

**T.C.
SAKARYA ÜNİVERSİTESİ
FEN BİLİMLERİ ENSTİTÜSÜ**

**BÜYÜK VERİ ANALİZİNDE HARİTALAMA
(MAPPING) İNDİRGEMESİ İÇİN YENİ BİR YÖNTEM
GELİŞTİRİLMESİ**

YÜKSEK LİSANS TEZİ

Suat ERDOĞAN

Enstitü Anabilim Dalı : ENDÜSTRİ MÜHENDİSLİĞİ

Tez Danışmanı : Doç. Dr. Safiye TURGAY

Haziran 2020

BEYAN

Tez içindeki tüm verilerin akademik kurallar çerçevesinde tarafımdan elde edildiğini, görsel ve yazılı tüm bilgi ve sonuçların akademik ve etik kurallara uygun şekilde sunulduğunu, kullanılan verilerde herhangi bir tahrifat yapılmadığını, başkalarının eserlerinden yararlanılması durumunda bilimsel normlara uygun olarak atıfta bulunulduğunu, tezde yer alan verilerin bu üniversite veya başka bir üniversitede herhangi bir tez çalışmasında kullanılmadığını beyan ederim.

Suat ERDOĞAN

19.06.2020

TEŞEKKÜR

Hayat boyunca desteğini benden esirgemeyen sevgili annem Hatice ERDOĞAN, kız kardeşim Gamze ERDOĞAN'a tüm kalbimle teşekkürlerimi sunarım. Ayrıca yüksek lisans eğitimim boyunca değerli bilgi ve deneyimlerinden yararlandığım, her konuda bilgi ve desteğini almaktan çekinmediğim, araştırmanın planlanmasından yazılmasına kadar tüm aşamalarında yardımlarını esirgemeyen, teşvik eden, aynı titizlikte beni yönlendiren değerli danışman hocam Doç. Dr. Safiye TURGAY'a teşekkürü bir borç bilirim.

İÇİNDEKİLER

TEŞEKKÜR.....	i
İÇİNDEKİLER	ii
SİMGELER VE KISALTMALAR LİSTESİ.....	iv
ŞEKİLLER LİSTESİ	v
TABLolar LİSTESİ	vii
ÖZET.....	viii
SUMMARY	ix
BÖLÜM 1.	
GİRİŞ	1
1.1. Amaç.....	3
1.2. Kapsam	3
BÖLÜM 2.	
KAYNAK ARAŞTIRMASI.....	5
BÖLÜM 3.	
VERİ ANALİZİ	12
3.1. Büyük Veri Yapısı ve Özellikleri	13
3.2. Büyük Veri ve Büyük Verinin Analizi	15
3.3. Hadoop Hdfs (Hadoop Distributed File System)	16
3.4. Veri Sınırlarının Tespiti.....	17
3.5. Örneklem ve Önemi.....	19
BÖLÜM 4.	
VERİLERİN HIZLI SIRALAMASI.....	24

4.1. Veri Haritalama ve İndirgeme.....	31
BÖLÜM 5.	
UYGULAMA	35
5.1. Giriş	36
5.2. Uygulama	36
5.3. Birinci Deneme.....	39
5.4. İkinci Deneme.....	44
BÖLÜM 6.	
SONUÇ	48
KAYNAKLAR.....	51
ÖZGEÇMİŞ	55

SİMGELER VE KISALTMALAR LİSTESİ

5V	: Büyük veri yapısının 5 özelliği.
HDFS	: Hadoop distributed file system.
YARN	: Hadoop yarn.



ŞEKİLLER LİSTESİ

Şekil 1.1. Veri indirgeme işleminde izlenen yol.....	2
Şekil 3.1. Büyük veri kavramına ait temel bileşenler.....	15
Şekil 3.2. Program içinde bulunan seçim, yüzde ve sınır değer isteme ekranı.	18
Şekil 3.3. Sınır değişkenleri.	19
Şekil 3.4. Örnek filtre kodu 1.....	20
Şekil 3.5. Örnek filtre kodu 2.....	20
Şekil 3.6. İstenilen verilerin örneklem içinde süzülmesi.	21
Şekil 3.7. Hızlı sırama algoritması.....	22
Şekil 3.8. Verilerin ortalama değerlerinin hesaplanması.....	23
Şekil 4.1. Hızlı sıralama çalışma görseli.....	24
Şekil 4.2.hızlı sıralama kod parçacığı.	27
Şekil 4.3. Hızlı sıralama algoritması.....	28
Şekil 4.4. Birleştirme metodu kod parçacığı.....	28
Şekil 4.5. Birleştirme algoritması.	29
Şekil 4.6. Kullanılan kütüphaneler.....	30
Şekil 4.7. Hadoop üzerinde iş parçaları tanımlanması.....	31
Şekil 4.8. Haritalama ve indirgeme için örnek gösterim şeması.....	32
Şekil 4.9. Haritalama indirgeme arasında olan sıralama işlemi.....	33
Şekil 4.10. İndirgeme metot görseli.....	34
Şekil 5.1. Uygulama ekranı.....	37
Şekil 5.2. Örnek kullanılan cvs dosya.....	38
Şekil 5.3. Çalışma ekranı.....	38
Şekil 5.4. Sonuç dosyası.....	39
Şekil 5.5. Sadece haritalama birinci deneme.....	40
Şekil 5.6. Haritalama çalışma grafikleri birinci deneme.....	41
Şekil 5.7.sonuç haritalama indirgeme birinci deneme.....	42

Şekil 5.8. Haritalama indirgeme çalışma grafikleri birinci deneme.....	43
Şekil 5.9. Hadoop log.....	43
Şekil 5.10. Sadece haritalama ikinci deneme.....	45
Şekil 5.11. Haritalama çalışma grafikleri ikinci deneme.	46
Şekil 5.12. Sonuç haritalama indirgeme ikinci deneme.	46
Şekil 5.13. Haritalama indirgeme çalışma grafikleri ikinci deneme.....	47
Şekil 5.14. En yüksek önceliğe sahip değer.	47
Şekil 6.1. Veri azaltım grafiği görseli.	48



TABLÖLAR LİSTESİ

Tablo 4.1. Hızlı sıralamamın diğer sıralamalar ile karşılaştırılması. 25

Tablo 4.2. Sıralama algoritmalarının özellikleri. 25



ÖZET

Anahtar kelimeler: Harita indirgemesi ve sıralanması, Büyük veri analizi, Harita sınırlama, Harita indeksleme, Harita verilerinin değer noktası

Günümüzde hızla gelişen teknoloji ile birlikte veri hacmi ve veri paylaşımları da her geçen gün artmaktadır. Büyük verilerin daha hızlı ve etkin işlenerek analiz edilmesi sürecinde, veri haritalama ve indirgenmesi oldukça önemlidir. Büyük veri analizinde veri haritalama dizisi ve küçültme, belirli bir algoritma yapısı kullanarak çalışır ve girdileri bir değer listesine, parametre olarak gönderir. Ara sonuç listesi için girilen sistemde yer alan listedeki tüm değerler dönüştürülerek oluşturulur. Haritalama (Map) işleminde, haritalama listesinin yapısında, veriler kapladıkları alan ve tekrar sayıları dikkate alınarak hızlı sıralama işlemlerine tabi tutulur. Az miktarda olan verinin işlenmesi daha az zaman tüketimi, bellek tüketimi, işlemci tüketimi ve disk tüketimi gibi konularda maliyet azaltıcı etki göstermektedir. Bu çalışmada, önerilen algoritma ile veri sıralaması, veri azaltımı işlemleri daha etkin bir şekilde gerçekleştirilmiştir. Algoritmanın analiz sonuç değerleri sonuç kısmında verilmiştir. Sistem içerisinde veri analizine hazırlanacak olan veriler öncelikle sıralanabilir özellikleri kontrol edilmiş ve daha sonra işlem uygulanmıştır. Çok miktarda veri olması durumunda, veriler çok daha fazla maliyet ile işlenecektir. Azaltım uygulanacak veriler, veri büyüklüğüne ve değerliğine sahip yapıları dikkate alınarak azaltma işleminin her veriye uygulanması sağlanmıştır. Bu işlem örneklerin seçimini kolaylaştırmıştır. Bu işlem ile aynı zamanda örneklerin seçim işlemi de kolaylaşmıştır.

Tez içinde amaçlanan yapıyı gerçekleştirecek bir yazılım geliştirilmiştir. Yazılım Hadoop içindeki sınıflar kullanılarak önerilen algoritma yapısına göre düzenlenerek yeni bir altprogram oluşmuştur. Bunun için Hadoop kütüphanesinden yararlanılmıştır. Çalışma içinde yazılıma aktarılan değer dosyası öncelikle belli sınırlara indirgenmiş ve ardından sıralama ve haritalama yapılmıştır. Haritalama çıktısı içerisinde indeksleme yapılmıştır. Paralel olarak çalışan haritalama ve indirgeme sınıflarında, indirgeme aşamasında veriler değerlendirilmiş ve değerlendirme sonucu oluşan bir çıktı dosyası üretilmiştir.

DEVELOPING A NEW METHOD FOR MAPPING DOWNLOAD IN LARGE DATA ANALYSIS

SUMMARY

Keywords: Map reduce and sorting, Big data analysis, Map bloced, Map index, Map's value point

Today, with the rapidly developing technology, data volume and data sharing are increasing day by day. Data mapping and reduction is very important in the process of analyzing the big data respect of the faster and more efficiently. In big data analysis, data mapping sequence and reduction is worked by using a certain algorithm structure and introduced then sended the inputs to a value list as a parameter. The intermediate result list is created by converting all values in the list included in the entered system. Time of the mapping (Map) process algorithm developed in the structure of the mapping list divide and obtain operations are performed. The sort depends on the bayt value that each data generates. In the case of small volumes of data, the data's result cost reduction, have less time consumption, memory consumption, processor consumption, and disk consumption. In this study, a more effective analysis process has been carried out with the proposed data sorting and reduction algorithm. The data to be prepared for data analysis in the system must have a sortable feature. If there is a large amount of data, the data processed at a much higher cost. The data mitigated must have data size and value, so that the reduction can be applied to each data. This can be facilitated the selection of examples. This process also facilitates the selection of the samples.

The software has been developed to perform the intended structure in this thesis. The software was formed by crushing, added new codes as a procedure and reorganizing classes in Hadoop. Hadoop's library was used for this purpose. In the study, the value file transferred to the software is reduced to certain limits, then sorting and mapping is performed. Indexing used in the printout with the map. Data is evaluated during the reduction phase with paralel running map and reduction classes, and an output file consisting of the evaluation result is produced.

BÖLÜM 1. GİRİŞ

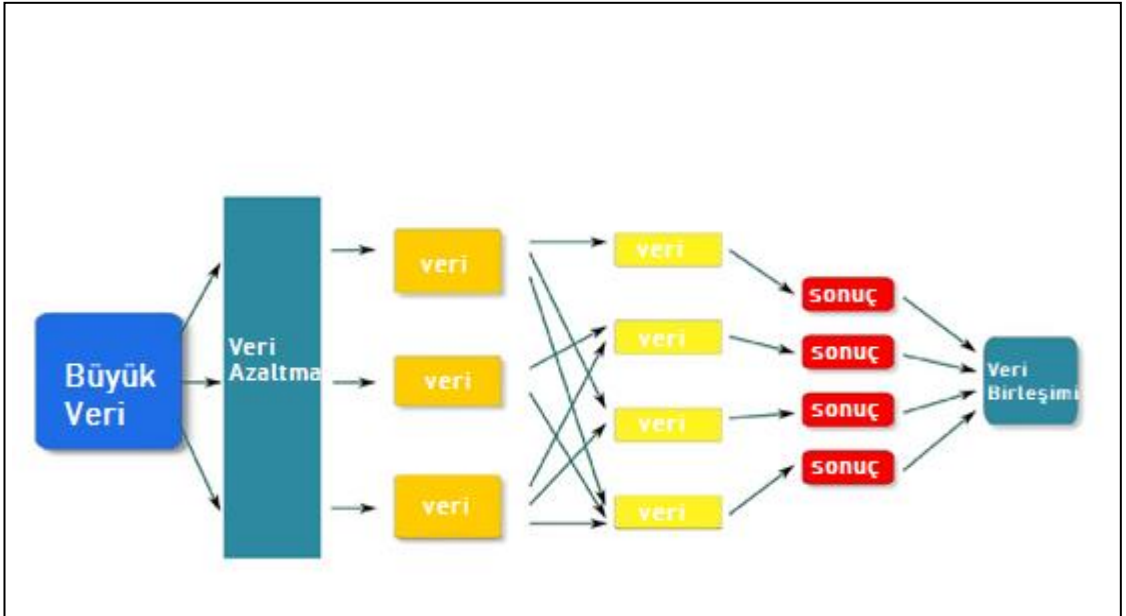
Yaşantımızda, bilgisayar ve internetin yaygın olarak kullanılması en zaruri, ihtiyaçlardan birisi haline gelmiştir. Beraberinde, veri yoğunluğu ve trafiği de çok hızlı bir biçimde, artmaya başlamıştır. Bu yoğun veri akışı içerisinde verilerin analiz edilmesi ve değerlendirilmesi oldukça önemlidir. Veri haritalama, verilerin özelliklerine göre sıralanması ve verilerin indirgenmesi, veri analizi içerisinde en önemli kavramlar arasında yer almakta ve geliştirilen sıralama algoritması ile veri analizine önemli bir katkı sağlamaktadır. Veri kavramı çok eskilere dayanıyor olsa da teknoloji bizlere bu kavramı kullanmamızda daha fazla zorunluluk oluşturmuştur. Günümüzde yaşanan teknolojinin hızlı gelişimi bu kavramında hayatımızın her alanında bizle birlikte olmasını sağlamıştır.

Günümüzde yoğunlaşan bu veriler ayrıca artan veri saklama imkanları ile veri kapasitelerinin inanılmaz boyutlara ulaşmasını sağlamıştır. Büyük veri, sosyal medyadan, bulut teknolojilerinden ve benzeri birçok noktadan sürekli, dinamik olarak durmadan artan yığındır. Verilerin saklanması zaman içinde maliyetinin azalmasının yanı sıra, kapasitenin her geçen gün geometrik olarak artması ile günümüzün sürekli çözümsel yaklaşımlarına karşı, teknolojinin etkin kullanılmasına rağmen, sürekli artan bir hızda çözümlenemeyen bir çözümdür.

Büyük veri, bilginin sınırlarından veya fazlalığından ziyade bu bilgiyi nasıl anlamlı hale getirileceği ile ilgilenmektedir. Kaynaklardan elde edilen verilere uygulanan bazı algoritma ve yöntemler ile daha anlamlı ve işlenebilir cevapların bulunması amaçlanmaktadır.

Haritalama ve indirgeme ilk olarak Google tarafından geliştirilmiş ve yapısal olarak iki temel parçanın paralel olarak uygulanmasından oluşan bir mimari modeldir. Bu

uygulamalar ise haritalama ve indirgeme (küçültme) işlemleridir. Kullanıcı tarafından sağlanan girdi, haritalama işlemi ile veri kümelerine dönüştürülür ve sonrasında analiz edilir. Uygulama işleminde ise haritalama işlemleri tarafından elde edilen çıktılar toplanır ve kullanıcı talebine göre şekillendirilir ve karar kuralı yapısı elde edilir. Hadoop çalışma zamanı, ortamı sağlar ve geliştiricilerin yalnızca giriş verilerini sağlaması, haritayı belirtmesi ve yürütülmesi gereken işlevlerin azaltmasını sağlamaktadır. Böylece verinin, zamana bağlı olan ve olmayan maliyetlerini azaltmaktadır. Apache Hadoop projesinin sponsoru Yahoo!, projeyi veri işlemek için ve hazır bir bulut bilişim platformuna dönüştürmek için büyük çaba harcamıştır. Hadoop bulut sistemi ile veri yığınlarının analizinde mükemmel bir araca dönüşmüştür. Yahoo! akademik kurumlar içinde mevcut olan dünyanın en büyük Hadoop kümesini elinde bulundurmaktadır. Bu çalışmada bu yüzden Yahoo'nun Hadoop yapısı incelenerek, indirgeme algoritması bu uygulama için geliştirilmiştir. Büyük veri ve veri analizi birbirini tamamlayan iki temel bileşendir. Veri sınırlama, analiz öncesi maliyet azalımı için kullanılmıştır. Genel yaklaşımımızda büyük verilerin işlenmesi sürecinde aşağıda yer alan hedefler dikkate alınmıştır. Veri azaltma işlemi, yapılan bu çalışmada ile Şekil 1.1.'de olduğu gibi kullanılmıştır.



Şekil 1.1. Veri indirgeme işleminde izlenen yol.

Bunlar ise;

1. Gerçek zamanlı verilerin temizlenmesi ve istenmeyen her türlü veriden ayıklanması ve bu sayede her türlü donanımda daha az maliyetle işlenmesi
2. İşlenmeden önce gereksiz verilerin tespit edilmesi
3. Daha önemli verilerin öncelikli olarak ele alınması
4. Daha az maliyet gerektiren verilerin öncelikle ele alınması.

1.1. Amaç

Bu tezde verilerin en az maliyetle işlenmesi ve sıralanması amaçlanmıştır. Literatür çalışmalarda ortaya çıkan haritalama ve indirgeme ve sıralanmasına yeni bir yaklaşımla, veri maliyeti göz önünde tutularak katkı sağlamasına çalışılmıştır. Büyük verinin hiç durmadan artan miktarı kaçınılmaz sorunları beraberinde getirmektedir. Veri analizi maliyetini düşürmek için uygun algoritmalar ile veri azaltım yöntemleri kullanılarak sıralanması ile değerli olan veriye ulaşım şansını arttırmaktadır. Ayrıca değeri yüksek olan verilerin, bir veri setinde öncelikli olarak ele alınması için indekslenmesi ve puanlanması, yüksek verilerin ön plana çıkmasında sıralama aşamasında çok önemli ve kritiktir. Bu çalışmada sıralama ve indirgeme işlemini etkin bir biçimde gerçekleştiren algoritma geliştirilmiş ve Hadoop kütüphanesinde yer alan alt program yazılmıştır. İşlem öncesinde haritalama aşamasından önce veriler indirgenmiş (azaltılmış) ve ardından sıralama işlemine geçilmiştir. Her veri kendi boyutuna göre sıralanmış ve tekrarlı veri sayısına göre bulunarak puanlanarak indekslenmiştir. Bu puanlama ve sıralama aşamasından sonra indirgenen devredilen veriler burada önem derecesine göre dört parçaya ayrılmış ve çıktılar elde edilmiştir.

