve test verisinin üretimi ve etkinliğinin araştırılmasında tüm faktörlerin kullanılması uygun görülmüştür. Sonuç olarak faktör çıkarımı işlemi yapılmamış ve mevcut tüm faktörler kullanılarak duyarlılık analizi yapılacaktır.

a. Eğitim ve Doğrulama Verisi Üretimi

Heyelan duyarlılık haritalarının üretilmesinde eğitim ve test verisi için 454 heyelan olan ve olası 67 adet heyelan olmayan alanlar poligon türünde tespit edilmiştir. Heyelan envanter haritası bünyesindeki heyelan olan alan sayının çalışma alanı için yeterli sayıda olması nedeniyle veri seti eğitim/test ile doğrulama verisi farklı sayılarda bölünmüştür. Çalışma amaçlarına uygun olarak eğitim veri seti boyutundaki değişimlerin heyelan

duyarlılık haritası doğruluğuna etkilerini incelemek amacıyla, envanter verisinin %10, %20, %30, %40, %50, %60, %70, %80 ve %90 oranında eğitim verisi olarak kullanılacağı 9 farklı eğitim veri seti oluşturulmuştur (Tablo 2). Heyelan envanter veri setindeki alanların 30m raster veri setine döndürülmesi sonrasında 59.628 adet piksel elde edilirken olmayan alan analizlerinde ise 50.105 adet 30m piksel üretilmiştir. Bu örnekleme oranlarına ek olarak 10p (10*12 faktör= heyelan olan sınıf için minimum 120 piksel ve heyelan olmayan sınıf için minimum 120 piksel), 30p ve çoklu terimli olasılık teorisine göre hesaplanan ve Delta olarak adlandırılan veri setlerine ilişkin örnek sayıları da tabloda verilmiştir.



Şekil 2. Pearson's korelasyon analizi sonrası elde edilen faktör ağırlıkları

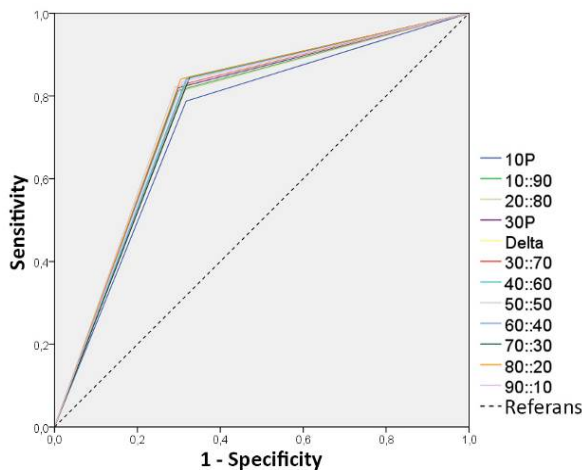
Tablo 2. Farklı oranlarda oluşturulmuş eğitim ve test verisindeki piksel sayıları.

Eğitim veri seti			Test veri seti		
	Olan	Olmayan		Olan	Olmayan
10P	120	120	10P	50938	58556
10%	5963	5011	90%	53665	45095
20%	11926	10021	80%	47702	40084
30P	360	360	30P	50698	58316
Delta	420	420	Delta	50638	58256
30%	17888	15032	70%	41740	35074
40%	23851	20042	60%	35777	30063
50%	29814	25053	50%	29814	25053
60%	35777	30063	40%	23851	20042
70%	41740	35074	30%	17888	15032
80%	47702	40084	20%	11926	10021
90%	53665	45095	10%	5963	5011

b. Lojistik regresyon ile heyelan duyarlılık haritası üretimi

Heyelan duyarlılık haritalarının üretilmesinde farklı eğitim veri seti kullanımı ve etkilerinin araştırılmasını esas alan bu uygulama kapsamında tahmin modeli oluşturulmasında ele alınan ilk algoritma geleneksel yaklaşım olarak bilinen LR metodudur. Belirlenen farklı boyutlardaki eğitim veri setleri kullanılarak lojistik regresyon metodu ile heyelan duyarlılık haritaları üretilmiştir. LR metodu literatürde ve birçok duyarlılık haritası çalışmasında sıklıkla kullanılan bir metot olmakla birlikte ayrıca literatüre önerilen yeni metotların performanslarının karşılaştırılmasında kullanılan bir metottur. Bu nedenle bu çalışmada son birkaç yıldır heyelan duyarlılık haritaları üretiminde kullanılan RO ve karar ağaçları gibi makine öğrenme algoritmalarının performanslarının karşılaştırılmasında LR metodu standart algoritma olarak değerlendirmeye alınmıştır.

ROC eğrileri model performansını test etmek için birçok alanda uygulanmaktadır. ROC eğrileri ile iki sınıf arasındaki ayrımı ve sınıflandırıcının performansını görselleştirmek mümkündür. Bu kapsamda tüm modellerin performanslarının değerlendirilmesi ve görsel olarak yorumlanması amacıyla ROC eğrilerinden yararlanılmıştır. Şekil 3'den de görüldüğü üzere çalışma alanına ait her bir eğitim-test veri seti ile model oluşumu sonrası LR ile üretilen modelin ROC çizgileri verilmiştir. Elde edilen ROC analizi sonrasında LR metodu için en iyi model performansı 80:20 (eğitim:test) oranı ile elde edilen model olduğu belirlenmiştir.



Şekil 3. LR algoritması ile elde edilen ROC eğrisi

Ayrıca üretilen duyarlılık haritalarını sonuç performans ve doğruluk ölçümleri için Genel Doğruluk (GD) ve kappa (K) değerlerine de bakılmıştır. Tablo 3'de verilen değerler

incelendiğinde hem %10 artan ve azalan boyutlarda eğitim-test verisi sonuçları hem de 10p ve 30p kuralına (çok terimli olasılık teorisi) göre elde edilen değerler için sonuçlar verilmiştir. Tablodaki sonuçlar incelendiğinde LR metodu 10p ve 30p verisi ile üretilen eğitim-test veri seti için sırasıyla %73,50 ve %76,07 GD ve 0,470 ile 0,521 kappa değerlerine ulaşmıştır. Bu sonuç artan veri seti boyutu ile birlikte doğruluğun arttığını göstermektedir. Diğer taraftan, eğitim seti boyutunun arttığı test veri seti boyutunun azaldığı durumda ise LR metodu artan eğitim seti boyunca GD değerlerinde artan bir eğilim gösterdiği tespit edilmiştir. Fakat Tablo 3'deki değerler incelendiğinde bu artan eğilimin 30:70 eğitim sonrasında sabit bir ivmede olduğu tespit edilmiştir. Performans sonuçları incelendiğinde LR algoritması ile en yüksek doğruluk değerlerine (%77,36 GD ile 0,521 kappa değeri), envanter verisinin %30'nun eğitim verisi olarak kullanıldığı, geriye kalan %70'lik kısmının test verisi kullanıldığı (yani 30:70 örnekleme oranı) durum için elde edildiği tespit edilmiştir.