1.2. Kapsam

Bu tez çalışması 6 bölümden oluşmaktadır. 2. bölümde daha önce bu alanda yapılmış çalışmalara değinilmiştir. Bu bölümde aynı zamanda, büyük verinin özellikleri ile

veri sınırlama ve bunun öneminden bahsedilmiştir. 3. bölümde tezde kullanılan ise sınırları tespit edilen verinin sıralanmasından bahsedilerek, veri haritalama ve indirgeme hakkında bilgi sunulmuştur. 4. bölümde sıralama yaklaşımı ele alınmıştır. 5. bölümde yapılan uygulamanın işlem basamakları ve çıktıları verilmiştir. 6. bölümde sonuç kısmında yer almaktadır.



BÖLÜM 2. KAYNAK ARAŞTIRMASI

Büyük veri analizi her geçen gün artan veri trafiğinde artan bir öneme sahiptir. Özellikle bu verilerin saklanması, depolanması ve analiz edilmesi, bu verilerden anlamlı sonuçların çıkartılması açısından oldukça önemlidir. Bununla birlikte, sosyal medyadan, bulut bilişim uygulamalarından sürekli ve dinamik bir biçimde her geçen gün üstel olarak artan bu verilerin saklanması da ayrı bir problemi beraberinde getirmektedir. Günümüzde teknolojinin hızlı bir şekilde gelişmesine karşın çok daha hızlı artan bu verilerin saklanması, depolanması ve analiz edilmesi de yakın zamanda çözümlenemeyen bir problem olarak karşımıza çıkmaktadır.

Veri haritalama ve indirgeme ile ilgili olarak yapılan bazı çalışmalara örnek olarak, Dean ve Ghemawat haritalama ve indirgeme uygulamasını büyük bir emtia makinesi kümeleri ile birlikte çalışarak gerçekleştirmiştir. Amaç yüksek ölçeklenebilirlik ile yapılan hesaplama işlemlerinde binlerce makinede olan verilerin işlenmesidir. Yapıda birçok haritalama ve indirgeme programı uygulanmış ve işlemler her gün Google'ın kümelerinde yürütülerek uygulanmıştır [1]. Yang, Parker ve arkadaşları haritalama-indirgeme ve birleştirme fonksiyonlarını içeren yeni bir model geliştirmiş ve haritalama-indirgeme işlemini de geliştirerek, harita birleştirme aşamasını azaltmış ve zaten bölümlenmiş, sıralanmış verileri verimli bir şekilde bir araya getirerek yeni bir harita yapısı oluşturmuşlardır [2]. Zaharia ve arkadaşları, Hadoop'un performansını, küme düğümlerinin homojen olması sayesinde görevler doğrusal ilerleme sağlamış ve bu varsayımların üzerinde çalışmışlardır. Homojenlik varsayımları her zaman tutmasa da bu yeni bir zamanlama algoritması ile yeni yapısal homojen ortamlarda çözümler önermişlerdir [3]. İndeksleme, arama motoru ve içine veri işleme fonksiyonlarını çözüm yollarını keşfetmek için Hadoop haritalama ve indirgeme ve sıralama yapısını kullanarak verileri verimli bir şekilde haritaya ayırabilen ve modülleri azaltabilen birleştirme aşamasını haritalama ve indirgeme aşamasına eklemişlerdir. Bu işlem

esnasında kullandıkları fonksiyonlar ise, çaprazlama, haritalama ve indirgeme ve birleştirme fonksiyonlarıdır. Bu verimli dizin dosyaları çapraz geçişler ile veri bölümlerini seçer, sayıları sınırlar ve haritalama için verinin azaltılması ve birleştirilmesini savunmuştur [4]. Kolb ve arkadaşları, haritalama-indirgeme programlama modelini, paralel nesne çözümü için kullanarak zorlukları ve olası çözümler için çok girişli sıralama ile en yakın komşu yaklaşımını kullanarak engellemeyi amaçlayan araştırmaları belirlemişlerdir [5]. Verma ve arkadaşları, haritalama-indirgeme modeli arasında bir engel kullanarak harita ve küçültme aşamalarında, iki parçalı programlama ve uygulama ile her ikisinde de basitlik ve kolaylık sağlamıştır [6]. Bununla birlikte birçok durumda, bu engel durumu performansı artırır. Haritalama ve indirgemedeki engelleri ortadan kaldırarak, verimliliğini artıracak şekilde bir yöntem geliştirmek amaçlanmıştır. Goodrich ve arkadaşları haritalama ve indirgeme işlemini algoritmik bir bakış açısıyla araştırmış ve en uygun çözümlerle yaklaşarak sıralama ve çoklu aramada en uygun benzetim modelini tasarlamıştır [7]. Herodotou ve Babu haritalama ve indirgeme işleminde maliyeti minimize etmek için (what if) tahmin motorunu önermiş ve tüm bileşenler için prototip elde etmiştir [8]. Hadoop haritalama-indirgeme sistemindeki her bileşenin etkinliği bir temsilci kullanılarak kapsamlı bir değerlendirme ile gösterilmiştir. Holmes [9] Hadoop yazılımını geliştirerek problem/çözüm biçimi ve kıyaslama paketi geliştirmiş ve haritalama-indirgeme programlarının çeşitli uygulama etki alanlarından bahsetmiştir. Bu uygulama özellikleri ile yüksek/düşük hesaplama ve yüksek/düşük karıştırma hacimlerini dikkate alarak değerlendirme kriterleri geliştirmiştir. Hadoop dağıtımının bir parçası olarak değil dağılımlar ve kriterler, gerekli görevleri azaltmak için geliştirmişlerdir. Slagter ve arkadaşları, yük dengeleme ve bellek tüketimine göre geliştirilmiş bir bölümlendirme işlemini ele alan, yeni bir algoritma geliştirmişlerdir. Önerilen algoritmanın sonucu olarak daha hızlı, daha fazla bellek ve daha doğru sonuçlar ortaya çıkmaktadır [10]. Gopalani ve Arora büyük veri yapısıyla ilgili dünya çapında bir araştırma yapmışlardır. Hadoop haritalama-indirgeme ve yakın zamanda tanıtılan Apache Spark karşılaştırılmıştır. Bu seçeneklerin her ikisi de büyük veri kavramına dayanmakta ve birlikte performansları uygulama altındaki kullanım durumuna göre önemli ölçüde değişiklik göstermektedir [11]. Bu araştırmacılar, bu iki seçeneği kümeleme için standart bir makine öğrenme algoritması (K-en yakın komşu)

kullanarak performans analizlerini karşılaştırmıştır. Gerçek zamanlı senaryoda, kullanılan verilerin hacmi zamanla doğrusal olarak artması durumu incelenmiştir. Sosyal medya başta olmak üzere kontrol edilemeyen verilerin büyümesi sorunsal durumunu yönetmek için büyük miktarda oluşan veri, önerilen yöntem ile dağıtılmış kümeler üzerinden küçük parçalara ayrılmıştır. Paralel olarak veri işlenecek ve sonrada işlenmiş veri kümeleri aralarında toplanmıştır. Haritalama, indirgeme, filtreleme, toplama görevini gerçekleştirmek ve verimli depolama yapısını korumak için kullanmıştır. Önerilen yöntem doğal dil yoluyla duygu analizi gibi teknikler kullanılarak geliştirilmiştir, veriler belirteçlere ve ifade tabanlı kümelere ayrıştırılmıştır. McCreadie ve arkadaşları, haritalama ve indirgemeyi yoğun veri dağıtımı için bir çerçeve olarak önermiştir. Haritalama-indirgeme endeksleme stratejisi değerlendirilmiştir [12]. Sonuç aktarımın en aza indirgenmesinin önemi kanıtlanmıştır.

Hadoop platformu çeşitli uygulama alanlarına ait büyük verilerin analiz sürecinde de kullanılmıştır. Bunlara örnek olarak Bakratsas ve arkadaşları Hadoop platformunu sosyal analizi için kullanmışlardır [13, 37]. Bununla birlikte, Auradkar ve diğerleri Mekansal Hadoop uygulaması ile büyük veri analizinde harita azatlımı uygulaması kullanılmıştır [14]. Lu ve Feng Hadoop'ta resim dosyalarının depolanması ve Hadoop'un resim dosyalarını işleyebilmesi için giriş ve çıkış formatlarının tanıtılması sürecini ele almıştır [15]. Babar ve arkadaşları Hadoop tabanlı akıllı kentsel veri yönetim sistemi önermişlerdir. Bu veri yönetim sistemi özellikle hava kirliliği verilerinin anlık işlenmesi ve hangi saatlerde havanın kirliliği olduğu hakkında bir uyarı vermesi şeklinde çalışmaktadır. Bu çalışmada aynı zamanda paralel algoritma kullanılarak verimli veri yükleme sağlanmıştır [16]. Tallada ve diğerleri Büyük kozmik veri kümelerinin etkileşimli keşfi ve dağıtımı açısından, Hadoop tabanlı bir web uygulaması geliştirilmiştir. CosmoHub ile astronomik verilerin analizi hedeflenmiştir [17].

Bazı çalışmalar ise tamamen Hadoop programının çalışması esnasındaki enerji tüketimi üzerine yoğunlaşmıştır. Bu çalışmalara örnek olarak Hadoop kullanımındaki

enerji tüketimini CPU kullanım sıklık aralığına göre analiz edilmiştir [18, 19]. Wu ve diğerleri, Hadoop'un enerji verimliliği incelenmiştir [20].

Bazı çalışmalar ise veri madenciliği uygulamaları ile birleştirmiştir. Bunlara örnek olarak, Pradeep ve Sundar QUAC mimari yapısını önererek sorgu, kullanıcı ve veri yönetimi modüllerini ele alınarak, bellek kullanımı ve işlem hızını artıran bir mimari yapı tasarlanmıştır. Hadoop'ta güvenli veri depolaması amacıyla bir dizi küme oluşturulmuş ve yönetilmiştir [21]. Bellek kullanımı, zamanlama ve işlem süresi açısından fonksiyonların değerlendirilmesi süreci ele alınmıştır.

Map-Reduce (MR), büyük veri kümelerini işleme yeteneği nedeniyle, tanıtımından bu yana çok popüler hale gelen dağıtılmış bir programlama çerçevesidir. MR, paralel programlamanın birçok yönetim yönü hakkında endişelenmeden dağıtılmış uygulamaları uygulamak için sağlam ve basit bir mekanizma sağlar (örn. İşleri başlatmak, veri dağıtımı, iş senkronizasyonu). Diğer yandan, basitlik ve esnekliği ile Kaynak Tanımlama Çerçevesi (RDF), web-semantiği temsil etmek için çok önemli olan yarı yapılandırılmış ve yapılandırılmamış verileri temsil edebilir. SPARQL, RDF formatında saklanan verileri almayı ve değiştirmeyi amaçlayan bir sorgu dilidir ve ayrıca "Büyük Veri" uygulamalarını desteklemektedir [22]. Büyük veriler, yakalamak, yönetmek, depolamak, dağıtmak ve yüksek hızda farklı yapılara sahip petabyte veya daha büyük boyutlu veri setlerini analiz edilebilir. Büyük veriler yapılandırılmış, yapılandırılmamış veya yarı yapılandırılmış olabilir [23].

Hadoop, büyük miktarda veriyi ucuz ve verimli bir şekilde işlemek için kullanılan açık kaynaklı bir çerçevedir ve iş planlaması, büyük veri işlemede yüksek performans elde etmek için önemli bir faktördür. Bu makale büyük verilere genel bir bakış sağlar ve büyük verilerdeki sorunları ve zorlukları vurgular. Daha sonra Hadoop Dağıtılmış Dosya Sistemi (HDFS), Hadoop Map Reduce ve İş Takibi, Görev Takibi, Ad Düzümü, Veri Düzümü, vb. gibi büyük verilerdeki iş çizelgeleme algoritmalarının performansını etkileyen çeşitli parametreler. Bu yazının temel amacı, iş çizelgeleme algoritmaları ile birlikte karşılaştırmalı bir çalışma sunmaktır. Hadoop ortamında deneysel sonuçlar. Gerçek zamanlı akış verilerinden bilgi keşfi ihtiyacı günümüzde

sürekli artmaktadır ve madencilik modelleri için işlemlerin işlenmesi verimli veri yapılarına ve algoritmalara ihtiyaç duymaktadır. Apache Hadoop kullanarak zaman açısından verimli bir Hadoop CanTree-GTree algoritması öneriyoruz. Bu algoritma, kayan pencere tekniğini kullanarak, gerçek zamanlı işlemlerden tüm sık kullanılan öge setlerini (kalıpları) madenler. Bunlar sık sık kapalı madde kümeleri için madencilik yapmakta ve daha sonra ilişkilendirme kuralları türetilmektedir. CanTree ve GTree olmak üzere iki veri yapısından yararlanır [24]. MapReduce çerçevesini kullanarak karakter tabanlı dize benzerlik birleşiminde veri çarpıklığını ele almak için yeni bir yaklaşım sunulmuştur.ve bu görevi yerine getirmek için MR-Pass Join adlı bir algoritma önerdik. Deneyleri yaptıktan sonra aday çifti oluşturma yönteminin ne yanlış pozitif ne de yanlış negatif sonuç üretmediğini gözlemlenmiştir. Önerilen algoritmanın, benzerlik veritabanı birleştirmesi için uygun sorgu işleme süresi ile ölçeklenebilirlik sağladığını tespit edilmiştir [25].Hadoop kullanımı son yıllarda büyük oranda artmaktadır. Hadoop'un benimsenmesi yaygındır. Yahoo, Facebook, Netflix ve Amazon gibi bazı önemli büyük kullanıcılar, Hadoop'u yapılandırılmamış ve yapılandırılmamış verilerle çok iyi çalıştığı için esas olarak yapılandırılmamış veri analizi için kullanır. Hadoop dağıtılmış dosya sistemi (HDFS) büyük dosyaları depolamak içindir, ancak çok sayıda küçük dosyanın depolanması gerektiğinde, HDFS'deki tüm dosyalar tek bir sunucu tarafından yönetildiği için HDFS'nin birkaç sorunla karşılaşması gerekir. HDFS'de küçük dosyalar sorunu ile başa çıkmak için çeşitli yöntemler önerilmiştir [26]. GOM Hadoop, veri kümelerinin sipariş özelliğinden verimli bir şekilde yararlanmak için yeni bir grup-sipariş-birleştirme mekanizması benimsenmiştir. GOM Hadoop, önerilen mekanizmayı alternatif bir karıştırma mekanizması olarak dahil etmek üzere Hadoop'u genişleterek inşa edilmiştir. GOM-Hadoop'un performansı hem resmi modelleme hem de gerçek dünya veri kümeleri üzerinde kapsamlı deneyler yoluyla değerlendirildi [27]. Lin ve diğerleri hiper bağlantılı üretim iş birliği sistemleri arasında bilgi entegrasyonu ve birlikte çalışabilirlik için bir dizi Hadoop aracında yeni bir yaklaşım önermektedir. Hadoop yaklaşımında veriler “Extracted” web kaynaklarından, “Loaded” olarak “NoSQL” Hadoop'un veri kümesi içine (HBase) tabloları ve “Transformed” kovan şema on okuma ile arzu edilen biçim modeline entegredir. Sözdizimi ve anlam bilim web veri entegrasyonu için Hadoop Özü Yük Dönüştürme (ELT) yaklaşımının küresel akıllı

telefon değer zincirinde benimsenebileceğini gösteren bir vaka çalışması yapılmıştır [28].

MapReduce, bulut bilişim ve büyük ölçekli veri paralel uygulamalarında kullanılan etkili bir programlama modelidir. Hadoop, MapReduce modelinin açık kaynaklı bir uygulamasıdır ve genellikle veri madenciliği ve web dizini oluşturma gibi veri yoğun uygulamalar için kullanılır. Geçerli Hadoop uygulaması, bir kümedeki her düğümün aynı bilgi işlem kapasitesine sahip olduğunu ve görevlerin veri yerel olduğunu varsayar; bu da ek yükü artırabilir ve MapReduce performansını düşürebilir. Bu makale, dengesiz düğüm iş yükü sorununu çözmek için bir veri yerleştirme algoritması önermektedir. Önerilen yöntem, her bir düğümde depolanan verileri, heterojen bir Hadoop kümesindeki her bir düğümün hesaplama kapasitesine göre dinamik olarak uyarlayabilir ve dengeleyebilir. Önerilen yöntem, gelişmiş Hadoop performansı elde etmek için veri aktarım süresini azaltabilir. Deneysel sonuçlar, dinamik veri yerleştirme politikasının yürütme süresini azaltabildiğini ve heterojen bir kümede Hadoop performansını artırabildiğini göstermektedir [29]. Hadoop ve enerji kullanımı ile ilgili olarak yapılan bazı çalışmalar ise [30, 31]'dir.

Yapılan çalışmalarda Pandu Sowkuntla veri azaltım ve sıralama ihtiyacını yönelik çalışması dikkate alınmıştır. Sowkuntla veri azaltmada kaba küme tabanlı yaklaşımla çalışmıştır. Özelliklerine göre verileri değerlendirerek indirgemeye gitmiştir [32]. Verileri indirgemedi maliyeti baz alarak algoritma çalıştırılmıştır. Böylece veri işlemede maliyet kazancını en üste tutmayı amaçlayarak bu noktada gelişme sağlanmıştır. Önerilen algoritmada kullanılan dikey bölümlene şeması kullanılarak veriler özellik alanı üzerinden dağıtılmıştır. Bu yapıda, veriler karıştırılır ve sıralanır. Maliyet kazancı, bu noktada en üst düzeyde ele alan bir yaklaşım geliştirilmiştir. Verileri bayt uzunluklarına göre sıralanmıştır. Sıralamada böl ve yönet yaklaşımını kullanılmıştır. Geliştirilen yapının her noktasında ön bellek kullanımını ve her türlü maliyete dikkat ederek yeni azaltım yaklaşımlarında bulunması amaçlanmıştır. İşlemlerde ve sonuç çıktılarında indeksleme yapılarak maliyet kazancı artırılmıştır. Geliştirmede sık kullanılan verileri ile az maliyetli verileri öncelikli hale getirilmiştir. Yaklaşım ileride eklenebilecek yazılım sınıflarını maliyetini de düşürmüş olmuştur.