Tablo 3. LR algoritması kullanarak farklı eğitim-test verisi için elde edilen GD, kappa (K) ve AUC değerleri

Oran	GD	K	AUC	İşlem süresi (sn)	Olmayan	Olan
10P	73,50	0,470	,735	0,012	0,720	0,748
10:90	75,75	0,515	,758	0,011	0,743	0,770
20:80	75,75	0,515	,758	0,015	0,744	0,769
30P	76,07	0,521	,761	0,015	0,746	0,774
Delta	75,54	0,511	,755	0,016	0,735	0,772
30:70	77,36	0,547	,755	0,012	0,756	0,789
40:60	76,18	0,524	,762	0,018	0,741	0,779
50:50	76,50	0,530	,765	0,013	0,751	0,778
60:40	76,18	0,524	,762	0,016	0,740	0,780
70:30	75,97	0,519	,760	0,016	0,737	0,779
80:20	76,82	0,536	,768	0,016	0,750	0,784
90:10	76,29	0,526	,763	0,017	0,745	0,779

Üretilen heyelan duyarlılık haritaları modellerinin performanslarının değerlendirilmesi amacıyla bir diğer ölçüt olan AUC (Area Under the Curve) değerleri incelenmiştir. Tablo 3'de verilen AUC değerleri incelendiğinde hesaplanan GD ve kappa değerlerine benzer şekilde LR algoritmasının tahmin performansının eğitim veri seti boyutundaki artışa paralel olduğu görülmektedir. Tablodaki tüm AUC performans değerleri incelendiğinde ise en yüksek değeri 80:20 eğitim seti ile %76,80 AUC değeri ile elde edildiği görülmektedir.

Performans değerlendirmeleri sonucunda her bir ölçüt için farklı sonuç ve performans etkileri ortaya çıktığı için optimum eğitim setinin tespitine karar verilmesi amacıyla Wilcoxon's Sıra Toplamı Testi uygulanmıştır. Böylelikle GD, kappa ve AUC değerlerinden sonra optimum setine dayalı model tespiti amacıyla modeller birbirleri ile karşılaştırılarak aralarındaki farkın istatistiksel anlamda farklı olup olmadıklarına karar verilmiştir (Tablo 4). İki doğruluk arasındaki fark için hesaplanan istatistik değerin kritik tablo değerinden (%95 güven aralığı için tablo değeri $z=1,96$) büyük olması durumunda doğruluklar arasındaki farkın istatistiksel olarak anlamlı olduğu, tersi durumda ise anlamsız olduğu söylenebilir. Aşağıdaki tabloda koyu olarak verilen değerler iki model sonucunun birbiri ile istatistiksel olarak anlamlı olduğunu göstermektedir. Diğer bir ifadeyle, tablo sonuçları incelendiğinde özellikle 30:70 eğitim seti ile diğer oranlar için ikili karşılaştırmalar sonucunda hesaplanan istatistik değerlerin, kritik tablo değerinden ($z=1,96$) az olduğu görülmüştür. Bu nedenle bu eğitim seti ile diğer modeller arasında istatistiksel anlamda bir fark olmadığı söylenebilir. Diğer bir deyişle bu model ile kurulan heyelan duyarlılık haritası ile eğitim veri seti oluşturulmasında artan sayıda oranların kullanıldığı diğer durumlarda üretilen haritalar arasında istatistiksel anlamda bir fark olmadığı ifade edilebilir.

Model performans sonuçlarına göre LR yöntemi ile üretilen optimum duyarlılık haritası

Şekil 4'de verilmiştir. Üretilen duyarlılık haritası 5 duyarlılık sınıfında incelendiğinde (çok düşük, düşük, orta, yüksek ve çok yüksek) çalışma alanının kuzey bölgelerinin heyelana karşı çok yüksek ve yüksek oranda duyarlılığı olduğu görülmektedir. Çalışma bölgesinin güneyine doğru duyarlılık alanlarında düşüşler olduğu özellikle güney bölgelerinde en düşük duyarlılığa sahip olan alanların olduğu tespit edilmiştir.

c. Rastgele orman ile heyelan duyarlılık haritası üretimi

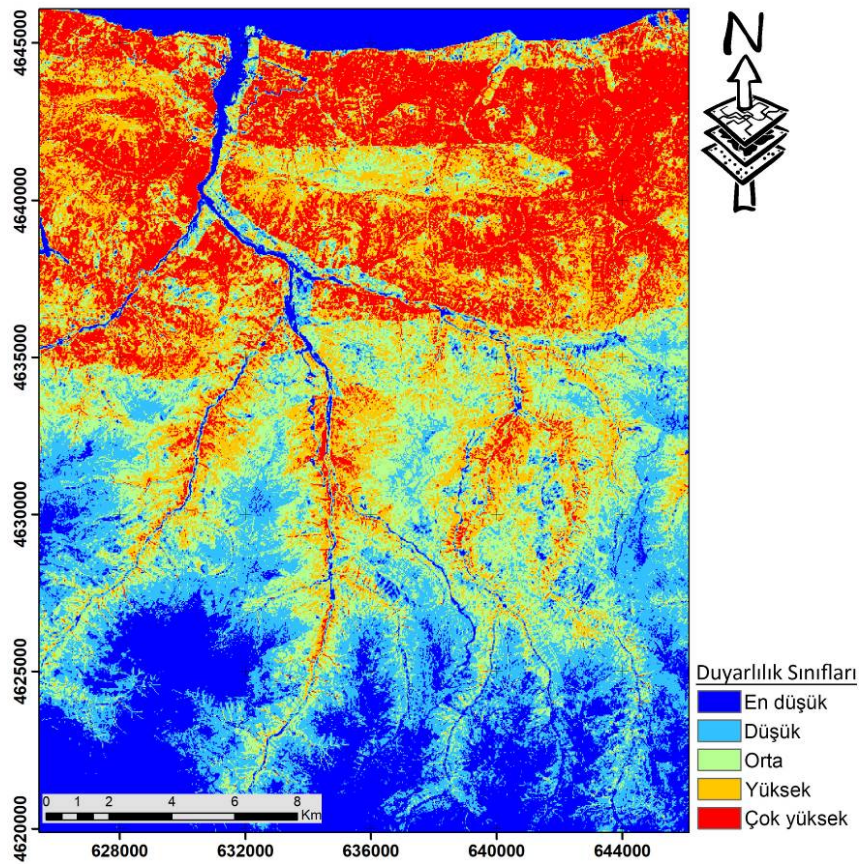
Heyelan duyarlılık haritası üretimi ve eğitim seti boyutunun harita doğruluğuna etkilerinin araştırılması amacıyla kullanılan bir diğer yöntem ise RO algoritmasıdır. Karar ağaçları, RO algoritmasında temel sınıflandırıcı olarak kullanılmaktadır. Çalışmada üretilen her bir eğitim veri seti için farklı boyutlarda ağaç sayısı belirlenmiştir (Tablo 5). Böylelikle daha doğru ve anlamlı sonuçlara ulaşabilmek mümkün olmuştur.

Çalışma kapsamında RO metodu ile üretilen tüm modellerin performanslarının değerlendirilmesi ve görsel olarak yorumlanması amacıyla ROC eğrisi analizi kullanılmıştır. Şekil 5'de görüldüğü üzere çalışma alanına ait her bir eğitim-test verisi sonrası RO ile üretilen modellerin ROC çizgileri verilmiştir. ROC analizi sonucunda RO metodu için 90:10 eğitim seti verisi ile en iyi model performansı elde edildiği tespit edilmiştir.

Tablo 4. LR algoritması için Wilcoxon's Sıra Toplamı Testi

	10:90	20:80	30P	Delta	30:70	40:60	50:50	60:40	70:30	80:20	90:10
10P	1,744	1,925	2,650	1,992	1,992	2,513	2,494	2,440	2,360	3,025	2,654
10:90		0,000	0,299	0,177	0,177	0,396	0,704	0,404	0,200	1,021	0,493
20:80			0,342	0,229	0,229	0,555	0,819	0,516	0,254	1,313	0,687
30P				0,585	0,585	0,122	0,442	0,122	0,124	0,911	0,258
Delta					0,000	0,905	1,152	1,000	0,667	1,852	1,121
30:70						0,905	1,152	1,000	0,667	1,852	1,121
40:60							0,469	0,000	0,426	1,342	0,218
50:50								0,522	0,762	0,522	0,333
60:40									0,535	1,604	0,258
70:30										1,633	0,728
80:20											1,291

Not: %95 güven aralığı için z değerinin 1,96'ten büyük olması tematik harita doğrulukları arasındaki farkın istatistiksel olarak anlamlı olduğunu ifade etmektedir.



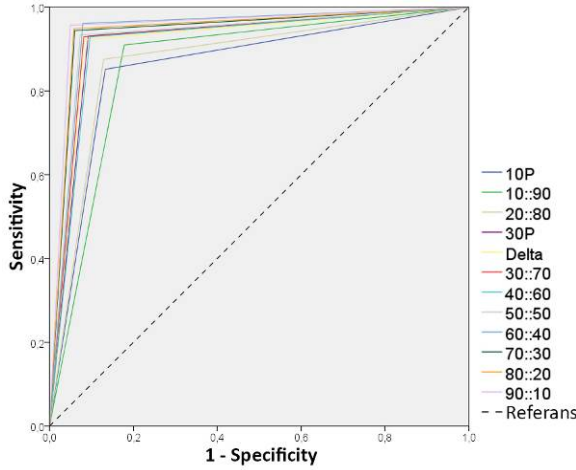
Şekil 4. 30:70 örnekleme oranı kullanılarak oluşturulan LR modeline ilişkin heyelan duyarlılık haritası.

Tablo 5. Rastgele orman algoritmasında her bir eğitim seti boyutu için tanımlanan ağaç sayısı

Eğitim-Test Veri Seti Oranı	Ağaç Sayısı (n)
10P	100
10:90	100
20:80	120
30P	120
Delta	130
30:70	150
40:60	150
50:50	150
60:40	200
70:30	200
80:20	210
90:10	250

Üretilen haritaların performanslarının değerlendirilmesi amacıyla RO metodu için ayrıca GD ve kappa sonuçlarına bakılmıştır. Tablo 6'de farklı oranlardaki eğitim-test veri

setleri için elde edilen doğruluk sonuçları ile 10p, 30p ve çok terimli olasılık teorisi (Delta) yaklaşımlarına göre elde edilen doğruluk sonuçları verilmiştir. Tablodaki sonuçlar incelendiğinde RO metodu 10p ve 30p kuralına göre oluşturulan veri setleri ile sırasıyla %86,59 ve %90,88 GD ve 0,867 ile 0,909 kappa değerlerine ulaşmıştır. Elde edilen bu sonuç artan veri seti boyutu ile birlikte hesaplanan doğruluğun da arttığını göstermektedir. Diğer taraftan, eğitim veri seti boyutunun arttığı test veri seti boyutunun azaldığı durumda ise RO metodunun performansı artan bir eğilim göstermiştir. Tablo 6'deki değerler incelendiğinde LR metoduna benzer şekilde örnekleme oranı olarak 70:30 kullanılması sonrasında, boyuttaki artışla birlikte GD sonuçlarının sabit bir ivmede olduğu, sadece 90:10 verisi için %0,11 oranında bir düşüş olduğu görülmektedir. Performans sonuçları bir bütün olarak ele alındığında en yüksek doğruluk değerlerine (%94,64 GD ile 0,946 kappa) envanter verisinin 70:30 oranında eğitim ve test veri seti olarak örnekleme ile ulaşıldığı tespit edilmiştir.