Yao ve arkadaşları [33] Uzamsal veri bölümlleme (SDP), uzamsal veriler için dağıtılmış depolama ve paralel hesaplamada yapısını önermişlerdir. Bununla birlikte, uzamsal verilerin çarpık dağılımı ve değişen uzamsal vektör nesneleri hacmi nedeniyle hem uzaysal işlemin optimum performansını hem de kümedeki veri dengesini sağlamak için oldukça güçtür. Bu sorunu çözmek için, Hadoop'ta büyük mekânsal verileri bölümlere ayırmak için mekânsal kodlama tabanlı bir yaklaşımı geliştirmişlerdir. Bu yaklaşım ilk olarak, uzamsal kod, boyut, sayım ve diğer bilgileri içeren bir algılama bilgi kümesi (SIS) oluşturmak için tüm büyük uzamsal verileri, uzamsal kodlama matrisine dayanarak sıkıştırmışlardır. Roy ve diğerleri [34] uydu görüntü işleme, büyük ölçekli görüntülerin tek bilgisayar analizine sınırlama getirmişlerdir. Hadoop Dağıtılmış Dosya Sistemi (HDFS) ise, bu dosyaları doğal veri paralellik tekniği ile ele almak için dikkate değer bir çözüm sunmuştur. Sharma ve Bagga [35] tarafından Hadoop resim işleme ara yüzü önerilmiştir.

Herodotou ve diğerleri [8] performans sağlama ve basitlik amaçlarını farklı şekilde hedef aldık. Yeni bir yaklaşım için çalıştık. Verilerin değerleri sahip oldukları bayt alanı ve veri yığnında oluşan tekrar adedi ile doğru orantılı kabul ettik. Burada kullanıcının en az maliyetli en önemli verilere ulaşması amaçladık. Yang ve arkadaşları haritalama ve indirgeme ve birleştirme eklemişlerdir. Bu noktada bizde sıralama yaklaşımı ekleyerek ek bir yapı koymayı başardık. Gopalani ve Arora [11] K en yakın komşu kullanmıştır performansta büyük artılar sağlamışlardır. Bizde verilerimizi minimize ederek en az maliyete sahip veri kümelerini en çok kullanılan veriler ile doldurduk. Bu noktada veri maliyetini verinin alanı ile ilişkilendiren ve sıralamada bu maliyeti göz önünde tutan bir yaklaşım olmuş olduk. Sıralama performans için bağlı listeler üzerinden gidilmiş. Aynı zamanda sonuçlar indekslenerek o noktada dahi maliyet kazancı göz önünde tutulmuştur.

BÖLÜM 3. VERİ ANALİZİ

Veri analizi ve verilerin etkin bir biçimde işlenmesi ancak veri sıralaması ve indirgemenin büyük veri analizinde etkin kullanımı ile mümkün olabilmektedir. Bulut bilişimde de veri analizi verilerin hacmi ve etkin bir biçimde işlenebilmesi açısından oldukça önemlidir. Bulut bilişim ve veri analizi aynı zamanda veri madenciliğinin temel alarak işlemleri gerçekleştirmektedir.

Veri madenciliğinde amaç anlamlı veya anlamsız verilerin arasında kalan, değerli anlamlı verilerin kullanıma hazır hale getirilmesidir. Verilerin öncelikle işlem maliyetleri çok net bir şekilde azaltmaktadır. Veri madenciliğinde gerçekleştirilen işlem basamakları ise:

1. Veri elde etme
2. Veri güvenliğini sağlama
3. Veri temizleme
4. Veri bütünleştirme
5. Veri indirgeme
6. Veri dönüştürme
7. Veri madenciliği algoritmaları uygulama
8. Yazılım dillerinde test aşamasına ve eğitim aşamasına sokma
9. Sonuçların değerlendirilmesi ve sunulması

Bulut bilişim ve büyük verilerin analizi son yıllarda artan teknoloji ile birlikte iş dünyası ve bilim dünyası olmak üzere uygulaması kaçınılmaz bir alan olarak karşımıza çıkmaktadır. İnternetin gelişmesi ile birlikte web 2.0, sosyal medya ile kullanıcılar arasında paylaşımın gelişmesi, çevrimiçi yayın yapan gazete ve yayın organlarında, dergi gibi kaynaklar ile birlikte bilgi miktarının artması ile birlikte yayınlanan bilgi

miktarı, daha önce görülmedik derecede büyük boyutlara çıkmıştır. Ancak veri miktarındaki bu artış verilerin anlamlandırılabilmesi, gerekli sonuç ve kuralların çıkartılması hususunda mevcut bilgilerin değerlendirilerek, kullanışlı veriler haline dönüştürülmesi gerekliliğini de beraberinde getirmektedir. Verilerin ait oldukları sınıfların belirlenmesi sürecinde geliştirilecek olan algoritma ile birlikte verilerin gruplandırılması ve ilerleyen zaman içerisinde uygun kuralların ortaya çıkartılabilmesi hedeflenmektedir. Makine öğrenmesi, veri madenciliği gibi algoritmaların kullanılması ile aynı zamanda büyük veri analiz işlemi de gerçekleştirilebilmektedir. Bu algoritmalar ancak büyük veri analizi için tasarlanmamışlardır. Dolayısıyla geliştirilecek algoritma yapısı ile birlikte verilerin dinamik yapıda sınıflandırılabilmesi, kümelendirilmesi gerekecektir.

Haritalama ve indirgeme ara bileşeni ile birlikte geliştirilecek program ve algoritma ile aynı zamanda verilerin analiz işlemleri de gerçekleştirilebilecektir. Bir algoritmanın işlerliği, veri miktarını analiz edebilme kapasitesine bağlı olarak değişebilmektedir.

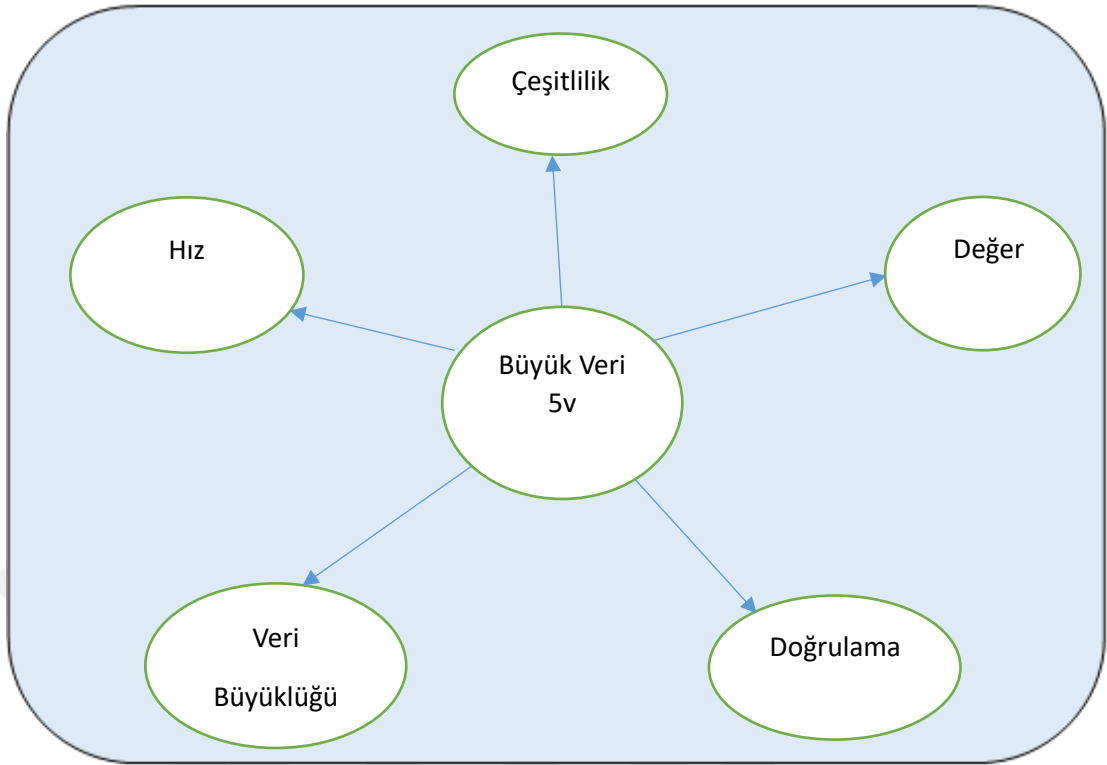
Büyük verilerin analizleri için aynı zamanda büyük kapasiteli bilgisayarlar ve işlemcilere de ihtiyaç vardır. Bulut bilişim mimarisi ile birlikte, veri haritalama ve azatlımı teknikleri ayrıca veri madenciliği algoritmaları ile birlikte bağlantılı olarak çalıştırılacaktır. Bu çalışmada Hadoop haritalama ve indirgeme uygulaması ele alınmış ve bu uygulamanın işlem ve algoritma yapısı ele alınarak incelenmiştir. Gerçekleştirilen uygulama ile bu yapının çalışma biçimi gösterilmiştir.

3.1. Büyük Veri Yapısı ve Özellikleri

Büyük veri, sürekli artan boyutsal miktarda, hacimde ve çeşitlilikle birlikte her geçen gün artan bir yapıya sahiptir. Bu karmaşık ve sürekli büyüyen veri kümelerini çözümlmek, işlemek ve anlamlı sonuçlar çıkartmak günümüzün en büyük problemlerinin başında gelmektedir. Bu tarz problemleri çözmek ve büyük verilerin analizi sürecinde geliştirilen yazılımlardan birisi olan Hadoop haritalama ve indirgemedir. Bu çalışmada Hadoop haritalama ve indirgeme ve performans analizi

uygulamalı olarak incelenmiştir. Bir veri kümesinin büyük veri olarak sınıflandırılması için, Demchenko ve arkadaşları tarafından belirlenen 5 V kavramı; çeşitlilik, hız, hacim, doğrulama ve değer özelliklerini sağlaması gerekmektedir [1]. Bu özellikler ise;

- Çeşitlilik: Verilerin yüzde 80'i yapısal değildir ve teknolojiler, farklı formatlarda veri üretebilmektedirler. Birçok teknoloji bunların yazılım sonuçlarını bir noktada birleşmek zorunda kalabilmektedir. Aynı zamanda veriler birleşebilir ve formatları uyarlanabilir olmalıdır.
- Hız: Veri hızı, her geçen gün artmakla birlikte, verilerin işleme hızı ve kullanım hızında ihtiyaca göre artış göstermekte ve beraberinde teknolojik ilerlemeyi de zorunlu hale getirmektedir. Bu sebepten dolayı verilerin doğru haritalandırması süreci öncesi verilerin sistem içerisine dahil edilmesi önem kazanmaktadır.
- Veri Büyüklüğü: IDC istatistiklerine göre 2020'de yeni ulaşılabilecek veri boyutu, 2009'un 44 katı kadar olacaktır. 2010'lu yıllarda ise dünyadaki toplam bilişim harcamaları yılda %5 artmakta, fakat üretilen veri miktarı %40 olarak artmaktadır. Durum bu şekilde devam ederken verilerin nasıl yapılandırılacağı ve çözümleneceği problemini şimdiden planlamalıyız.
- Doğrulama: Bilgi yoğunluğu içerisinde veri akışı, verinin depolanması ve işlenmesi sürecinde doğrulama sürecinin de gerçekleştirilmesi gerekmektedir. Güvenli bir ortamda verinin analiz edilmesi oldukça önemlidir. Akış sırasında, doğru noktadan kullanıcılar tarafından talep edilen güvenlik kurallarına bağlı olarak alınması gerekmektedir. Bu süreç sonrasında aynı zamanda verinin doğruluğundan emin olunması gerekmektedir.
- Değer: Değer yaratması, yani verilerden belirli bir sonuç ve kural yapısının elde edilebilmesi sürecidir. Eğer veriler değer içermez ise bunları yorumlamak, işlemek veya karşılaştırmak kuşkusuz imkânsız olacaktır.



Şekil 3.1. Büyük veri kavramına ait temel bileşenler.

3.2. Büyük Veri ve Büyük Verinin Analizi

Büyük veri analizini, herhangi bir şey veya bir durum hakkında ne kadar çok şey bildiğinizle hakkında, daha güvenilir bir şekilde yeni yaklaşımlar geliştirebileceğimiz ve gelecekte neler olacağı konusunda tahminlerde bulunabileceğinizi ilkelere dayanır. Çağımız, sürekli artan oranda bir veri yoğunluğu ile karşı olduğumuzu göstermektedir. Sadece bir günlük kısıtlarda daha önceki günden daha fazla veriye sahip oluğumuz kuşku götürmez bir gerçektir. Oluşan bu veri miktarı sadece artmakla kalmayacak kullanılan teknolojileri ileriye taşıyacaktır. Daha hızlı veri analiz eden ve işleyen yapılar oluşturmamız bu konuda gereklidir. Ayrıca veri yoğunluğunda tüm verinin işlenmesi gerekliliği sorusunu kendimize sormalıyız. Örneğin analiz edilecek değerler belli özellikleri içeriyor ise haritalandırma aşamasında diğer verilerin dışarıda bırakılmasını sağlamak zorunda olduğumuz gerçeğinden hareket etmeliyiz. Ayrıca bir büyük veri havuzunun tamamını analiz etmek zorunda mıyız? Bu soruda şartlara bağlı olarak değişecektir. Verileri sınıflandırmamız ve sınırlandırmamız analiz konusunda ihtiyaç duyacağımız kaynak miktarını azaltacaktır. Ayrıca bu alınan veriler ile

karşılaştırma yapılması için oldukça fazla miktarda ve az boyutta olmaları very analizde kaynakların doğru kullanılmasını sağlayacaktır.

Verilerin analiz öncesi ve analiz esnasında dikkate alınması gereken işlemler;

1. Verileri sınıflandırmalı ve sınıf dışında kalan verileri işleme almamalı,
2. Verilerin içinde örneklem alınacağı zaman, çok miktarda olan ve az yoğunluktaki verilere öncelik verilmeli,
3. Veri işlemeye başlamadan önce bu işlemler yapılmalı ve haritalandırılmış veriler ile çalışmalar yapılmalıdır.

3.3. Hadoop Hdfs (Hadoop Distributed File System)

HadoopCommon: Büyük veri işleminde kullanılan tüm Hadoop modüllerini kullanımında olan ortak modülleri kapsamaktadır (<https://hadoop.apache.org>) [38].

Hadoop Distributed File System (HDFS): Büyük oranda olan verilere yüksek iş/zaman oranı işlem sağlayan dağıtık bir dosya yönetim sistemidir. Burada amaç bir çok düğüm üzerinde olan dosya sistemlerini birbiriyle bağlayarak tek bir dosya sistemi gibi düğümlerin çalışmasını sağlamaktadır. HDFS, düğüm noktalarının her zaman çalışmayacağını ve belli zamanlı kesintiler olabileceğini göze ardı etmez. Bu yüzden veri güvenliğini, verinin birçok ta paylaşılarak tutması sağlanmalıdır.

HadoopYarn: İş zaman ve kaynak paylaşımı yapan bir dizi kütüphaneden oluşur.

HadoopMapReduce: Yarn temelli, çok miktarda veriyi paralel olarak işlemeye yarayan bir sistemdir. İş yükünü belli parçalara ayırarak işlemektedir, burada haritalama ve işleme isimden anlaşılacağı gibi iki ayrı parçadır.

Cassandra: Ölçeklenebilir ve çok parçadan oluşan iki parçalı, kümelenmiş veri tabanıdır.

Chukwa: Büyük ve parçalı sistemleri yönetmek veriveri toplama yapabilen sistemidir.

Hive: Veri özetleme ve “ad hoc” sorgu yazma ortamı olan bir veri ambarı yapısıdır.

Pig: Paralel işlemlerin tasarlanması için hazırlanmış, veri akışlarını ve çalıştırmasının planlanması için oluşturulmuş bir dildir.

Veriler sürekli arttığı bir ortamda verilerin iş bölümü ile işlenmesi sağlanmalıdır. Tek bir makine yerine birden fazla makinede işleme ile oluşan çözümler ortaya atılmıştır. Hadoop yazılımı bulut bilişim olanakları ile birlikte, istenilen düzeyde verilerin direk olarak haritalaması, azatlımı konularını da ele almaktadır. Ancak veri azaltıcı sınıflandırma kümeleri standart bir işlem uygulamaktadır. Önermiş olduğumuz çalışma ile veri madenciliği algoritması sisteme ilave edilerek sınıflandırma işleminin gerçekleştirilmesini sağlayan algoritma oluşturulmuştur. Hadoop yazılımında, veriler birden fazla makinede saklanmakta ve HDFS kütüphanesi kullanılmaktadır. Haritalandırılan verilerin paralel işlemlere tabi tutulması ile gerekli analiz işlemleri gerçekleştirilmektedir(<https://hadoop.apache.org/>)[39]. HDFS tarafından geliştirilen haritalama ve indirgeme yöntemi kullanılmış ve iki parçalı işlem sağlanmıştır.

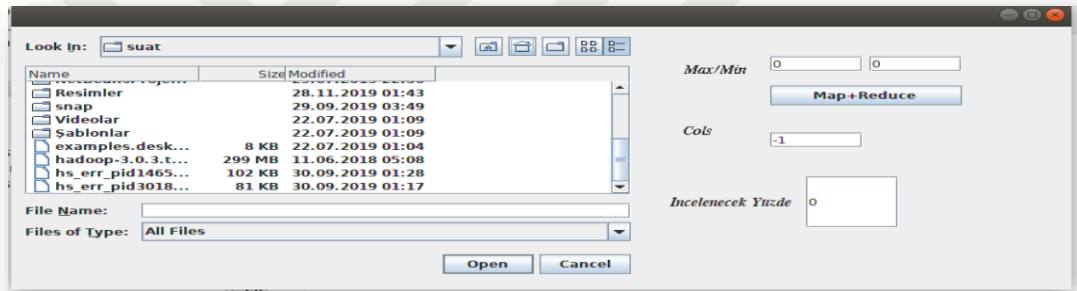
1. Devasa oranda büyük verileri saklayabilir ve işleyebilir.
2. Kaynak kullanımında çok karlı bir minimize kullanılmaktadır. Yatırım maliyetini düşürmektedir.
3. HDFS içerisine veriler aldıktan sonra birçok kaynak üzerinden aynı anda çalışmayabilir.
4. Verilere hızlı bir erişim sunar.
5. Veriler parçalandığı için kolay ve ufak boyutlarda saklanabilir.