Şekil 5. RO algoritması ile elde edilen ROC eğrisi

Tablo 6. RO algoritması kullanarak farklı eğitim-test verisi için elde edilen GD, kappa (K) ve AUC değerleri

Oran	GD	K	AUC	İşlem süresi (sn)	Olmayan	Olan
10P	86,59	0,732	0,859	0,466	0,867	0,865
10:90	87,34	0,747	0,866	0,579	0,868	0,879
20:80	87,88	0,758	0,873	0,710	0,877	0,881
30P	90,88	0,818	0,920	0,626	0,909	0,909
Delta	92,38	0,848	0,920	0,669	0,923	0,924
30:70	90,34	0,807	0,924	0,674	0,902	0,905
40:60	91,09	0,822	0,916	0,781	0,911	0,911
50:50	93,03	0,861	0,930	0,890	0,930	0,930
60:40	93,56	0,871	0,941	0,947	0,934	0,937
70:30	94,64	0,893	0,942	1,046	0,946	0,947
80:20	94,64	0,893	0,945	1,144	0,946	0,947
90:10	94,53	0,891	0,954	1,224	0,946	0,945

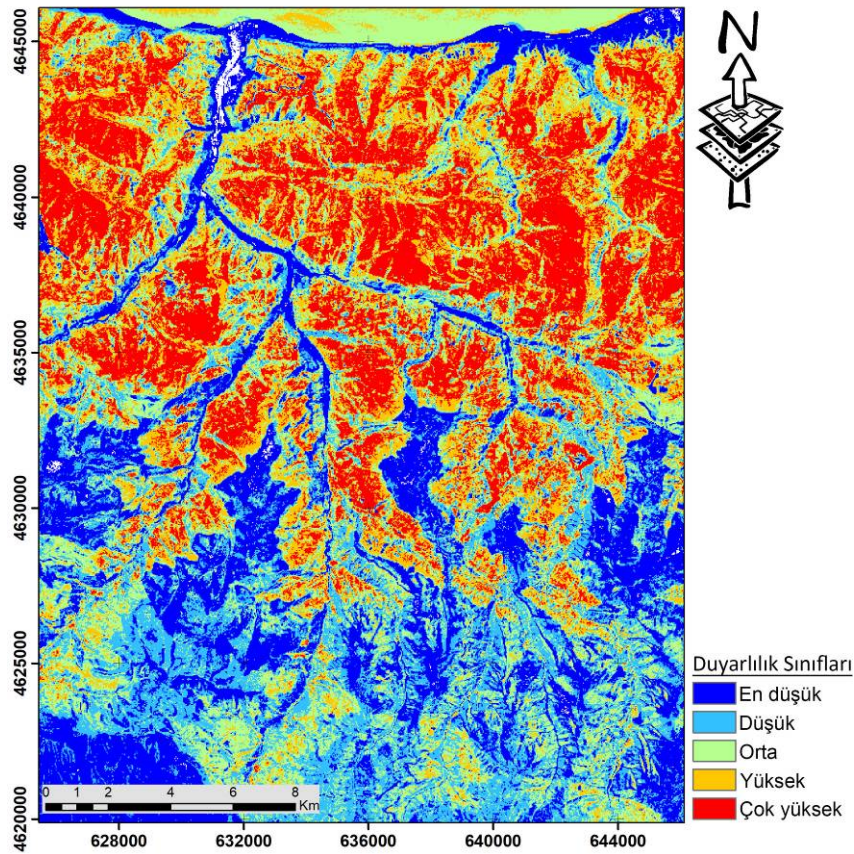
Tablo 7. RO algoritması için Wilcoxon's Sıra Toplamı Testi

	10:90	20:80	30P	Delta	30:70	40:60	50:50	60:40	70:30	80:20	90:10
10P	0,600	1,307	6,184	6,184	6,189	5,381	6,736	7,452	7,514	7,921	8,547
10:90		0,664	5,522	5,213	5,511	5,222	6,124	7,462	7,526	7,633	8,041
20:80			4,279	4,279	4,587	4,472	5,682	6,533	6,746	6,948	7,777
30P				0,000	0,508	0,378	1,270	2,462	2,605	2,954	3,825
Delta					0,894	0,372	1,336	2,500	2,734	3,151	4,202
30:70						0,868	0,775	0,066	2,213	2,673	3,810
40:60							0,982	3,221	3,618	4,117	4,808
50:50								1,622	2,117	2,646	3,889
60:40									2,180	0,756	2,191
70:30										0,688	0,117
80:20											1,706

Not: %95 güven aralığı için z değerinin 1,96'ten büyük olması tematik harita doğrulukları arasındaki farkın istatistiksel olarak anlamlı olduğunu ifade etmektedir.

RO algoritmasının ROC analiziyle görsel olarak değerlendirilen performansı ayrıca AUC skorları dikkate alınarak da değerlendirilmiştir. Tablo 6'de verilen AUC değerleri incelendiğinde GD ve kappa sonuçlarına benzer şekilde RO algoritmasının performansının eğitim veri setindeki artışa paralel olarak iyileştiği görülmektedir. En yüksek AUC değerinin 90:10 olarak adlandırılan eğitim seti ile %95,40 olarak hesaplandığı görülmektedir. Tablo 7'de verilen Wilcoxon's Sıra Toplamı Testi sonuçları incelendiğinde 70:30 örnekleme oranı esas alınarak oluşturulan RO tahmin modeli ile daha küçük boyutta eğitim veri seti boyutları için üretilen RO modelleri arasındaki doğruluk farklılıklarının istatistiksel olarak anlamlı farklar olduğu görülmektedir. Diğer taraftan, 70:30 örnekleme oranından bir sonraki örnekleme boyutları (80:20 ve 90:10) için hesaplanan Wilcoxon's testi sonuçları farkların istatistiksel olarak anlamsız olduğunu göstermektedir. Diğer bir ifadeyle, RO algoritması için 70:30 örnekleme oranı ile üretilen eğitim veri seti ile optimum RO tahmin modeli oluşturulmuştur.

Optimum RO tahmin modeli tüm veri setine uygulanarak çalışma alanı için heyelan duyarlılık haritası üretilmiştir (Şekil 6). Söz konusu heyelan duyarlılık haritası eşit dağılım (quantile) prensibi ile 5 duyarlılık bölgesine ayrılmıştır. Duyarlılık haritası incelendiğinde, bölgenin özellikle kuzey ve orta kısımlarda heyelan olma riskinin çok yüksek ve yüksek olduğu görülmektedir. Diğer taraftan çalışma alanının güney kısımları çok düşük ve düşük derecede heyelan olma riski taşıyan alanlar olarak sınıflandırılmıştır.



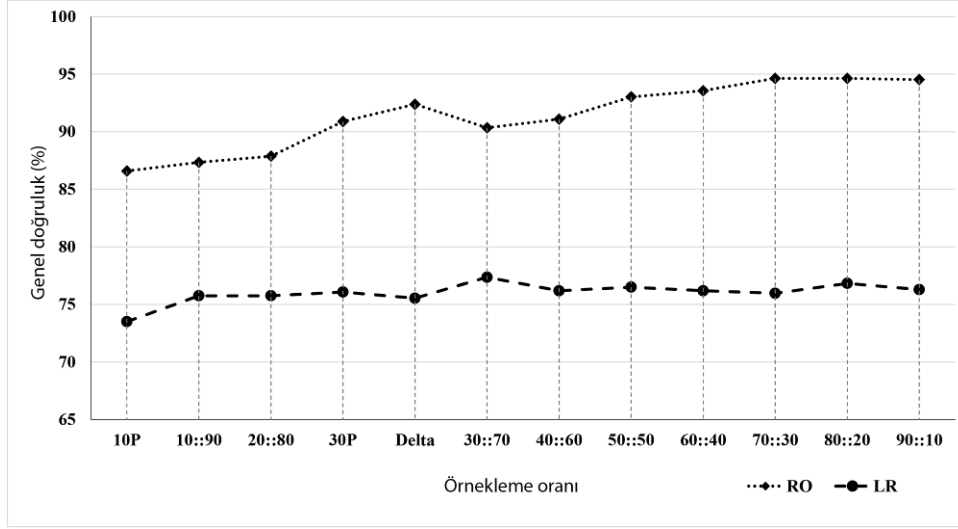
Şekil 6. RO algoritması ile 70:30 örneklem oranı için üretilen heyelan duyarlılık haritası.

ç. LR ve RO performanslarının karşılaştırılması

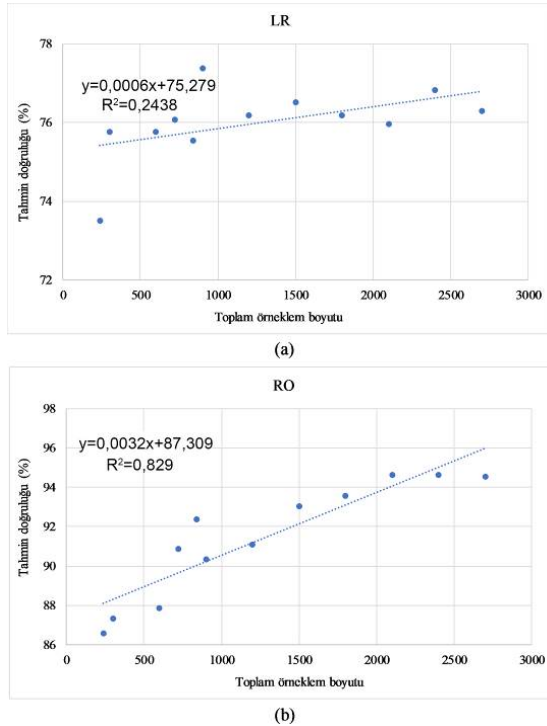
Çalışma kapsamında ele alınan LR ve RO algoritmalarının farklı örnekleme oranları kullanımına ilişkin performanslarını bir arada değerlendirmek amacıyla her iki algoritma için hesaplanan genel doğruluk değerleri Şekil 7'de gösterilmiştir. Şekilden de görüleceği üzere çalışma kapsamında ele alınan 12 farklı örnekleme oranı için literatürde ileri tahmin algoritması olarak nitelendirilen RO algoritmasının klasik tahmin algoritması olan LR yöntemine göre sınırlı sayıda eğitim verisi kullanıldığı durumlarda (10P, 30P, Delta ve 10:90) tahmin doğruluğu açısından %13 daha iyi olduğu, özellikle fazla sayıda örneğin eğitim verisi olarak kullanıldığı (60:40, 70:30, 80:20 ve 90:10) durumlarda RO algoritması ile tahmin doğruluğundaki iyileşmenin yaklaşık %19 seviyelerine ulaştığı görülmektedir. Tabloda verilen performans değişimleri incelendiğinde, artan eğitim veri seti boyutuna karşı LR algoritmasının performansında önemli bir

değişiklik görülmezken, RO algoritmasının performansında %8'e varan seviyelerde artış olduğu görülmektedir. Bu durum LR algoritmasının eğitim veri seti boyutuna karşı daha az duyarlı olduğunu gösterirken, RO algoritmasının performansının istatistiksel olarak belirli bir noktaya kadar (70:30) arttığını, bu noktadan sonra performansının değişmediğini göstermektedir.

Örneklem boyutundaki değişimler ile tahmin algoritmaları arasındaki ilişki hesaplanan doğrusal ve logaritmik eğilim çizgileri kullanılarak ayrı ayrı incelenmiştir. Doğrusal ve logaritmik regresyon denklemleri ile bunlara ilişkin R^2 değerleri hesaplanmıştır. Şekil 8'de LR ve RO algoritmaları için hazırlanan dağılım grafiği ve doğrusal regresyon eğrisi gösterilmiştir. Şekilden de görüleceği üzere LR algoritması için hesaplanan R^2 değeri 0,244 iken RO algoritması için hesaplanan istatistik değeri 0,829'dur.



Şekil 7. Farklı örneklem oranları için LR ve RF algoritmalarının performansı.



Şekil 8. LR ve RO algoritmaları için hesaplanan doğrusal eğilim grafikleri.

Elde edilen bu sonuç, RO algoritmasının tahmin performansı ile örneklem boyutu arasındaki doğrusal ilişkinin oldukça güçlü olduğunu göstermektedir. Örneğin RO ile hesaplanan doğru denklemi (Şekil 8b) $y = 0,0032x + 87,309$ şeklindedir. Burada, x modelin oluşturulması için kullanılacak eğitim veri setinde heyelan ve heyelan olmayan örneklerin toplam sayısını; y ise RO algoritması ile elde edilebilecek doğruluğu göstermektedir. Örneğin

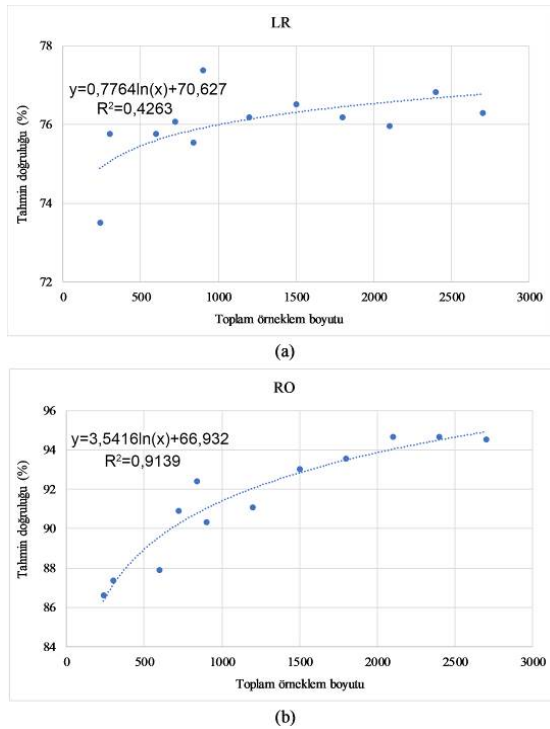
10P kuralına göre oluşturulan eğitim veri setinde toplam 240 örnek bulunmaktadır. Denklem yardımıyla bu veri seti için RO algoritmasının tahmin performansı yaklaşık %88 olarak hesaplanırken, bu örneklem boyutunda algoritmanın hesaplanan performansı %87'dir. LR algoritması için hesaplanan doğru denklemi $y = 0,0006x + 75,279$ şeklindedir.