3.4. Veri Sınırlarının Tespiti

Veri analizi günümüzde sınırsız verilerin içinde boğulmamızı engelleyerek gerekli analizlerin yapılması ve istenilen sonuç değerlerinin elde edilmesini sağlayan bir

yaklaşımıdır. Aranan verilerin belli sınırları tespit edilebiliyor ise, bu alt ve üst sınırlara uymayan verilerin elenmesi ve haritalandırmadan önce sistem dışına itilmesi ve haritalama aşamasında, hatta özellikle indirgeme aşamasında bize çok sağlıklı bir veri yığını oluşturabilmektedir.

Burada kullanıcı verilerin belli bir yüzdesini incelemek isteyebilir. Sorun en az ve en öncelikli olan verilerin analizi için kullanıcıya verilmesi gerekecektir. Burada en az bayt sahip ve en çok tekrara sahip veriler en önce gelecektir. Ayrıca kullanıcı verilerin bayt olarak aralığını bilebilir ve bu durumda gene kendi tercihi olan bayt aralığında olan verileri öncelikli olarak getiren bir haritalama yapısı bizi uygulamada daha tutarlı ve daha az zaman kaybına dönüşecektir. Şekil 3.2 de ise program içinde seçim, yüzde değer isteme ekranı verilmiştir.



Şekil 3.2. Program içinde bulunan seçim, yüzde ve sınır değer isteme ekranı.

Kullanıcı kolay kullanımı için grafiksel bir ara yüz tasarlanmıştır. Kullanıcı inceleyeceği dosya seçebilir. Dosya seçiminden önce maksimin ve minimum değerleri yazması gerekmektedir. Değerler değiştirilmez yapıda, sıfır olarak kalırsa listedeki tüm elemanlar kabul edilir. Eğer değerler değişmiş ise her eleman bayt olarak bu sayıların arasında olma şartı ile kabul edilir. Şekil 3.3’de ise sınır değişkenlere yer verilmiştir.



Şekil 3.3. Sınır değişkenleri.

Tarafımızdan geliştirilen yapıda verilerin sınırlandırılması bayt oranlarına bağlanmıştır. Bu bilgileri Şekil 3.3. gösterildiği gibi kullanıcıdan alınmıştır. Sınırlama için incelenecek en yüksek ve en alçak veri sınırı ile birlikte örneklem seçilmiş ve bu örneklem toplam verinin yüzde kaçını oluşturduğuna bakılmıştır. Örneklem seçilirken en değerli ve az maliyetli veriler, örneklem için yerleştirecek bir yaklaşım bu tez çalışması ile birlikte geliştirilmiştir. Bundan dolayı sınırlama esnasında en yüksek bayt oranına sahip verileri dışarıda tutulmuştur. Bu işlem ile belli bir sıralamaya şu an için sahip olmayan fakat en değerli verilerden oluşan bir örneklem elde edilmiş olacaktır. Çalışma kod parçasında öncelikle dosya seçimi olması şartı dikkate alınarak kontrol edilmiştir. Daha sonra, ekranda kullanıcı tarafından girilmiş olan ve yazılımda istenmeyen verilerin engellenmesi ile kullanılacak verilerin maksimum ve minimum uzunluğu alınarak işleme devam edilir. Son olarak ise veri yığınının oluşturulması istenen örneklem yüzdesi ele alınmaktadır. Bu yüzde bazında sıfır ile yüz arasından alınan değer oranında örneklem alınacaktır. Bu yüzdesel ifade bize seçeceğimiz örneklem boyutu hakkında bilgi verecektir. Program kod parçası ile hata oluşması durumuna karşı bir tedbir alınarak, blok içine alınarak hataya karşı korunmuştur.

3.5. Örneklem ve Önemi

Tüm veri kitlesi ile çalışmak büyük veri analizinde, büyük verilerden sonuç elde etme sürecinde bazen oldukça zor, zaman alıcı ve maliyetli olabilir. Bu durumlarda, ana kütleyi en iyi yansıtacak bir örneklem seçilmesi oldukça önemlidir.

Örneklem seçiminde, seçilen örneklem hacmi ne kadar sağlıklı olursa, elde edilecek sonuçlar yani ana kütleyi temsil etme yeteneği de o denli kuvvetli olur. Dolayısıyla

ana kütlenin tamamının analiz edilmesine gerek kalmadan, örneklem analizi ile ana kütle hakkında tutarlı bir biçimde tahmin edebilme yeteneğine sahip olmuş oluruz.

Amaç tüm veriyi analiz etmek yerine tüm veriyi en iyi temsil eden örneklemi seçerek, maliyet ve zaman açısından tasarruf sağlamaktır Şekil 3.4 ve Şekil 3.5 örnek filtre kodunu göstermektedir.

```
wordList.stream().filter(item->item.length()<20==true&&
item.length()>0==true).map(item->item).collect(Collectors.toList());
```

Şekil 3.4. Örnek Filtre Kodu 1.

Bu kodlardan da görüldüğü gibi, sınır değerlerin uygulanmasından sonra, uygun örneklem seçimi esnasında, veri örneklemelerinin seçim işleminde yardımcı olmakla birlikte aynı zamanda maliyet açısından da avantaj sağlamaktadır.

```
wordList.stream().map(item->item).filter(item->item.length()<20==true
&&item.length()>0==true).collect(Collectors.toList());
```

Şekil 3.5. Örnek Filtre Kodu 2.

Kod yapısında veri işlemede uygun işlemler sıra ile yapılmalıdır. Tabi bu durumlara göre değişebilir. Öncelikli işlem Şekil 3.4.'de işlemde öncelikle analiz yapılmıştır ve ardından ise haritalama işlemi sonucu aktarılmıştır. Ardından gelen işlem Şekil 3.5.'de işlemler tam tersi bir sırada yapılmıştır. Öncelikle haritalama işlemi gerçekleştirilmiş ve daha sonra analiz işlemi gerçekleştirilmiştir. İkinci kod parçası da maliyet açısından büyük farklar oluşturarak avantajlar sağlamıştır.

Geliştirilen bilgisayar program kodu ile Java'da yazılıp geliştirilen örnek kod yazılımları lambda programı ile yazılmıştır. İki kod parçasında da çalıştırılma sonucu elde edilen çıktı aynı sonuç olacaktır. İlk işlemde, önce filtreme yapılmış ardından haritalama işlemi gerçekleştirilmiştir. İkinci işleminde haritalama sonrası filtreme uygulanmıştır. Çalışma süreleri farklıdır. Sonuç aynı olsa da buradan açıkça anlaşılacağı gibi eğer liste haritalama aşamasında uygun bir filtreleme ile gerekli örneklemelerin arasında oluşturulur ise indirgeme aşamasında işlenecek veri de o denli

azalacaktır. Bu noktada geliştirilen yapıda, bu sebepten dolayı, öncelikle maliyeti yüksek verileri sistem dışına çıkartacağız ve ardından maliyeti ile değeri fazla olan veriler ön plana çekilmiştir. Bununla birlikte maliyetleri daha dip bir noktaya çekmek için verileri daha az saklamaya ve maliyetine uygun bir indeksleme yapılması gerekmektedir. Filtreleme ve sınırlama veri analizinde bu işlem gerekmektedir. Böylece en fazla verim ve az maliyet ile gerekli olan çalışma ortamı sağlanmıştır.

Verilerin önünde ve arkasında olan boşluklar Şekil 3.6.'da görüldüğü gibi temizlenmiş ve her verinin bayt uzunluğu dikkate alınarak işlenmiştir. Bu program parçasında verilerin bayt uzunluğu dikkate alınmıştır. Kullanıcıdan alınan en büyük ve en küçük bayt sınırları ile karşılaştırarak kontrol edilmiştir. Bu verilerin bayt uzunlukları kullanıcı tarafından istenen aralıkta olan veri uzunluğuna içerisinde olmadığında dışlanır yani kullanmak üzere bu veriler tercih edilmez. Eğer giriş yapılan için veri kabul edilen aralıkta ise program içinde olan global bir harita içine eklenir. Eğer burada listede aynı verilerden tekrar eden başka verilerde var ise bu veriler herhangi bir karışıklığa sebep vermemesi için indeks numarası ile kayıt altında tutulur. İndeks adreslemesi ile verinin işlenmesi ve sonuç çıktılarında en az maliyet sağlanması amaçlanmıştır. Ayrıca bu işlemler haritalama işlemi ile birlikte yapılacak sıralamada işlem maliyetini de azaltmıştır.

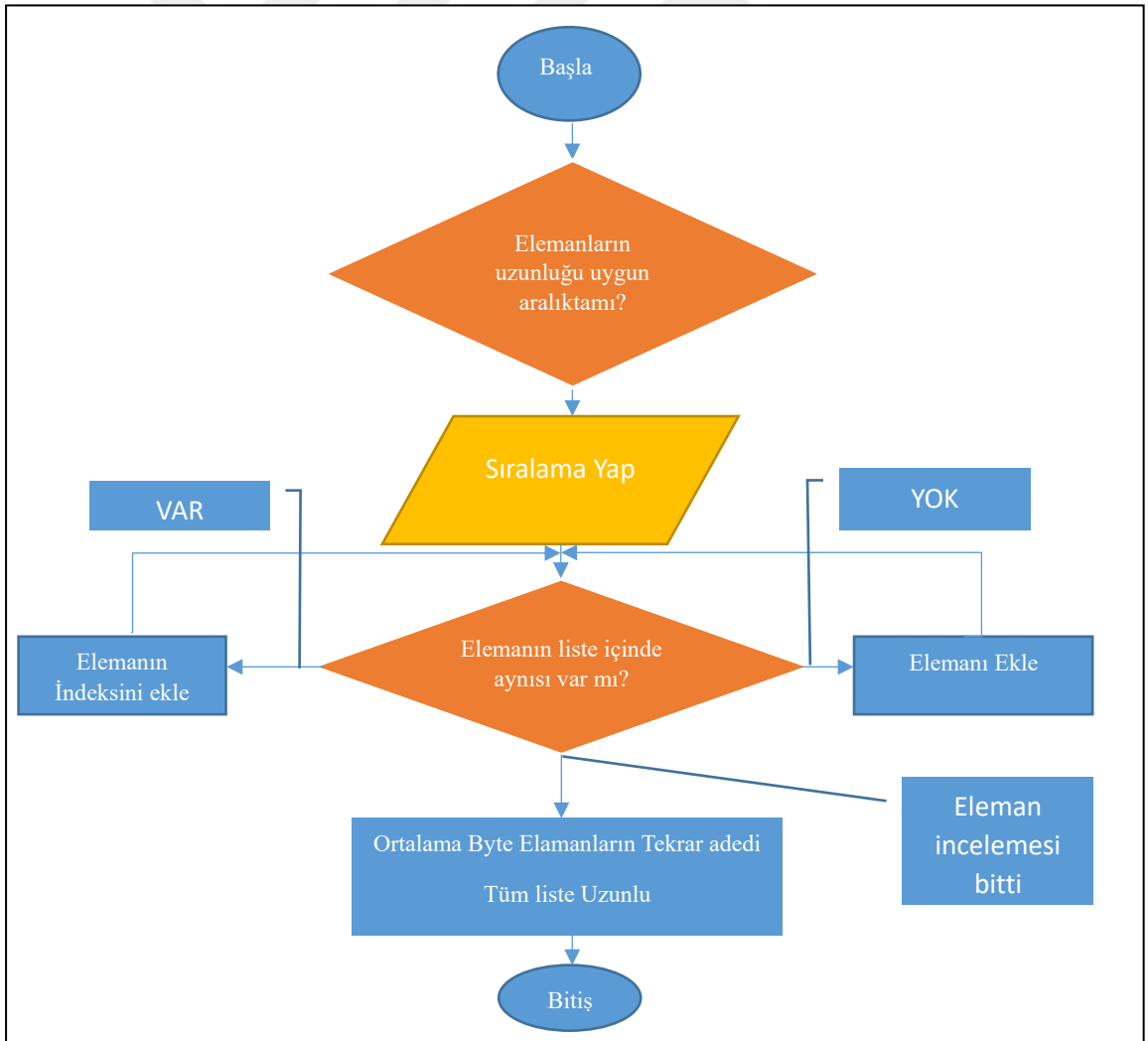
```
while (itr.hasNextToken()) {
    String temp = itr.nextToken();
    int sizeValue = temp.trim().getBytes("UTF-8").length;
    elements.add(temp.trim());
    logicalIndex++;
    if (HadoopFrame.getMaxValue() != 0 && HadoopFrame.getMinValue() != 0) {
        if (temp.trim().length() == HadoopFrame.getMaxValue()) {
            if (elements.contains(temp.trim())) {
                if (mapPermit) {
                    logicalIndex.put(temp.trim(), logicalIndex.toString());
                } else {
                    mapPermit = true;
                    for (Map.Entry<String, String> entry : indexList.entrySet()) {
                        if (entry.getKey().equals(temp.trim())) {
                            entry.setValue(entry.getValue() + 1 + logicalIndex.toString());
                        }
                    }
                }
            } else if (elements.contains(temp.trim())) {
                elements.add(temp.trim());
                indexList.put(temp.trim(), logicalIndex.toString());
            } else {
                mapPermit = true;
                for (Map.Entry<String, String> entry : indexList.entrySet()) {
                    if (entry.getKey().equals(temp.trim())) {
                        entry.setValue(entry.getValue() + 1 + logicalIndex.toString());
                    }
                }
            }
        }
    }
}
```

Şekil 3.6. İstenilen verilerin örneklem içinde süzülmesi.

Aynı veriler, sıralamaya dahil olmayacak böylece işlem daha az yapılacak fakat aynı değerlerin yerleri tutulan indeks adreslerinden kolaylıkla bulunabilecektir. Bir veri dizisinde istenilen değerlerin tekrar sayısının değerin önemi ile orantılıdır. Bu sebepten elemanların işlem frekansı bize örneklem seçiminde kolaylık sağlamıştır. İstenilen şey en önemli, en az maliyetli ve en çok veriyi analiz etmektir. Bu sebepten dolayı

ortalama tekrar adedi ve ortalama bayt uzunluđu baz alınarak sıralama sonuçlanmalıdır. Ortalama bulma işleminde aritmetik ortalama uygun olmuştur.

Bu noktada verilerin analizinde ilgili birçok çalışmada verilerin değeri tespit edilmeye çalışılmıştır. Şekil 3.7’de hızlı sıralama algoritmasına yer verilmiştir. Veri sıralamaya tabi tutulduktan sonra verinin tekrar adedi ile tüm yığın içinde olan değeri tespit edilmiştir. Bu işlem Şekil 3.8.’de gösterildiğı gibi liste içinde olan elemanların hepsini incelenmiş ve ortalama bayt uzunluđu alınmıştır. Her eleman için bir frekans değeri ile puanlama yapılmıştır. Bu puanlama elemanın tekrar sayısının yüzdelik dilimde karşılığıdır. Ardında bu bulunan sonuçlar, kendi geliştirdiğimiz sınıflarımızın yapısında çıktıyazdırılır. Bulunan değeri ile birlikte tüm elemanların uzunluđu sonuç verisi olarak tutulmuş ve son olarak derleyici çıktısına sistem zamanı yazdırılmıştır.



Şekil 3.7. Hızlı sırama algoritması.

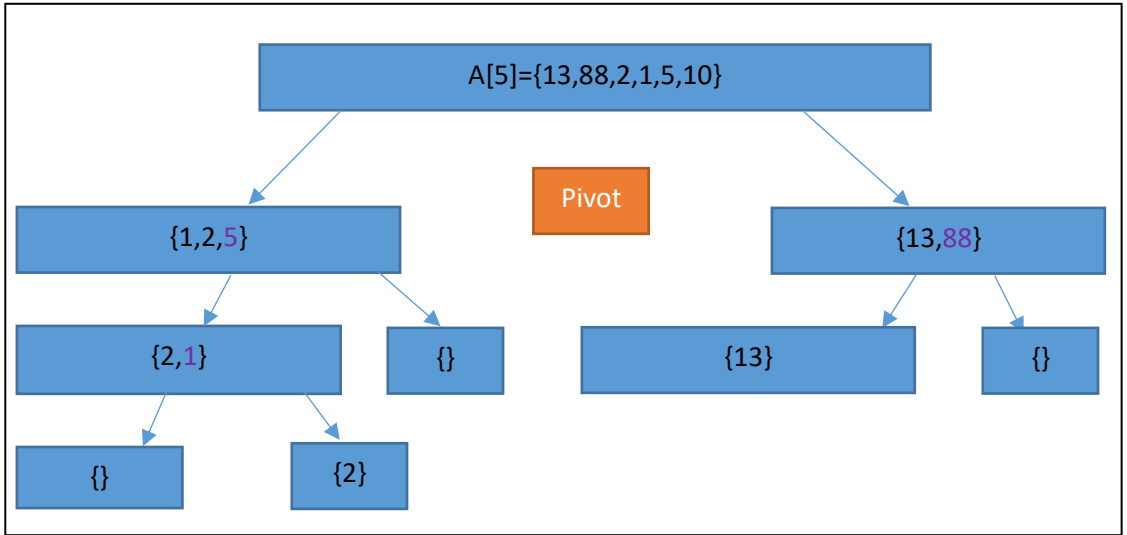


Şekil 3.8. Verilerin ortalama değerlerinin hesaplanması.

Bu çalışmada, bu verilerin analizinde verinin işlem maliyeti ile birlikte değerinin hesaplanması amaçlanmıştır. Bu işlemler esnasında maliyet ve değerin birbirleri ile ters orantılı olduğunu ve birbirleri ile bağlantılı olduğunu gösterilmiştir. Verilerin daha hızlı bir algoritma ile işlenerek sonuçlanması aynı zamanda daha az maliyet sağlayacağı da gösterilmiştir. Geliştirdiğimiz yapıda indeksleme yaklaşımını bu noktada bize katkı sağlamıştır. Zaten indeksleme yapısı ile işlemlere devam etmemiz işlem yapacak diğer sınıfların yük maliyetini ve sonuç çıktılarının oransal boyutunu düşürecektir.

BÖLÜM 4. VERİLERİN HIZLI SIRALAMASI

Verilerin hızlı sıralanması özellikle büyük veri analizinde verilerin işlenmesi ve verilerden anlamlı sonuçlar çıkartılması sürecinde oldukça önemlidir. Böl ve Fethet yapısı ile birlikte algoritma veri kümesini aşamalı olarak bölmekte ve her böldüğü parça üzerine algoritma uygulanarak anlamlı sonuçlar çıkartılmaya çalışılır. Kısaca Böl ve Fethet sıralama sorunun çözmeye yönelik yapılan bir algoritma yaklaşımıdır. Ortaya çıkan problemi benzer problem yapılarına göre parçalıyoruz. Daha sonra bu parçaları kendi içerisinde ayrı ayrı anı işlemlerle çözümledikten sonra bu parçaları birleştirerek gerçek çözüme ulaşıyoruz. Klasik yapıda veriler parçalanır karşılaştırılır ve sıralanır. Bu karşılaştırma örneği Şekil 4.1.'de verilmiştir.



Şekil 4.1. Hızlı sıralama çalışma görseli.

Hızlı sıralama (böl ve yönet) sürecinde veriler seçilen bir anahtar eleman ile işlem gerçekleştirilmekte ve sıralama işlemi bölümlendirme ile özyinelemeli olarak oluşturulmaktadır. Burada değişken durumu anahtar seçimi ile yakından ilişkilidir.

Pivot durumu tamamen verinin yapısına bağlı olarak başta veya sonda bir değer alabilmektedir. Seçilen pivot değerine bağlı olarak veriler parçalanır ve her parça birbiri ile karşılaştırılarak sıralanır. Bu bölme işleminde dolayı hızlı sıralamanın bulundu alana böl ve yönet algoritmaları denir. Hızlı sıralamalı işlemlerde tamamen verinin dinamik yapısı dikkate alınarak veri sıralaması işlemi gerçekleştirilir. Hızlı sıralama, daha kolay uygulanabilmesi için çeşitli veri tipleriyle çalıştırılabilmesi ayrıca da genel olarak diğer algoritmalarından hızlı çalışması sebebiyle oldukça sık kullanılan sıralama algoritmasıdır. Ortalama olarak $O(N \log N)$ yapısı ele alınmıştır. Birleştirme algoritması da hız olarak iyi olmasına rağmen alansal tasarrufta hızlı sıralama kadar başarılı değildir. Bu sebepten sıralamada en uygun olan yapı hızlı sıralama algoritması olacaktır. Tablom 4.1’de hızlı sıralamanın hızlı, birleştirme, eklemeli ve seçmeli tipleri ile birlikte karşılaştırma sonucu verilmiştir. Tablo 4.2’de ise sıralama algoritmalarının özellikleri verilmiştir.

Tablo 4.1. Hızlı sıralamamın diğer sıralamalar ile karşılaştırılması.

	En İyi Zaman	Ortalama Zaman	EnKötü Zaman	Alan
Hızlı	$O(n \log n)$	$O(n \log n)$	$O(n^2)$	$n \log n$
Birleştirme	$O(n \log n)$	$O(n \log n)$	$O(n \log n)$	1
Eklemeli	$O(n)$	$O(n^2)$	$O(n^2)$	1
Seçmeli	$O(n^2)$	$O(n^2)$	$O(n^2)$	1

Tablo 4.2. Sıralama algoritmalarının özellikleri.

İsimlendirme	Özellikler
Sabit	$O(1)$ tekrarsız arama
Loglog	$O(\log \log N)$ sezgisel arama
Logaritmik	$O(\log N)$ İyi azaltılmış bir argoritma bulur ve problemleri parçalıdır
Doğrusal	$N \log N$ $o(N \log N)$ Hızlı bir algoritmadır. N tane veriyi küçük verilere böler.
Karesel	$O(N^2)$ Veri miktarı az olduğu zamanlarda kullanılmalıdır
Kübik	$O(N^2)$ Veri miktarı az olduğu zamanlarda kullanılmalıdır
Üssel	$O(N^2)$ Veri miktarı çok olduğu zamanlarda kullanılmalıdır

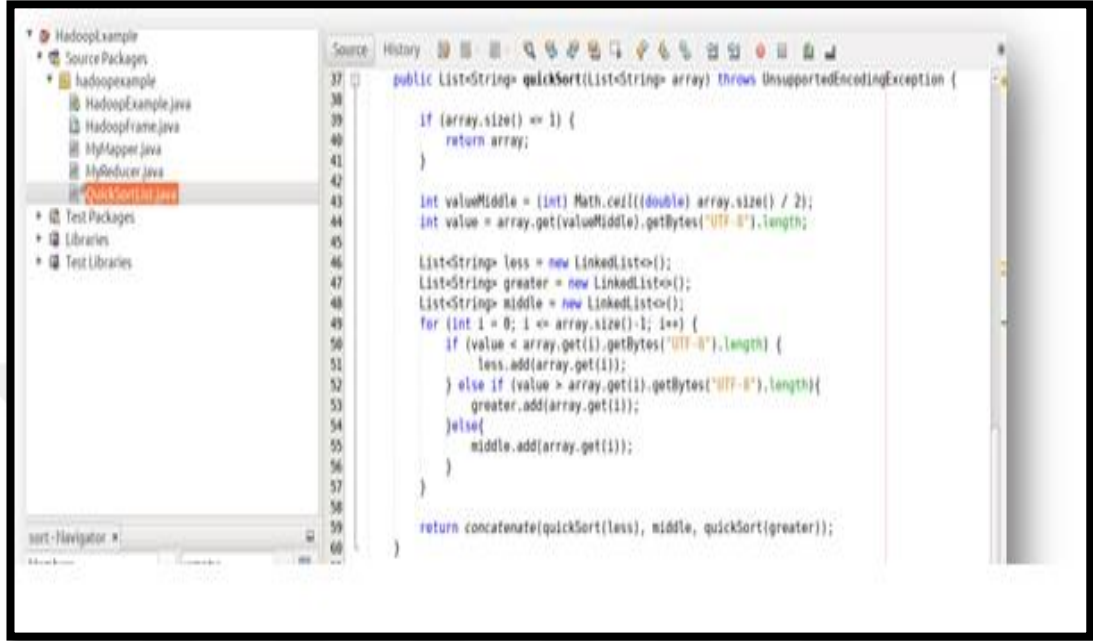
Hızlı sıralamada bayt uzunluğu önemlidir ve bu yapıda pivot ile aynı olan değerlerin sistemden çıkarılarak tekrar sıralamaya gönderilmemesi zaman ve maliyet kazancına yol açacaktır. Bu sebepten sıralamada üç adet veri tutucu geçici diziler oluşturulmuş ve bu dizilerden eşit olanlar sıralamaya gönderilmezken diğer iki dizi sıralanmak için tekrar işleme alınmıştır. İşlemler tekrarlanır ve üç adet yeni dizi oluşturulur, bu döngü sıralama bitene kadar devam etmiştir.

Bu dizi parçalama işlemleri ve sıralama işlemleri tamamlanmasının ardından sırada bu yapıları birleştirmek yer almaktadır. Veride bayt olarak daha az olan değerler başa yazılarak sıralamada daha az maliyete sebep olacak verilerin ön tarafa alınması, incelemeleri kolaylaştıracağı gibi ön görülen zamanda daha çok veriyi değerlendirmemizi sağlayacaktır.

Bu algoritmada $n \log n$ veri alanı kullanılmaktadır. Burada veri kullanımını, minimum süre değerlerine indigeme amaçlanmaktadır. Bu nokta çözüm yöntemi ile adres gösterme olacaktır. Kısaca aktarmak gerekir ise bir değeri listeye eklemek yerine ilgili değerın hafıza üzerinde tutulduğu veri yolu adresinin tutulmasıdır. Bu sebepten Java ile geliştirilen bu yapı içinde dizi listelerinden bağlı listeler kullanılmıştır. Liste içinde her bir ögesi, listede her iki komşu öğeye işaret eden iki işaretçi (adres) içerir. Bu nedenle kaldırma işlemi, yalnızca düğümün iki komşu düğümünde (öge) işaretçi konumunda değişiklik yapılmasını gerektirir. Kısaca verileri parçalar iken sadece düğümlerin içerdiği adresler bağlanıp koparılır. Bu adresleme ve düğümlleme ile çalışan yapı bize bellek tüketimi yüksek yapsa da sürekli olarak değişen sıralamada çok yüksek hızlı çözüm avantajı sağlamıştır.

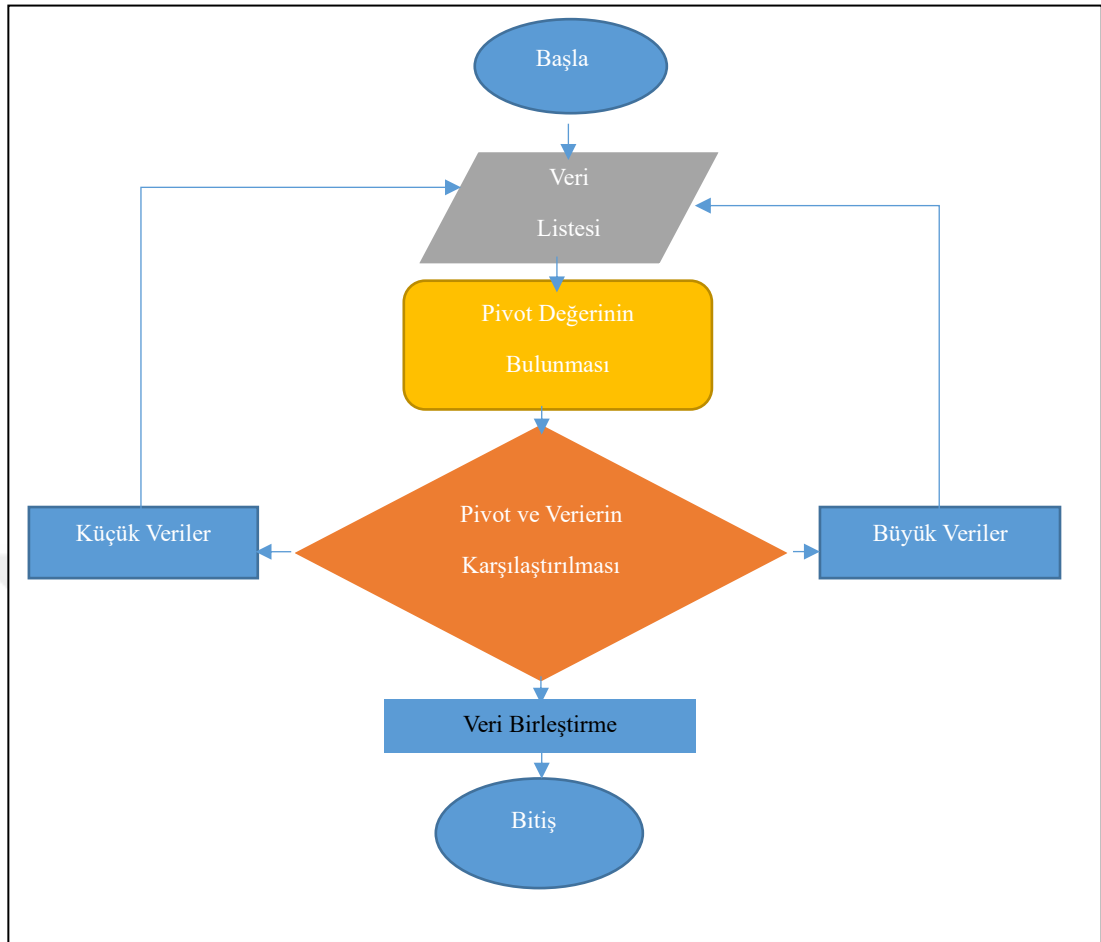
Geliştirdiğimiz bu yaklaşımda veri işleme zaman kazanımına ihtiyaç duyulmaktadır. Bu sebepten dolayı, sıralama esnasında işlem yükünü azaltmamız gerekmektedir. Bu sebep ile işlemde bölme işlemini üç parçalı gerçekleştirilmiştir. Bu parçalardan ikisi tekrar işlem içine girer iken bir parçada sürekli dışarı çıkmaktadır. Bu parçalar verilerin bayt olarak karşılaştırılması ile oluşur. Kısaca pivot değerinden büyük olanlar ile küçük olanlar iki ayrı liste olarak tekrar metot içine alınır. Eşit bayt olarak eşit olan

veriler ise tekrar metod içine alınmayan üçüncü bir listede tutulur. Şekil 4.2’de ise hızlı sıralama kod yazılımı yer almaktadır.



Şekil 4.2.Hızlı sıralama kod parçasığı.

Hızlı sıralamada bağlı listeler yaklaşımı geliştirilmiştir. Bunun sebebi, bize gereken veri çözümleme hızıdır. Eğer bağlı listeler kullanamaz isek sistem daha az bellek kullanmasına karşın daha çok program çalışacaktır. Sebebi ise, bağlı listelerde ön ve arka düğüm bilgilerinin işlem düğümünde saklanmasıdır. Şekil 4.2. Kod yazılımında, öz yenilemeli bir metod uygulanmıştır. Öncelikle liste içinde olan her elemanın adet toplamalarını birden büyük olması şartı kontrol edilmektedir. Ardından veri listesinde olan elemanların ortanca elemanı seçilmekte ve bu elemanın bayt karşılığı pivot değer olarak tutulmaktadır. Kullanılması amacı ile üç ayrı bağlı liste oluşturulmuştur. Listenin her bir elemanı sırayla pivot değer ile karşılaştırılmıştır. Bu karşılaştırılan her bir değer sonucuna göre ilgili listeye eklenmiştir. Bundan sonra büyük ve küçük bayt uzunluğuna elemanların olduğu listeler ele alınır ve tekrardan sıralanması için işleme gönderilmiştir. Her parça sıralaması tamamlısı ile birlikte birleştirme için yeni bir metod içine listeler gönderilmiştir. Şekil 4.3’de hızlı sıralama algoritması yer almaktadır. Şekil 4.4’de ise birleştirme metoduna ait kod yazılımı yer almaktadır.



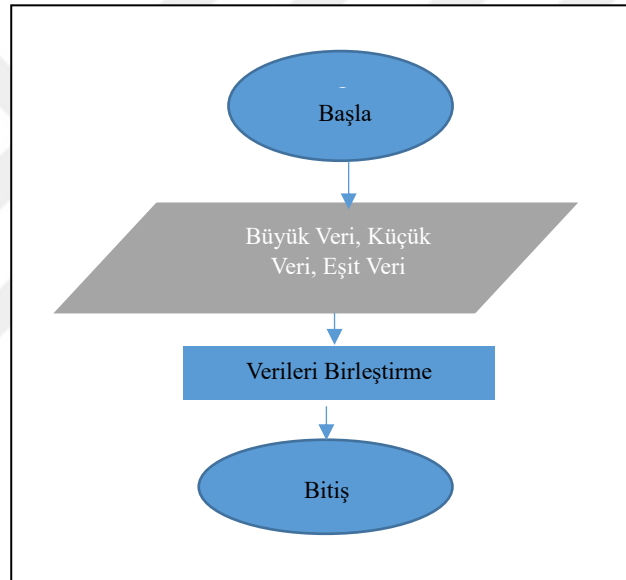
Şekil 4.3. Hızlı sıralama algoritması.



Şekil 4.4. Birleştirme metodu kod parçasığı.

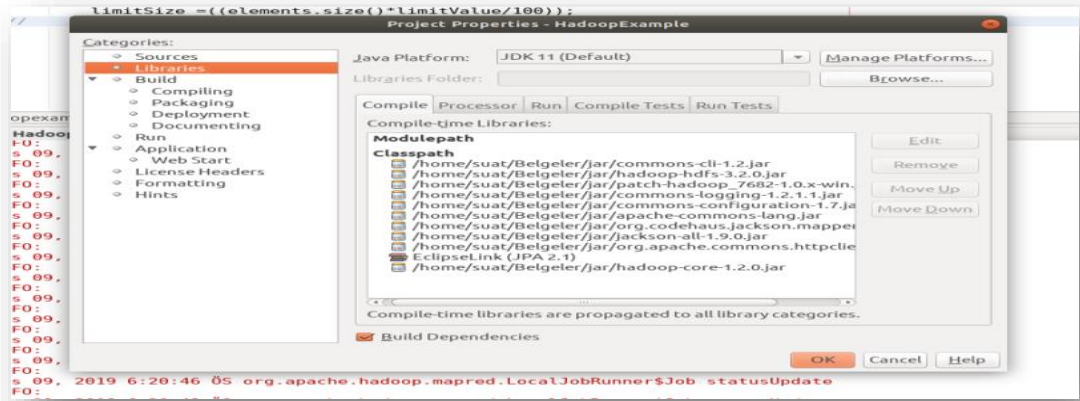
Birleştirme işlememi Şekil 4.4.'te görüldüğü gibi ayrı bir metot içinde yapılmıştır. Bu metot içine gelen listede ise küçük bayt sahibi veriler dizinin en önüne alınmaktadır. Sıralı bu listede en küçükten başlayarak en büyük verilerin sona alınması ile oluşmaktadır.

Java kütüphaneleri Hadoop yazılımı geliştirmek için bize kolaylık sağlamaktadır. Bu kütüphaneleri sisteme ekler ve gerekli noktalarda bu yapılar kullanılmaktadır. Geliştirilen yapıda Hadoop sistemine ait kütüphaneler dışında veri incelemek için kullandığımız başka ek kütüphaneler de vardır.



Şekil 4.5. Birleştirme Algoritması.

Haritalama ve indirgemeye ekleyeceğimiz yeni yaklaşım, Hadoop'a ait sınıflar dönüştürülerek gerçekleştirilmiştir. Bu sebepten dolayı, Hadoop'a ait kütüphaneleri kendi yazımıza ekledik. Şekil 4.6. Geliştirilen yazılım içine eklenen kütüphaneler görülmektedir. Bu kütüphanede hazır metotlar yer almaktadır. Geliştirilen yazılımda bu metotlardan haritalama ve indirgeme metotları yeniden gözden geçirilmiş ve önermiş olduğumuz algoritmaya göre ezilerek, yeniden yazılmıştır.