Örneğin 10P kuralına göre oluşturulan eğitim veri setinde toplam 240 örnek bulunmaktadır. Denklem yardımıyla bu veri seti için LR algoritmasının tahmin performansı yaklaşık %75 olarak tahmin edilirken, bu örneklem boyutunda algoritmanın hesaplanan performansı %73'tür. LR algoritması için hesaplanan doğru denklemi olması gerekenden istatistiksel olarak farklı sonuçlar vermektedir.

Şekil 9'da LR ve RO algoritmaları için dağılım grafiği ve hesaplanan logaritmik regresyon eğrisi gösterilmiştir. Şekilden de görüleceği üzere LR algoritması için hesaplanan R^2 değeri 0,426 iken RO algoritması için hesaplanan istatistik değeri 0,914'tür. Bu durum RO algoritmasının tahmin performansı ile örneklem boyutu arasındaki logaritmik ilişkinin de güçlü olduğunu göstermektedir. Örneğin RO ile hesaplanan doğru denklemi $y = 3,5416 \ln(x) + 66,932$

şeklindedir. Burada, x modelin oluşturulması için kullanacak eğitim veri setinde heyelan ve heyelan olmayan örneklerin toplam sayısını; y ise RO algoritması ile elde edilebilecek doğruluğu göstermektedir. Örneğin 10P kuralına göre oluşturulan eğitim veri setinde toplam 240 örnek bulunmaktadır. Denklem yardımıyla bu veri seti için RO algoritmasının performansının %86,34 olacağı tahmin edilirken, bu örneklem boyutunda

algoritmanın hesaplanan performansı %86,59'dur. LR algoritması için hesaplanan doğru denklemi $y = 0,776 \ln(x) + 70,627$ şeklindedir. Örneğin 10P kuralına göre oluşturulan eğitim veri setinde toplam 240 örnek bulunmaktadır. Denklem yardımıyla bu veri seti için LR algoritmasının tahmin doğruluğunun yaklaşık %75 olacağı kestirilirken, bu örneklem boyutunda algoritmanın hesaplanan performansı %73'tür. LR algoritması için hesaplanan doğru denklemi olması gerekenden istatistiksel olarak farklı sonuçlar vermektedir



Şekil 9. LR ve RF algoritmaları için hesaplanan logaritmik eğilim grafikleri.

5. SONUÇLAR VE TARTIŞMA

Çalışma kapsamında heyelan duyarlılık haritalarının üretilmesinde ve heyelana duyarlı alanların tahmin edilmesinde kullanılacak eğitim veri seti boyutu ve eğitim veri seti boyutundaki değişimlerin tahmin doğruluğu üzerindeki etkilerinin detaylı şekilde analiz edilmesi ve problem çözümüne yönelik bir metodoloji önerilmiştir. Bu amaçla ülkemizde heyelan alanlarının sıklıkla gözlemlendiği Türkiye'nin Orta Karadeniz bölgesinin Sinop ili sınırları içinde yer alan bölge için envanter çalışması yapılmış ve heyelana etki eden faktörler tespit edilmiştir. Ayrıca optimum faktör kümesinin tespiti için Pearson's korelasyonu metodu uygulanmıştır. Bunlara ilave olarak özellikle uzaktan algılama

alanında sınıflandırma işlemi için gerekli olan minimum örnekleme nokta sayısının tespitinde kullanılan 10P, 30P ve Delta kabulleri de esas alınarak örneklem boyutu tespit edilmiştir. Toplamda 12 farklı örneklem boyutu ile eğitim ve test veri setleri oluşturulmuş ve tahmin doğruluğuna etkilerinin araştırılması maksadıyla LR ve RO algoritmaları değerlendirmeye alınmıştır. Çalışma kapsamında elde edilen bulgular aşağıda maddeler halinde sıralanmıştır:

- Heyelan duyarlılık haritalarının üretilmesinde birbirinden farklı ve karmaşık yapıdaki birçok faktörün bir arada değerlendirilmesi ve tahmin modeli oluşumu söz konusudur. Çalışma kapsamında Pearson's korelasyonu ile yapılan analiz sonucunda değerlendirmeye alınan 12 faktörden NDVI, AÖAK, litoloji, eğim ve TNİ faktörlerinin model oluşumunda en etkin faktörler olduğu görülmüştür. Bu faktörlerin yanında bakı, STİ, plan eğriliği, AGİ, profil eğriliği, AOA ve yükseklik faktörleri ise Pearson's analizine göre en az etkin faktörler olarak tespit edilmiştir. Ancak tüm faktörlerin model oluşumuna etkisi olduğundan faktör çıkarımı veya en etkin faktör kümesi belirlemesi yoluna gidilememiş tüm faktörler model oluşumunda değerlendirmeye alınmıştır. Elde edilen bu sonuç literatürdeki birçok araştırma ile paralellik göstermektedir.

- Heyelan duyarlılık haritalarının üretilmesinde eğitim veri seti boyutu ile tahmin doğruluğu arasındaki ilişkinin incelenmesi amacıyla 12 farklı eğitim veri seti boyutu tespit edilmiştir. Söz konusu eğitim veri seti boyutları için LR ve RO ayrı ayrı tahmin modelleri oluşturulmuş ve sonuçlar detaylı şekilde analiz edilmiştir. Uygulama sonucunda özellikle veri seti boyutunun tahmin doğruluğuna etkisi olduğu, özellikle gelişmiş tahmin algoritması olan RO algoritmasının performansının veri seti boyutundaki artışla dayalı olarak paralellik göstererek belirli bir seviyeye kadar arttığı görülmüştür. İstatistiksel analiz sonuçları da özellikle 70:30 oranına kadar RO algoritmasının performansında anlamlı değişimler olduğunu, bu noktadan sonra model performansının istatistiksel olarak benzer olduğunu göstermiştir. Elde edilen bu sonuç literatürde sıklıkla değerlendirmeye alınan 70:30 örnekleme boyutu ile de örtüşmektedir.