Şekil 4.6. Kullanılan Kütüphaneler.

Öncelikle Hadoop için sonuç ve veri alımı yapacağımız dosya yollarını belirtmemiz gerekmektedir. Kendi geliştirdiğimiz haritalama ve indirgeme sınıflarını iş yapısı içine aktarmaktayız. İndirgeme işlemi, her durumda çalışması gerekli değildir. Bu sebepten dolayı, indirgeme olmaması için sıfır atanması imkanı hale getirilmiştir. İş yapısına göre anahtar değerleri çıktı tipleri olarak atanmıştır. Hadoop benzer tipten verilerin işlenmesi esnasında çok hassas davranmaktadır. Buralarda gerekli dikkati göstermemiz sağlıklı işleyen bir modül için gereklidir.

Haritalama sıralama daha öncede birçok çalışmada denenmiştir. Biz bu denemelerin bayt olarak sıralanması yaklaşımın verimli olacağını göstermeyi amaçladık. Bu sıralamadan oluşan ek zaman farkını azaltmak için bahsettiğimiz üç yapıllı liste üzerinde bağlı listeler kullanılmıştır. Böylece daha fazla hızlı işlem yaparak maliyet kazancımızı arttırdık. Sıralamada öncelikle maliyeti az olan verileri ön tarafa taşıyarak ileride yapılacak işlemlerimize de maliyet kazancı sağlamış olmaktadır.

Hadoop kısaca yönlendirebileceğimiz ve geliştirebileceğimiz bir yazılımdır. Kendi geliştirdiğimiz sınıflarımızı Hadoop sitemine aktararak geliştirdiğimiz yeni yaklaşımı uygulama fırsatı bulmaktayız. Şekil 4.7. üzerinde gösterilen kod sayfasında öncelikle Hadoop içine belgelin alınacağı ve yazılacağı klasörlerin adresleri verilmiştir. Hadoop'a ait iş yapısına kendi yazdığımız MyMapper sınıfını yerleştirdik ardından

çıkış verilerinin türlerini belirttik. Haritalamada anahtar değer InWritablesonuç değerleri ise Text olacak şekilde ayarlar yapılmıştır. İndirgeme işlemi için MyReduce sınıfını geliştirdik ve ardından Hadoop iş yapısı içine ekledik. Son olarak indirgeme işlemi her defasında yapılmak zorunda değildir. Burada kullanıcı kendi ekranından haritalama sonrasında indirgeme işlemine izin vermez ise Hadoop indirgeme iş yapısını sıfır değeri atanır. İndirgemeye izin verilmiş ise zaten otomatik olarak sistem burada bir sayısını kullanır.

```
try {
    FileInputFormat.setInputPaths(job, new Path("/home/suat/Belgeler/inputDir"));
} catch (IOException ex) {
    Logger.getLogger(HadoopFrame.class.getName()).log(Level.SEVERE, null, ex);
}
FileOutputFormat.setOutputPath(job, new Path("/home/suat/Belgeler/outputDir"));

job.setMapperClass(MyMapper.class);

job.setOutputValueClass(Text.class);
job.setOutputKeyClass(IntWritable.class);

job.setReducerClass(MyReducer.class);
if (!TgBReduce.isSelected()) {
    onlyMap = false;
    job.setNumReduceTasks(0);
}

System.exit(job.waitForCompletion(true) ? 0 : 1);

if (job.isSuccessful()) {
    System.out.println("Job was successful");
} else if (!job.isSuccessful()) {
    System.out.println("Job was not successful");
}
```

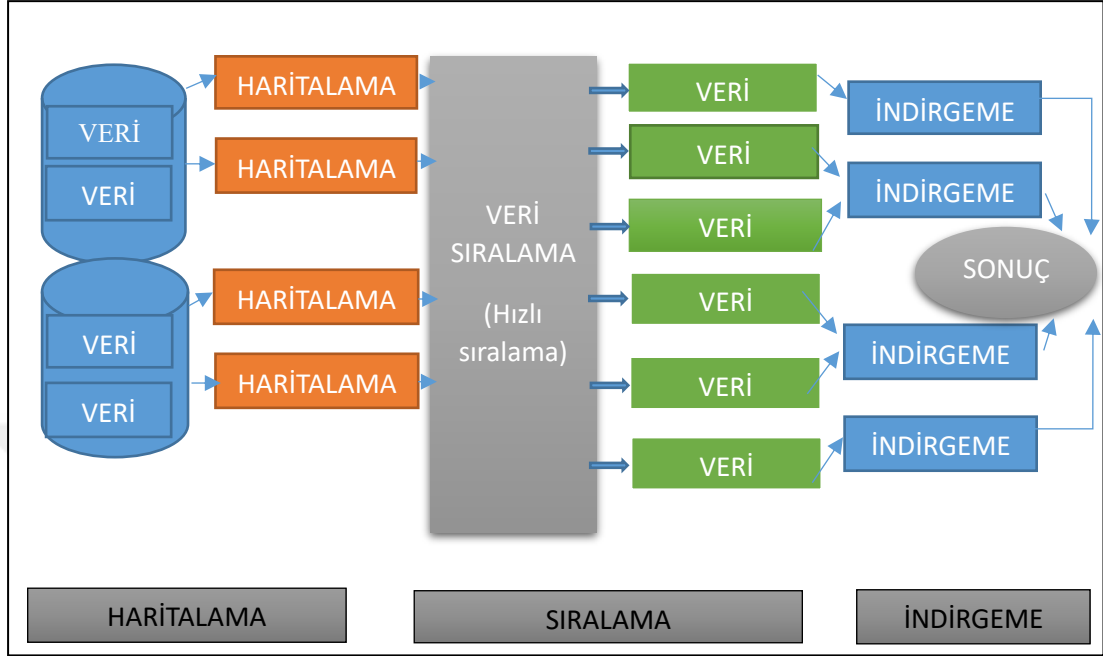
Şekil 4.7. Hadoop Üzerinde İş Parçaları Tanımlanması.

4.1. Veri Haritalama ve İndirgeme

Haritalama ve indirgeme dağıtık mimari yapıda olan ve büyük verilerin kolay bir şekilde analizini sağlayan bir haritalama ve işleme yapan sistemdir. 2004 yılında Google tarafından kullanıma sunulan bu sistem gerçekte 1960'lı yıllarda geliştirilmiş bir yapıdır. Fonksiyonel programlamadaki haritalama ve indirgeme yapısından esinlenerek oluşturulmuştur.

Büyük veri haritalama aşamasında ana düğüm tarafından alınıp daha ufak parçacıklara ayırarak işçi düğümlerine dağıtır. Bu işçi düğümler bu işleri tamamlar ve sonucunu ana düğüme tekrar geri yollarlar. İndirgeme aşamasında ise tamamlanmak istenilen işler tekrar birleştirilerek sonuçlar elde edilir. Şekil 4.8'de haritalama ve indirgeme

olarak deęiřir. Haritalama yapısında aktarılan veriler indirgemed e paralel olarak listelenmiřtir.



řekil 4.9. Haritalama indirgeme arasında olan sıralama iřlemi.

Haritalama iřlemi, daęınık halde bulunan aęda, sunucuda veya disklerdeki verilerin her birinin nerelerde olduęunu önemsemeksizin belli bir metodu o verilere uygular ve kendi diski üzerinde tutmaktadır. Daha sonra indirgeme iřlemi her bir veriyi alır ve belli bir kurala göre analiz eder ve analiz edilmiř veriyi çıktı olarak verir.

Bizim amacımız, haritalama ile tutulan verileri belli bir hizaya kavuřturmaktır. Harita dzenini veri boyutlarına göre yapılırsa istendięi taktirde daha az boyutta ve daha çok veri ile iřlem geręekleřtirilebilir. İndirgeme iřlemini kullanıcı istedięi taktirde tüm verilere uygulamayabilecektir. Bu iřlemlerin sonrasında örneklem seçmek gerekmektedir. Veri içindeki doęal sınırlar, veri analizi vb. iřlemeye da algoritmaları da dikkate alınarak, az kaynak ve az maliyetle çok fazla veri analizi yapabilmektedir.

Verilerin sürekli artıęı bu çağda kaynaklar aynı oranda artmadıęı, maliyetlerin yükseldięi kuřku götürmez bir gerçektir. Örneklem seçimi ve analizinde haritalama sonrası yapılacak bu ara iřlemler ile uygun çözüm sağlanacaktır.

Ayarlar oluşturulur ve Apache Hadoop üzerinden alınan ayarlar bir işleme dönüştürülür. Kütüphane oluşturacak sınıfı çağırılır ve sonrasında haritalamaindirgeme işleme eklenir. Bundan sonra çağırılan sınıflar sonuçları yazan sınıflardır. İşlemin planlamasını haritalama, gerçekleşmesini indirgeme sağlamaktadır. Şekil 4.10'da ise indirgeme alt program çıktısı yer almaktadır.

```

@Override
public void reduce(IntWritable key, Iterable<Text> value, Context context) throws IOException, InterruptedException {
    List<String> newList = Arrays.asList(context.getCurrentValue().toString().substring(1, context.getCurrentValue().toString().length()));
    String tempValue = newList.get(2).trim();
    Double tempFreq = Double.valueOf(newList.get(1).trim());

    if (tempValue.trim().getBytes("UTF-8").length >= HadoopFrame.getAveLenght() && tempFreq >= HadoopFrame.getAveRepeat()) {
        HadoopFrame.truefalse.add(" lenght = "+tempValue.length()+" Value = "+ tempValue);
    } else if (tempValue.trim().getBytes("UTF-8").length <= HadoopFrame.getAveLenght() && tempFreq <= HadoopFrame.getAveRepeat()) {
        HadoopFrame.falsefalse.add(" lenght = "+tempValue.length()+" Value = "+ tempValue);
    } else if (tempValue.trim().getBytes("UTF-8").length <= HadoopFrame.getAveLenght() && tempFreq >= HadoopFrame.getAveRepeat()) {
        HadoopFrame.falsefalse.add(" lenght = "+tempValue.length()+" Value = "+ tempValue);
    } else if (tempValue.trim().getBytes("UTF-8").length >= HadoopFrame.getAveLenght() && tempFreq <= HadoopFrame.getAveRepeat()) {
        HadoopFrame.truefalse.add(" lenght = "+tempValue.length()+" Value = "+ tempValue);
    }

    if (Objects.equals(Integer.valueOf(key.toString()), HadoopFrame.getAllValueListSize())) {
        int newIndex = 0;
        for (String values : HadoopFrame.truefalse) {
            newIndex++;
            context.write(new IntWritable(newIndex), new Text(values));
        }

        for (String values : HadoopFrame.falsefalse) {
            newIndex++;
            context.write(new IntWritable(newIndex), new Text(values));
        }

        for (String values : HadoopFrame.truefalse) {
            newIndex++;
            context.write(new IntWritable(newIndex), new Text(values));
        }

        for (String values : HadoopFrame.falsefalse) {
            newIndex++;
            context.write(new IntWritable(newIndex), new Text(values));
        }
    }
}

```

Şekil 4.10. İndirgeme metod görseli.

İndirgeme işleminiHadoop'a ait sınıfın ezilmesi ile oluşturduk. Burada verilerin önem değerlerine göre bir sıralamaya sahip olmasını sağladık. İşlenen değerlerini ortalama bayt ve ortalama tekrar sayısı ile oluşturduk. Bu değerlere göre dört parçalı bir birleşim listesi de oluşturduk ve çıktı olarak aktardık. Şekil 4.10.'de oluşturduğumuz bu çıktı indirgeme metodu sonucudur. Öncelikle global değişkenler tutulan ortalama baytuzunlu ve ortalama veri tekrar sayısı metod içerisine alınmıştır. İndirgeme işlemine gelen her değerın tekrar sayısı ve sahip olduğu bayt uzunluğu hesaplanır. Öncelik ile sıralama işleminde ortalama tekrar adedini sağlayanlar ve bayt uzunluğu sağlayanlar sonrasında sırası ile, ortalama tekrar adedini sağlayanlar ve bayt uzunluğu sağlamayanlar, ortalama tekrar adedini sağlamayanlar ve bayt uzunluğu sağlayanlar son olarak iki şartı sağlamayanlardan oluşturulmuştur.

BÖLÜM 5. UYGULAMA

Bu çalışmada ana yapı Hadoop üzerine kurulduğu için en verimli alt yapının Linux üzerinde sağlanacağı düşüncesi ile öncelikle Linux üzerinde kurulumlar yapılmıştır. Bu noktada Linux versiyonlarından Ubuntu tercih edilmiştir. Ubuntu <https://ubuntu.com/> adresinden elde edilebilir ve ayrıca Hadoop ise <https://hadoop.apache.org/> üzerinden indirilebilir. Bu iki ana kurulum sonrasında Java sanal makinesi ve Hadoop çalışmasında sorun yaşamamak için ssl kurulumu da yapıldı. Hadoop'a ait kütüphane dosyalarına ihtiyacımız vardır. Bunun sebebi Hadoop sisteminde olan metotları ezip kendi metotlarımız sisteme yerleştirecek olmamızdır. Bu jar dosyaları <https://jar-download.com/> elde edilebilir. Üzerinde çalışma yapılan hdfs-jar dosyası 3.2.0'dır. Bunun dışında birçok yerden bu gerekli yazılım kurulum dosyalarını ve kütüphane dosyalarını elde edebiliriz. Java versiyonlarından JDK 8 kullanılmıştır.

Çalışmada kullanılan örneklem dosyası 2012 yılında agresif bir büyüme planı uygulayan bir sigorta firmasının 2012 yılına ait 36634 kayıttan oluşan veri kümesi Sample insurance portfolio (Örnek sigorta portföyü) isimli dosyadır. Bu veri yığnında on dokuz kolon vardır. Özellikle bu dosya coğrafi kodlamanın seçilebileceği adres bilgisine sahiptir ve dosyadaki mevcut enlem / boylam kullanılabilir. Dosya boyutu 4.124.255 bayt. Dosya formatı cvs'dir. Dosya cvs formatı verilerin virgülle ayrılmasıyla oluşur. Bu ilgili cvs dosyayı <https://ckan.coegss.hhrs.de/> [41] üzerinden elde edebiliriz. Bu dosya numara ve metinlerden oluşmaktadır. Bu dosyada iki uygulama yapılmıştır. Bunlardan birincisi, verileri azaltacağız sadece veri azalttım ile birlikte işlem esnasında yeni değerler elde edilecektir. İkinci denemede ise aynı şekilde veri azaltımı amaçlanmaktadır ve aynı yöntemleri uygulayarak fakat ek olarak daha

fazla çıktı dosyasında olan verilerin indekslemesi veri sonucunda oluşan dosyanın küçültülmesi ile maliyet kazınıcıda arttırılacaktır.

Bu bölümde önerilen hızlı sıralama algoritmasının Sample insurance portfolio üzerindeki uygulamasına yer verilmiştir. Özellikle Sample insurance portfolio tercih edilmesinin nedeni büyük veri yapısını en iyi bir biçimde temsil ettiği ve hızlı bir şekilde gelişen bir bölgeye ait emlak bilgilerinin yer alması sebebiyle aynı zamanda dinamik yapısı da büyük veri için iyi bir örnek oluşturmuş olmasından dolayı seçilmiştir.

5.1. Giriş

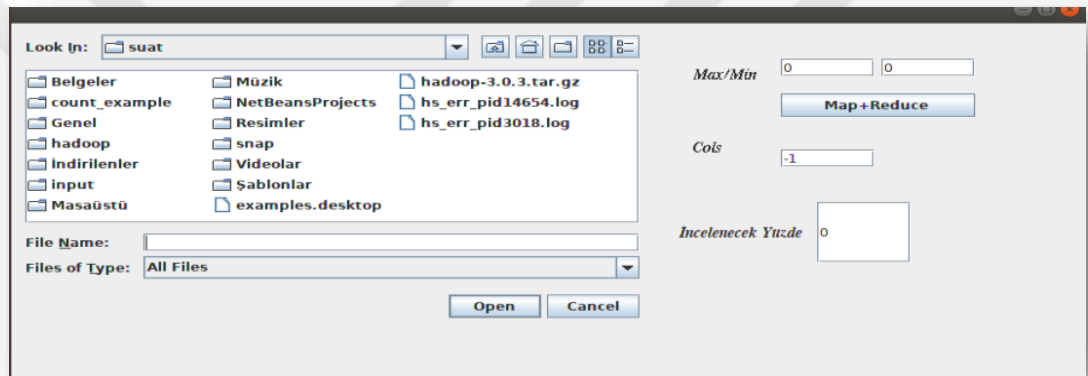
Java dilinde grafiksel kullanıcı arayüzü geliştirilmiştir. Kullanım esnasında Java ait Swing kütüphanesi kullanılarak Hadoop'a ait kütüphaneler ve dosya yönetim kütüphaneleri oluşturulmuştur. Önerilen işlemler için geliştirilen kullanıcı grafik arayüzü Şekil 5.1.'de gösterilmiştir. Bu geliştirmelerde geliştirme Java olarak yapılmıştır. Şekil 5.2'de programın çalışması esnasındaki ekran yer almaktadır. Ön yüzde olan alanların açıklamasını maddeler olarak halinde belirtilmiştir.

5.2. Uygulama

Uygulama çalışma boyunca öne sürülen önerilerin soyut bir noktadan somut bir noktaya taşımıştır. Bu noktada verilerin haritala öncesi sınırlanması ve sonrasında sıralanması incelenebilecektir. Geliştirilen uygulama özelliklerinin kullanıldığı ve kullanılmadığı durumlarda gözlenmesi için iki ayrı deneme yapılmıştır. Hadoop yapısında çok yararlı bir yapı olarak geliştirilen haritalama ve indirgeme yeni yaklaşımımız ile daha verimli ve az maliyetli bir yapıya kavuşulacaktır.