- Çalışmada LR algoritmasının performansının, gelişmiş veri madenciliği yöntemlerinden biri olan RF algoritmasına göre istatistiksel olarak oldukça düşük olduğu görülmüştür. Özellikle model oluşumunda sınırlı sayıda örnek kullanıldığı durumda %13, yüksek boyutlu örneklem

kullanılması durumunda yaklaşık %19'lara varan seviyelerde RO algoritmasının daha iyi sonuçlar verdiği görülmüştür. Elde edilen bu sonuç, özellikle son yıllarda literatürde heyelan duyarlılık analizi ve heyelan duyarlılık haritalarının üretilmesi araştırmalarında veri madenciliği yöntemlerinin yaygın bir şekilde kullanılmaya başlamasının temel nedenlerinden biri olup, sonuçlar literatür çalışmalarını destekler niteliktedir.

- Örneklem boyutundaki değişimler ile tahmin algoritmaları arasındaki ilişki hesaplanan doğrusal ve logaritmik eğilim çizgileri kullanılarak ayrı ayrı incelenmiştir. Doğrusal ve logaritmik regresyon denklemleri ve R^2 değerleri hesaplanmıştır. Elde edilen sonuçlar, RO algoritması için doğrusal regresyon esas alındığında hesaplanan R^2 değerinin 0,829, LR algoritması için hesaplanan R^2 değerinin 0,244 olduğunu göstermektedir. Logaritmik regresyon analizi esas alındığında RO ve LR algoritmaları için hesaplanan R^2 değerleri sırasıyla 0,914 ve 0,426'dır. Elde edilen bu bulgular, RO algoritması ile örneklem boyutu arasında güçlü bir doğrusal ilişki olduğunu, bunun aksine LR algoritması ile doğrusal bir ilişkinin olmadığı ifade edilebilir.

- Literatürde uzaktan algılama alanında istatistiksel metotların kullanımında örneklem nokta sayısının tespitinde kullanılan 10P, 30P ve Delta kabulleri ile belirlenen örneklem boyutlarının heyelan duyarlılık analizinde harita doğruluğunun artışına bir etki sağlamadığı görülmüştür.

TEŞEKKÜRLER

Gebze Teknik Üniversitesi, BAP 2018-A101-11 nolu Bilimsel Araştırma Projesi ile desteklenmiştir.

ORCID

Emrehan Kutluğ ŞAHİN  <https://orcid.org/0000-0002-9830-8585>

İsmail ÇÖLKESEN  <https://orcid.org/0000-0001-9670-3023>

KAYNAKLAR

Bai, S. B., Lu, G. N., Wang, J. A., Zhou, P. G. ve Ding, L. A. (2011). GIS-based rare events logistic regression for landslide-susceptibility mapping of Lianyungang, China. *Environmental Earth Sciences*, 62(1), 139-149. doi:10.1007/s12665-010-0509-3

Biesiada, J. ve Duch, W. (2007). *Feature Selection for High-Dimensional Data — A Pearson Redundancy Based Filter*. Berlin: Heidelberg.

Breiman, L. (2001). Random forests. *Machine Learning*, 45(1), 5-32. doi: 10.1023/A:1010933404324

Brenning, A. (2005). Spatial prediction models for landslide hazards: review, comparison and evaluation. *Natural Hazards and Earth System Sciences*, 5(6), 853-862. doi:10.5194/nhess-5-853-2005

Bui, D. T., Ho, T. C., Pradhan, B., Pham, B. T., Nhu, V. H. ve Revhaug, I. (2016). GIS-based modeling of rainfall-induced landslides using data mining-based functional trees classifier with AdaBoost, Bagging, and MultiBoost ensemble frameworks. *Environmental Earth Sciences*, 75(14). doi:110110.1007/s12665-016-5919-4

Carrara A. (1983). Multivariate models for landslide hazard evaluation. *Journal of the International Association for Mathematical Geology*, 15(3), 403-426. doi: 10.1007/Bf01031290

Chen, W., Li, W. P., Hou, E. K., Zhao, Z., Deng, N. D., Bai, H.Y. ve Wang, D. Z. (2014). Landslide susceptibility mapping based on GIS and information value model for the Chencang District of Baoji, China. *Arabian Journal of Geosciences*, 7(11), 4499-4511. doi:10.1007/s12517-014-1369-z

Chen, W., Xie, X., Wang, J., Pradhan, B., Hong, H., Bui, D. T., Duan, Z. ve Ma, J. A. (2017). A comparative study of logistic model tree, random forest, and classification and regression tree models for spatial prediction of landslide susceptibility. *Catena*, 151, 147-160. doi:10.1016/j.catena.2016.11.032

Chung, C. J. F., Fabbri, A. G. ve Van Westen, C. J. (1995). *Multivariate regression analysis for landslide hazard zonation*. A. Carrara, F. Guzzetti (Ed.) Geographical information systems in assessing natural hazards (s. 107–133) içinde. New York: Springer.

Colkesen, I., Sahin E. K. ve Kavzoglu, T. (2016). Susceptibility mapping of shallow landslides using kernel-based Gaussian process, support vector machines and logistic regression. *Journal of African Earth Sciences*, 118, 53-64. doi:10.1016/j.jafrearsci.2016.02.019

- Congalton, R. G. ve Green, K. (2008). *Assessing the Accuracy of Remotely Sensed Data. Principles and Practices*. (2. Baskı) Boca Raton, FL: CRS Press/Taylor & Francis.
- Dai, F. C. ve Lee, C. F. (2002). landslide characteristics and slope instability modeling using GIS, Lantau Island, Hong Kong. *Geomorphology*, 42(3-4), 213-228. doi: 10.1016/S0169-555x(01)00087-3
- Foody, G. M. ve Mathur, A. (2004). A relative evaluation of multiclass image classification by support vector machines. *IEEE Transactions on Geoscience and Remote Sensing*, 42(6), 1335-1343. doi:10.1109/Tgrs.2004.827257
- Gomez, H. ve Kavzoglu, T. (2005). Assessment of shallow landslide susceptibility using artificial neural networks in Jabonosa River Basin, Venezuela. *Engineering Geology*, 78(1-2), 11-27. doi:10.1016/j.enggeo.2004.10.004
- Heckmann, T., Gegg, K., Gegg, A. ve Becht, M. (2014). Sample size matters: investigating the effect of sample size on a logistic regression susceptibility model for debris flows. *Natural Hazards and Earth System Sciences*, 14(2), 259-278. doi:10.5194/nhess-14-259-2014
- Hong, H. Y., Pradhan, B., Xu, C. Ve Tien Bui., D. (2015). Spatial prediction of landslide hazard at the Yihuang area (China) using two-class kernel logistic regression, alternating decision tree and support vector machines. *Catena*, 133, 266-281. doi:10.1016/j.catena.2015.05.019
- Kavzoglu, T. ve Colkesen, I. (2013). An assessment of the effectiveness of a rotation forest ensemble for land-use and land-cover mapping. *International Journal of Remote Sensing*, 34(12), 4224-4241. doi:10.1080/01431161.2013.774099
- Kavzoglu, T. ve Colkesen, I. (2012). *The effects of training set size for the performance of support vector machines and decision trees*. Proceedings of 10th International Symposium on Spatial Accuracy Assessment in Natural Resources and Environmental Sciences (ACCURACY2012), Florianópolis, SC, Brazil.
- Kavzoglu, T., Sahin E. K. ve Colkesen, I. (2014). Landslide susceptibility mapping using GIS-based multi-criteria decision analysis, support vector machines, and logistic regression. *Landslides*, 11(3), 425-439. doi:10.1007/s10346-013-0391-7
- Kavzoglu, T., Sahin, E. K. ve Colkesen, I. (2015a). An assessment of multivariate and bivariate approaches in landslide susceptibility mapping: a case study of Duzkoy district. *Natural Hazards*, 76(1), 471-496. doi:10.1007/s11069-014-1506-8
- Kavzoglu, T., Sahin, E. K. ve Colkesen, I. (2015b). Selecting optimal conditioning factors in shallow translational landslide susceptibility mapping using genetic algorithm. *Engineering Geology*, 192, 101-112. doi:10.1016/j.enggeo.2015.04.004
- Kuncheva, L. I. ve Whitaker, C. J. (2003). Measures of diversity in classifier ensembles and their relationship with the ensemble accuracy. *Machine Learning*, 51(2), 181-207. doi:10.1023/A:1022859003006
- Lee, S. (2005). Application of logistic regression model and its validation for landslide susceptibility mapping using GIS and remote sensing data journals. *International Journal of Remote Sensing*, 26(7), 1477-1491. doi:10.1080/01431160412331331012
- Mather, P. M. (2004). *Computer Processing of Remotely-Sensed Images* (3. baskı). Chichester: Wiley.
- Petschko, H., Brenning, A., Bell, R., Goetz, J. ve Glade, T. (2014). Assessing the quality of landslide susceptibility maps- case study Lower Austria. *Natural Hazards and Earth System Sciences*, 14(1), 95-118. doi:10.5194/nhess-14-95-2014
- Pham, B. T., Bui, D. T., Dholakia, M. B., Prakash, I., Pham, H. V., Mehmood, K. ve Le, H. Q. (2016). A novel ensemble classifier of rotation forest and Naïve Bayer for landslide susceptibility assessment at the Luc Yen district, Yen Bai Province (Viet Nam) using GIS. *Geomatics, Natural Hazards and Risk*, 8(2), 649-671. doi:10.1080/19475705.2016.1255667

- Pradhan, A. M. S. ve Kim, Y. T. (2014). Relative effect method of landslide susceptibility zonation in weathered granite soil: a case study in Deokjeokri Creek, South Korea. *Natural Hazards*, 72(2), 1189-1217. doi:10.1007/s11069-014-1065-z
- Suzen, M. L. ve Doyuran, V. (2004). A comparison of the GIS based landslide susceptibility assessment methods: multivariate versus bivariate. *Environmental Geology*, 45(5), 665-679. doi:10.1007/s00254-003-0917-8
- Şahin, E. K. (2017). *Özellik Seçimi Algoritmaları Kullanılarak Heyelanda Etkili Faktörlerin Belirlenmesi ve Heyelan Duyarlılık Haritalarının Üretilmesi* (Doktora Tezi). YÖK Ulusal Tez Merkezi veri tabanından erişildi. İTÜ, Fen Bilimleri Enstitüsü, İstanbul.
- Tien Bui, D., Tuan, T. A., Klempe, H., Pradhan, B. ve Revhaug, I. (2016). Spatial prediction models for shallow landslide hazards: a comparative assessment of the efficacy of support vector machines, artificial neural networks, kernel logistic regression, and logistic model tree. *Landslides*, 13(2), 361-378. doi:10.1007/s10346-015-0557-6
- Tsai, F. ve Philpot W. D. (2002). A derivative-aided hyperspectral image analysis system for land cover classification. *IEEE Transactions on Geoscience and Remote Sensing*, 40(2), 416-425. doi: 10.1109/36.992805
- Van Westen, C. J. (1993). *Application of Geographic Information Systems to Landslide Hazard Zonation*. Enschede: ITC-Publication.
- Webb, A. R. ve Copsey, K. D. (2011). *Statistical Pattern Recognition* (3. baskı). Chichester: John Wiley and Sons.
- Yalcin, A., S. Reis, Aydinoglu, A. C. ve Yomralioglu, T. (2011). A GIS-based comparative study of frequency ratio, analytical hierarchy process, bivariate statistics and logistics regression methods for landslide susceptibility mapping in Trabzon, NE Turkey. *Catena*, 85(3), 274-287. doi:10.1016/j.catena.2011.01.014
- Yang, Z. H., Lan, H. X., Gao, X., Li, L. P., Meng, Y. S. ve Wu, Y. M. (2015). Urgent landslide susceptibility assessment in the 2013 Lushan earthquake-impacted area, Sichuan Province. China. *Natural Hazards*, 75(3), 2467-2487. doi:10.1007/s11069-014-1441-8
- Yao, X., Tham, L. G. ve Dai, F. C. (2008). Landslide susceptibility mapping based on support vector machine: A case study on natural slopes of Hong Kong, China. *Geomorphology*, 101(4), 572-582. doi:10.1016/j.geomorph.2008.02.011
- Yilmaz, I. (2009). Landslide susceptibility mapping using frequency ratio, logistic regression, artificial neural networks and their comparison: A case study from Kat landslides (Tokat-Turkey). *Computers & Geosciences*, 35(6), 1125-1138. doi:10.1016/j.cageo.2008.08.007
- Yilmaz, I. (2010). Comparison of landslide susceptibility mapping methodologies for Koyulhisar, Turkey: conditional probability, logistic regression, artificial neural networks, and support vector machine. *Environmental Earth Sciences*, 61(4), 821-836. doi:10.1007/s12665-009-0394-9
- Youssef, A. M., Pradhan, B., Pourghasemi, H. R. ve Abdullahi, S. (2015). Landslide susceptibility assessment at Wadi Jawrah Basin, Jizan region, Saudi Arabia using two bivariate models in GIS. *Geosciences Journal*, 19(3), 449-469. doi:10.1007/s12303-014-0065-z