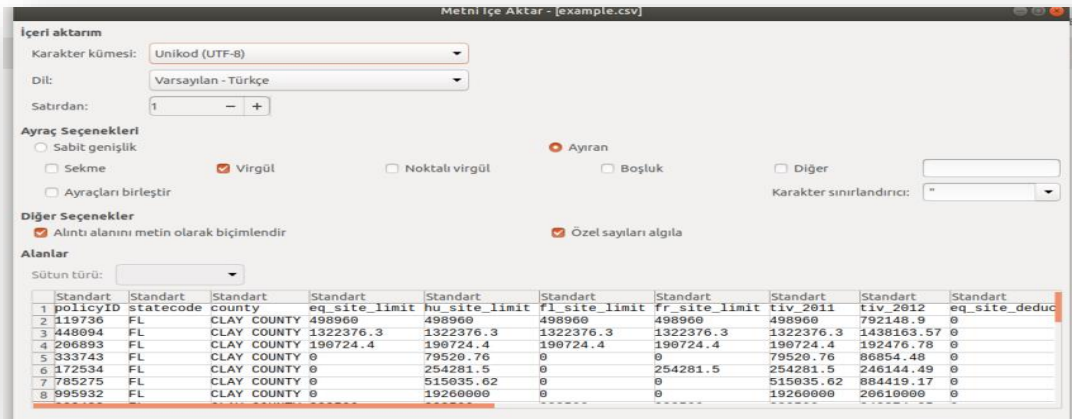
1. Open: Dosya seçimini onaylar.
2. Max/Min: İncelenecek verilerin bayt sınırlarını tespit eder.

3. Map+Reduce: Basılı durumda ise haritalama-indirgeme basılı değil ise sadece haritalama işlemi yapar.
4. Cols: Cvs ve benzeri dosya uzantılarında open'a basıldığında dosya içinde ilgili kolonu seçmek ve buna uygun bir indeksleme oluşturmak için kullanılır. Bu seçimde amaç verilerin homojenliğini artırır. İstenilen kolon girilir veya pasif halde tutulmak için -1 değeri atanır
5. İncelenecek Yüzde: Girilen veri setinde verilen oran kadar en az ve en yoğun olan veri miktarını almak için belirtilir.



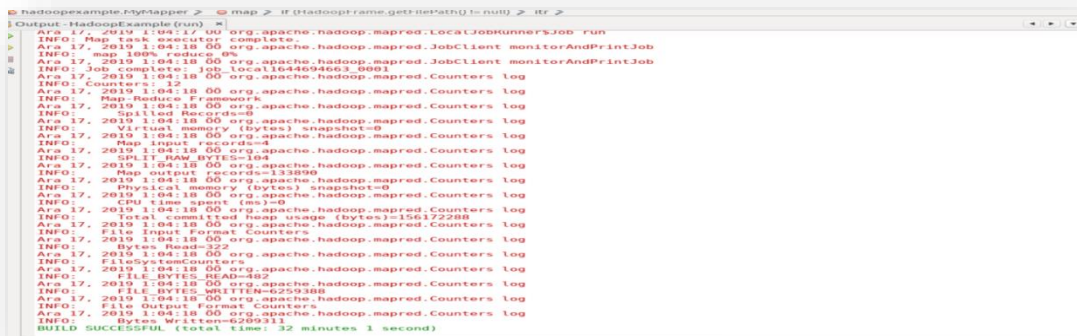
Şekil 5.1. Uygulama Ekranı.

CSV (Virgülle Ayrılan Değerlerden Oluşur) dosyadır. Bu dosya verilerini sütunlarda yerleşik olarak sunmak yerine virgülle ayrılmış olarak depolar. İçindeki Metin ve sayılar bir CSV dosyası olarak kaydedildiğinde kolaylıkla parçalanarak analiz edilebilir.



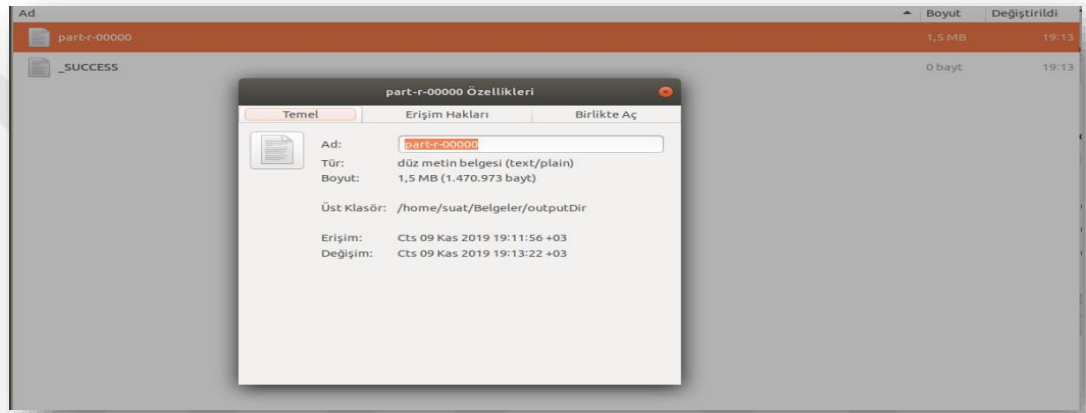
Şekil 5.2. Örnek kullanılan cvs dosya.

Örnek çalışmasında haritalandırılan değerler virgül ile ayrılmıştır veri değerlerinin toplamı kısaca hücre adedi 659.430 olmasına rağmen sınır değerimiz alınmamış harita boyutu kısa tutulmamıştır. Ayrıca indirgeme aşamasında çok ve doğru sonuç hızlıca elde edilebilmesi amaçlandığı için sıralama azdan çoğa doğru olmuştur (Şekil 5.3). Sıralama sonucunda oluşan sistemsel zaman farkı içine sıralama algoritması dahildir. Pivot değeri ortanca değerden alınarak seçilmiştir, ortanca değeri analiz etmeye başlayarak sırama yapılacaktır. Burada kullanılan yapıda pivot seçimi ve kullanılan donanımsal özellikler süre durumu değiştirebilir. Amaç yazılım olarak iyileştirme sağlamak olduğu için bu konu göz ardı edilmiştir.



Şekil 5.3. Çalışma ekranı.

Uygulama çalıştıktan sonra part ismi ile başlayan bir sonuç dosyası oluşturur (Şekil 5.4). Çalıştırılan haritalama ise part-m veya haritalama ardından indirme çalıştırılmış ise part-r olarak başlayan çıktı sonuç dosyası alınmaktadır. Sonuçdosyası sistemde belirtilen çıktı yolu üzerinde oluşturulur. Bu uygulama, verilen çıktı yolunda outputDir klasörü tanımlanmıştır. Sonuç dosyaları burada otomatik olarak oluşmaktadır. Bu sonuç dosyaları bize üzerinde çalışacağım verileri düzenli bir şekilde vermektedir.



Şekil 5.4. Sonuç dosyası.

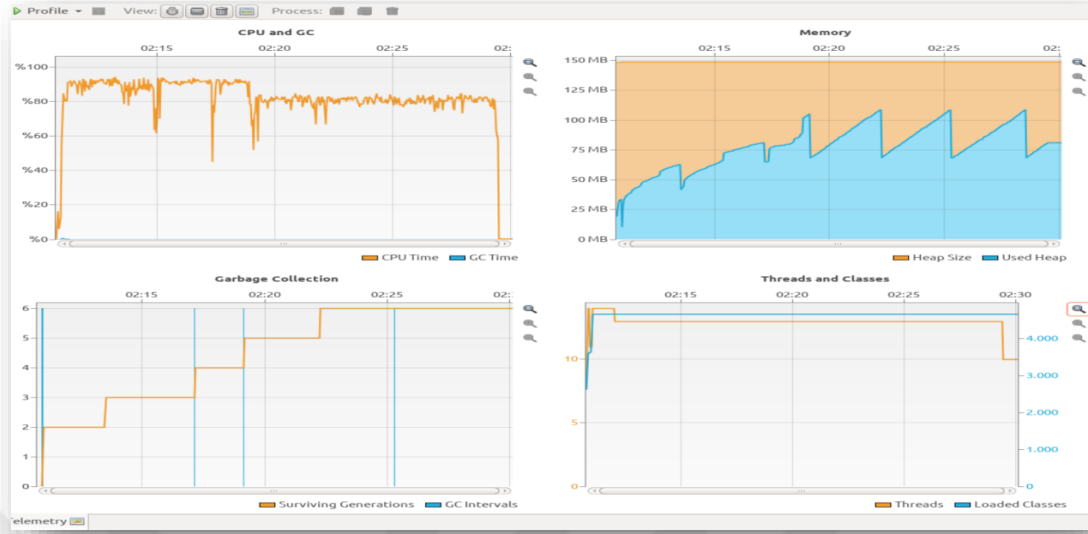
5.3. Birinci Deneme

1. Veri yığını veri sınırlama yapılmamıştır.
2. Verilerde ilk işleme girdiği indeks değerleri tutulmuş. Haritalamada gösterilmiştir
3. Veri Boyutu bayt bazında sınırlanmamıştır.
4. Veri sınırlama ve indeksleme olmadığı durumun sonucunu göstermeyi amaçlamıştır.

140	[140,0,0.0016681073047935338,12,6,]
141	[141,0,0.0024263178978815037,16,2,]
142	[142,1,5164611861759398E-4,1037,]
143	[143,0,0.00803724428673248,1350,]
144	[144,7,582385938879699E-4,6087,]
145	[145,0,0.004094445202675038,4950,]
146	[146,0,0.002123845660646316,3510,]
147	[147,1,5164611861759398E-4,70,2,]
148	[148,0,0.000186800405150075,6750,]
149	[149,9,0.09876711705564E-4,6730,]
150	[150,0,0.0027260301351160915,7290,]
151	[151,0,0.002577984016499098,9270,]
152	[152,0,0.0013648150675583458,1020,]
153	[153,0,0.0056109063888850977,5508,]
154	[154,0,0.00118450480909094754,3681,]
155	[155,0,0.004701029677145413,2205,]
156	[156,9,0.09876711705564E-4,44,2,]
157	[157,0,0.006975721456409323,2754,]
158	[158,0,0.00338761415161579,1830,]
159	[159,0,0.0125866278452603,2880,]
160	[160,0,0.002746017792039096,57,0,]
161	[161,0,0.005844744703759E-4,6075,]
162	[162,1,5164611861759398E-4,6099,]
163	[163,0,0.004549383558527819,1377,]
164	[164,6,0.005844744703759E-4,13,5,]
165	[165,0,0.003942799084057444,1935,]
166	[166,7,582385938879699E-4,80,7,]
167	[167,0,0.0015164611861759399,3735,]
168	[168,1,0.03292372318797E-4,74,7,]
169	[169,0,0.0033621460995870676,1107,]
170	[170,0,0.0016681073047935338,5079,]
171	[171,0,0.004701029677145413,6550,]
172	[172,9,0.09876711705564E-4,8532,]
173	[173,0,0.0015164611861759399,2952,]
174	[174,0,0.0016681073047935338,3720,]
175	[175,0,0.0037611329654308495,733,]
176	[176,7,582385938879699E-4,2943,]
177	[177,0,0.0037611329654308495,1197,]
178	[178,0,0.0022746017792639096,7608,]
179	[179,7,582385938879699E-4,5601,]
180	[180,0,0.000658447447037595,6714,]
181	[181,0,0.0012131609480407519,8200,]
182	[182,0,0.0015164611861759399,2259,]
183	[183,0,0.0015164611861759399,2484,]
184	[184,0,0.00538761415161579,3177,]
185	[185,0,0.0033621460995870676,3485,]
186	[186,0,0.0015164611861759399,8406,]
187	[187,7,582385938879699E-4,8181,]
188	[188,9,0.09876711705564E-4,3222,]
189	[189,7,582385938879699E-4,1566,]
190	[190,0,0.003942799084057444,6210,]
191	[191,0,0.0015164611861759399,1050,]
192	[192,0,0.003032923723518798,4599,]
193	[193,1,5164611861759398E-4,7620,]
194	[194,7,582385938879699E-4,2889,]
195	[195,0,0.0015164611861759399,6705,]
196	[196,9,0.09876711705564E-4,2016,]
197	[197,0,0.0015164611861759399,1701,]
198	[198,0,0.0037911529654398495,3384,]

Şekil 5.5. Sadece haritalama birinci deneme.

Haritalama işlemi çıktısında sıra numarası, verinin yığın içinde olan tekrar yüzdesi ile puanlanmış ve işlenen verinin bilgisi köşeli parantezler içinde gösterilmiştir. Şekil 5.5. haritala işlemi sonrası elde edilmiştir. Bu işlemde 659.430 veri ele alınmış ve haritalama aşaması sonrası 113.389 olarak azaltılmış bu veriler tekrar eden değerlerdir. İlk noktaları ve toplam adetlerine göre frekans değerleri hesaplanmıştır. İşlem sonrası uzunluğa bağlı elimizde veri yığını oluştu. Haritalama işlemi 11 dakika 25 saniye sürmüş ve bir harita dosyası oluşturulmuştur. 6.161.167 bayt büyüklüğünde ve part-m-00000 isimli dosya oluşmuştur. Haritalama işlemi sonunda sistem nano time olarak 6577001052929 olarak kayda alınmıştır (Şekil 5.6). İndeks listesi ve listedeki her elemana ait frekans puanı ve haritalamada analizi edilmiştir.



Şekil 5.6. Haritalama çalışma grafikleri birinci deneme.

1. (GC) uygulamanın hafıza kullanımını ile (CPU) İşlemci kullanımı gösterir
2. (Heap Size) uygulamanın hafıza konuştığı alan ile UsedHeap kullandığı alanı gösterir
3. GarbageCollector üzerinden çöplerin toplanması işlemi grafiğidir. Hafızanın geri alınması görevini GarbageCollector yapar. (GC) aralığı iki çöp toplama arasını işlemleri ve (SurvivingGeneration) en az bir çöp koleksiyonundan kurtulanları objeleri gösterir.
4. (Threads) İş parçacıkları ile Yüklenmiş (LoadedClasses) Yüklenmiş sınıfları gösterir

İndirgeme işlemi yapılması için kolon seçimi yapılmamıştır. Bundan dolayı tekrarda öncelikli ve ortalama tekrar sayılarını verenler başta olmak üzere tüm veriler puanlanmış ve sıralı bir sonuç çıktısı vermiştir. Sonuç çıktısı Şekil 5.7.'de gösterildiği gibi indirgeme işleminde sıralı olarak 113.389 veriden oluşan 4 ayrı listenin birleşim sonucunu vermiştir. Sıralama ortalama uzunluk ve ortalama frekans göz önünde tutularak elde edilmiştir. 6.707.094 bayt büyüklüğünde ve part-r-00000 isimli dosya oluşmuştur (Şekil 5.8). Liste ortalama uzunluk ve ortalama frekans şartını

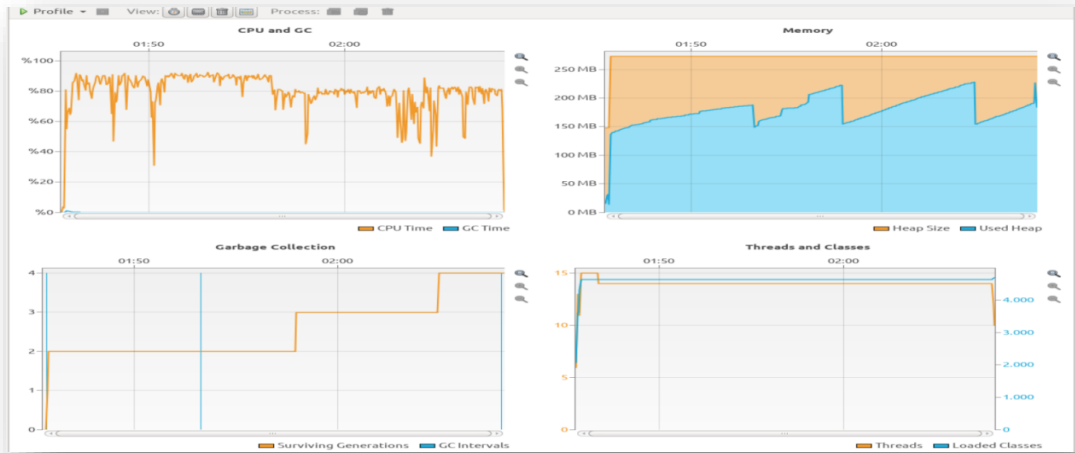
sağlayanlardan başlar ve sıra ortalama uzunluk sağlayan fakat ortalama frekans şartını sağlamayanlar, ortalama uzunluk sağlamayanlar fakat ortalama frekans şartını sağlayanlar son olarak da her iki şartında sağlamadığı değerlerin listesi olarak sonuçlanır.



ID	length	Value
64954	length	Value = 263618.76
64955	length	Value = -81.36739
64956	length	Value = 27.298871
64957	length	Value = 141996.88
64958	length	Value = 287975.07
64959	length	Value = 27.363315
64960	length	Value = -81.47523
64961	length	Value = 398315.93
64962	length	Value = 478436.88
64963	length	Value = 27.287487
64964	length	Value = -81.36747
64965	length	Value = 27.323424
64966	length	Value = -81.32348
64967	length	Value = 27.273764
64968	length	Value = 27.258395
64969	length	Value = 627116.14
64970	length	Value = 27.348835
64971	length	Value = 287187.33
64972	length	Value = 242717.82
64973	length	Value = -81.39859
64974	length	Value = 27.275054
64975	length	Value = 252669.47
64976	length	Value = 264554.43
64977	length	Value = 5844574.7
64978	length	Value = 6652858.7
64979	length	Value = 27.267025
64980	length	Value = -81.37363
64981	length	Value = -81.31592
64982	length	Value = 162688.64
64983	length	Value = -81.37268
64984	length	Value = 873878.54
64985	length	Value = 188838.46
64986	length	Value = -81.41664
64987	length	Value = 161347.77
64988	length	Value = -81.58035
64989	length	Value = 169410.25
64990	length	Value = -81.57743
64991	length	Value = 27.963429
64992	length	Value = 27.893525
64993	length	Value = 163439.76
64994	length	Value = 3361972.5
64995	length	Value = -81.58784
64996	length	Value = 879321.77
64997	length	Value = 858884.82
64998	length	Value = 27.894022
64999	length	Value = 858882.72
65000	length	Value = 163938.76
65001	length	Value = -81.59358
65002	length	Value = 317793.25
65003	length	Value = 582927.25
65004	length	Value = -81.57928
65005	length	Value = 138089.38
65006	length	Value = 115291.59
65007	length	Value = 6577905.5
65008	length	Value = 412515.12
65009	length	Value = 568837.41
65010	length	Value = 996617.96
65011	length	Value = 3248181.7
65012	length	Value = 949432.17

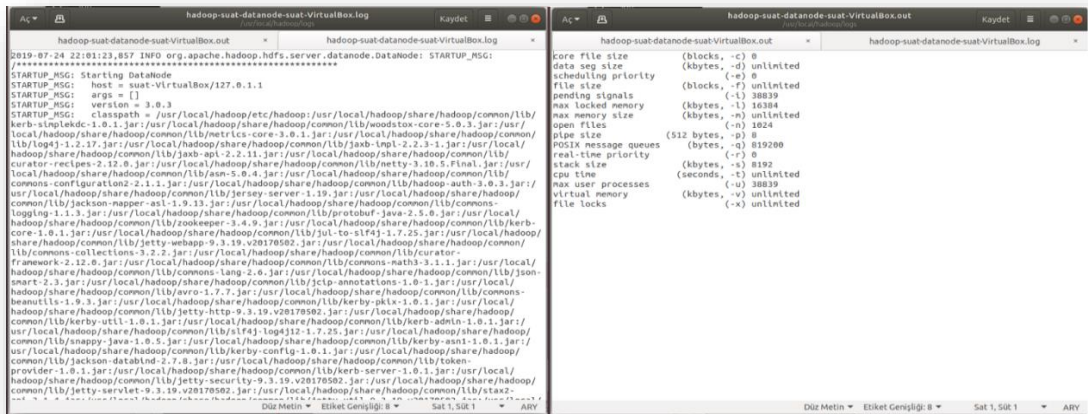
Şekil 5.7. Sonuç haritalama indirgeme birinci deneme.

Tutuğumuz veri miktarı azalmasına karşın dosya boyutu artmıştır. Bunun sebepleri işlenen değerler dışında işlemde oluşan değerlerin çıktı olarak yazdırılmasıdır. Ayrıca Hadoop üzerinde çıkan ve yazdırılan sonuç dosyasının dosya boyutunda artış sağlamıştır. Buradan ortaya çıkan sonuç. Maliyetleri azaltmak için sadece veri adedini düşürmemiz yetmez. Her veriler için uygun sınırlama ve indeksleme kullanılmalıdır. İkinci yapacağımız denemede bu noktalarında geliştirdiğim yazılımda sonucunu olumlu olduğunu gösterilmiştir.



Şekil 5.8. Haritalama indirgeme çalışma grafikleri birinci deneme.

Hadoop yazılımında işlenen verilerin ve oluşan işlemlerin takip edildiği log çıktıları oluşturulmuştur. İşlemler sonrasında Hadoop işlem dosyalarında olan log dosyasını inceleyebiliriz. Bu log dosyasına ait çıktı Şekil 5.9. gösterilmiştir. Bu noktadan anlaşıldığı gibi işlemler devam ederken Hadoop'da arka taraftan log işlemine devam edilmiştir. İşlemin toplam zamanı 23 dakika 35 saniye olarak log düşürülmüştür. Okunan dosya 4.911.589 bayt yazılan dosya ise 1.482.473 bayt olarak kayda geçmiştir. Okunan dosya cvs formatındadır.



Şekil 5.9. Hadoop log.

5.4. İkinci Deneme

İkinci deneme için gerçekleştirilen işlem basamakları aşağıda verilmiştir.

1. Veri yığnında bir kolon seçilmiş böylece veri sınırlama yapılmıştır.
2. Verilerde tüm işleme girdiği indekslerin değerleri tutulmuş.
3. Veri Boyutu bayt bazında sınırlanmamıştır.
4. Veri sınırlama ve indeksleme olduğu durumun sonucunu göstermeyi amaçlamıştır.

Sıralama işlemi sonrası indirgeme işlemi yapılırsa sonuç çıktısında veriler sıralanmıştır. Burada köşeli parantez içinde verinin sahip olduğu değer puanı ikinci sırada sonrasında verinin kendisi ardından bulunduğu indeks adresleri sıralı bir şekilde verilmiştir. Yapılan deneme sonucu Şekil 5.10'da bize gösterilen sonuçta çok büyük kazanım elde edilmiştir. İşlemde 659.430 veri ele alınmış ve haritalama aşaması sonrasında sonuç 68 adet veriden oluşmuş ayrıca bu veriler tekrar eden değerlerdir. Toplam adetlerine göre frekans değerleri hesaplanmıştır. Bunun dışında ortalama verilerin bayt uzunluğu hesaplanmamıştır. İşlem sonrası uzunluğa bağlı elimizde veri yığını oluştu. Haritalama işlemi 28 saniye sürdü ve bir harita dosyası oluşturdu. 211.54 bayt büyüklüğünde ve part-m-00000 isimli dosya oluşmuştur. Haritalama işlemi sonunda sistem nano time olarak 46106917359782 olarak kayda alınmıştır. İncelenen değerlerin hepsine ait indeks listesi değerinin yanında kayda alınmıştır.

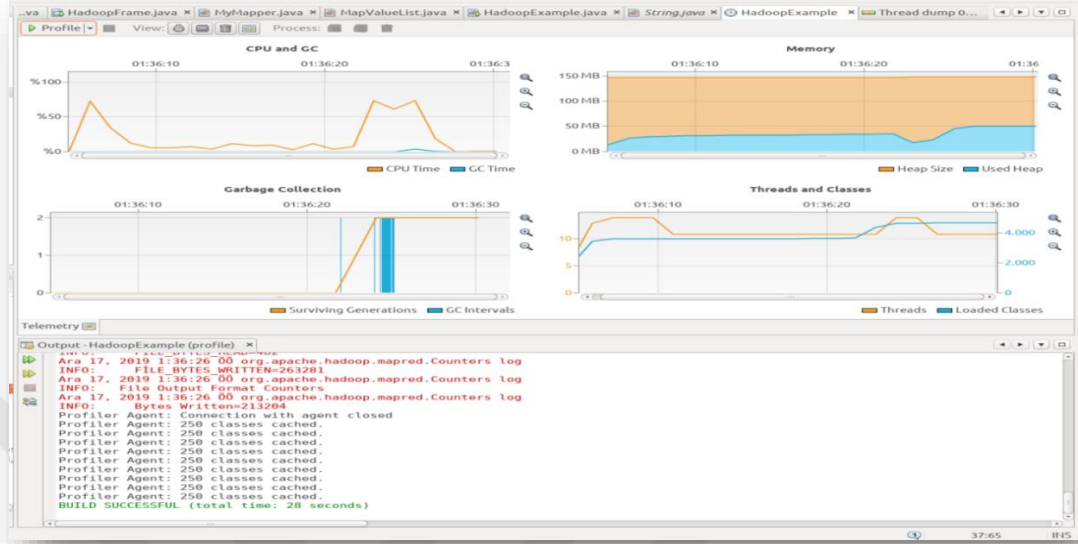
The image shows a screenshot of a text editor window titled "Metin Düzenleyici". The editor contains a large block of text, which appears to be a dataset or a list of coordinates. The text is organized into rows and columns, with some rows starting with labels like "Ac =", "part-m-00000", and "KAYDET". The editor window has a title bar with "Metin Düzenleyici" and a status bar with "Satır 01336" and "part-m-00000".

Şekil 5.10. Sadece haritalama ikinci deneme.

İki denemde sadece haritama çizilme grafikleri sonuçlarının (Şekil 5.11 ve Şekil 5.6) görselleri bize çözüm konusunda yardımcı olmaktadır. İkinci denemede veri sınırlama ve indeksleme kullanımı yoğun tam anlamı ile sağladığı için istatistikler olumlu yöne geçmiştir. Kullanılan geçici hafıza büyük oranda düşmüştür. Uygulanılan hafıza ve işlemci kullanımı çok azalmıştır ve ayrıca oluşan artık, kullanılmayan çöpl verilerinde azalığın görebiliriz. İki deneme için oluşan bu grafik farkları haritalama sonrası indirgeme yapılan işlemlerin grafiklerinde de oldukça kolay bir şekilde aynı olumlu yönde gözlenebilir.

Haritalama işlemi sonrasında eğer indirgeme yapılmaz ise oluşan sonuç listesinin bir bölümü Şekil 5.12’de gösterilmiştir. Burada sonuç olarak indirgemeişleminde sıralı olarak 68 veriden oluşan 4 ayrı listenin birleşim sonucunu vermiştir. 2.547 bayt büyüklüğünde ve part-r-00000 isimli dosya oluşmuştur. Sıralama ortalama uzunluk ve ortalama frekans göz önünde tutularak elde edilmiştir. Liste ortalama uzunluk ve ortalama frekans şartını sağlayanlardan başlar ve sıra ortalama uzunluk sağlayan fakat ortalama frekans şartını sağlamayanlar, ortalama uzunluk sağlamayanlar fakat

ortalama frekans şartını sağlayanlar son olarak her iki şartında sağlamadığı değerlerin listesi olarak sonuçlanır.



Şekil 5.11. Haritalama çalışma grafikleri ikinci deneme.

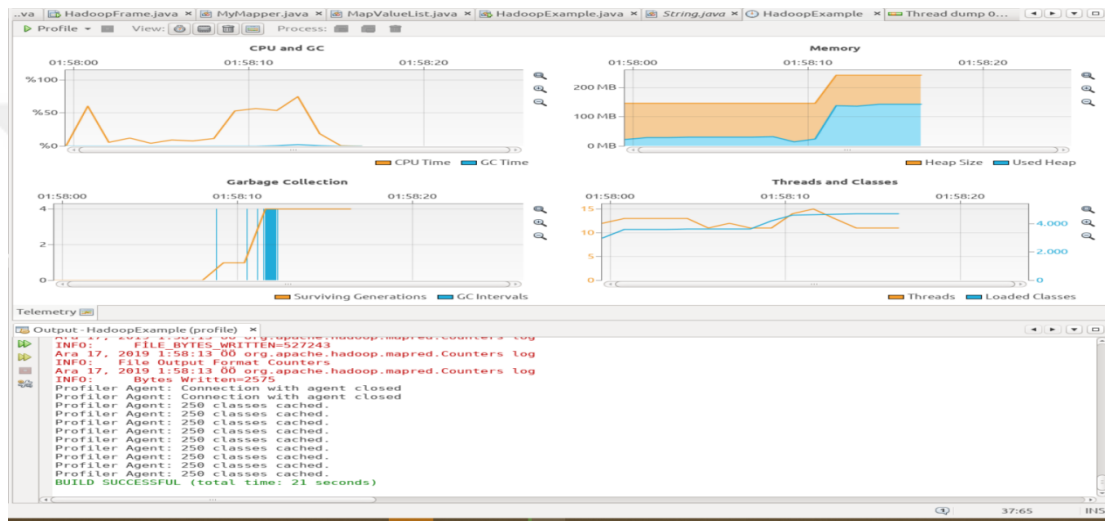
part-m-00000	part-m-00000	part-r-00000	part-r-00000
1 length = 13 value = KOLUSACOUNTY			
2 length = 13 value = ALACHUACOUNTY			
3 length = 13 value = BREVARDCOUNTY			
4 length = 13 value = BROWARDCOUNTY			
5 length = 13 value = MANATEECOUNTY			
6 length = 13 value = COLLIERCOUNTY			
7 length = 14 value = ESCAMBIACOUNTY			
8 length = 14 value = OKALOOSACOUNTY			
9 length = 14 value = SEMINOLECOUNTY			
10 length = 14 value = PINELLASCOUNTY			
11 length = 14 value = SARASOTACOUNTY			
12 length = 13 value = STJOHNSCOUNTY			
13 length = 15 value = HIGHLANDSCOUNTY			
14 length = 15 value = CHARLOTTECOUNTY			
15 length = 15 value = SANTIAGOACOUNTY			
16 length = 15 value = FLORIDACOUNTY			
17 length = 15 value = PALMBEACHCOUNTY			
18 length = 17 value = INDIANRIVERCOUNTY			
19 length = 16 value = HILLSBOROUGHCOUNTY			
20 length = 14 value = SUGARCREECOUNTY			
21 length = 14 value = COLUMBIACOUNTY			
22 length = 14 value = BRADFORDCOUNTY			
23 length = 14 value = HAMILTONCOUNTY			
24 length = 14 value = FRANKLINCOUNTY			
25 length = 14 value = HERNANDOCOUNTY			
26 length = 15 value = LAFAYETTESCOUNTY			
27 length = 15 value = JEFFERSONCOUNTY			
28 length = 15 value = GILCHRISTCOUNTY			
29 length = 14 value = NorthForMyers			
30 length = 16 value = WASHINGTONCOUNTY			
31 length = 9 value = BAYCOUNTY			
32 length = 10 value = LEECOUNTY			
33 length = 10 value = POLKCOUNTY			
34 length = 11 value = DUVALCOUNTY			
35 length = 11 value = PASCOCOUNTY			
36 length = 12 value = HARRISCOUNTY			
37 length = 12 value = ORANGECOUNTY			
38 length = 12 value = CITRUSCOUNTY			
39 length = 6 value = county			
40 length = 7 value = Orlando			
41 length = 10 value = CLAYCOUNTY			
42 length = 10 value = LAKECOUNTY			
43 length = 10 value = LEONCOUNTY			
44 length = 10 value = GULFCOUNTY			
45 length = 10 value = LEVYCOUNTY			
46 length = 11 value = BAKERCOUNTY			
47 length = 11 value = UNIONCOUNTY			
48 length = 11 value = DIXIECOUNTY			
49 length = 12 value = NASSAUCOUNTY			
50 length = 12 value = PUTNAMCOUNTY			
51 length = 12 value = SUFTCOUNTY			
52 length = 12 value = TAYLORCOUNTY			
53 length = 12 value = MALTOUNCOUNTY			
54 length = 12 value = HOLMESCOUNTY			
55 length = 12 value = PONDERCOUNTY			
56 length = 12 value = MARTINCOUNTY			
57 length = 12 value = HENRYCOUNTY			
58 length = 12 value = GLADESCOUNTY			
59 length = 12 value = HARDEECOUNTY			
60 length = 12 value = DECATACOUNTY			

Şekil 5.12. Sonuç Haritalama indirgeme ikinci deneme.

İki deneme arasında oluşan farkın sebebi verilerin homojen olmasıdır. Analizde frekans bazında puanlamayı çok daha verimli hale getirmiştir. İkinci denemede sadece

bir kolondaki veriler puanlamıştır. Bu Veri değerlendirmede sınırlar belirtildikçe değerlendirme verimi artıracağıının kanıtıdır (Şekil 5.13 ve Şekil 5.14).

2012'de Florida'da oluşan bu agresif büyüme analiz etmemek istendiği takdirde yapılan bu denemeler bazında VolusiaCounty analizinde en az maliyet ile en değerli verilerin olduğu yerdir. Bu veri yığını üzerinde yapılacak en az maliyetli ve en değerli veri çalışmasında VolusiaCounty'e ait veriler öncelikle kullanılmalıdır.

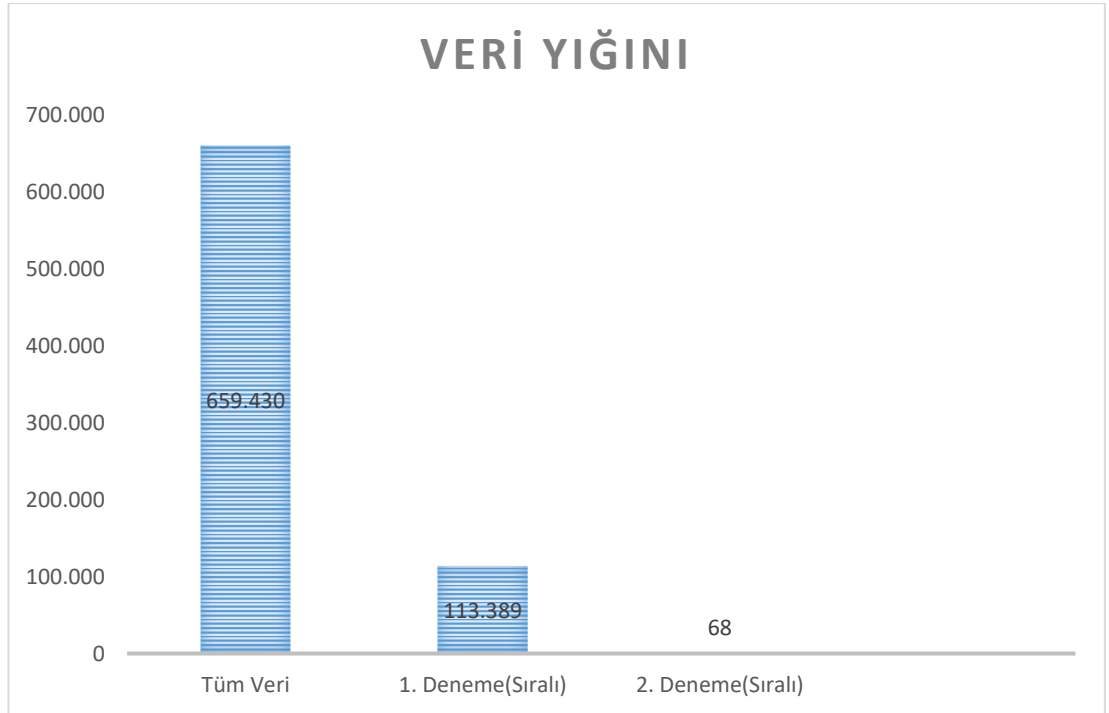


Şekil 5.13. Haritalama indirgeme çalışma grafikleri ikinci deneme.

12037	12038	12039	12040	12041	12042	12043	12044	12045	12046	12047	12048	12049	12050	12051	12052	12053	12054	12055	12056	12057	12058	12059	12060	12061	12062	12063	12064	12065	12066	12067	12068	12069	12070	12071	12072	12073	12074	12075	12076	12077	12078	12079	12080	12081	12082	12083	12084	12085	12086	12087	12088	12089	12090	12091	12092	12093	12094	12095	12096	12097	12098	12099	12100	12101	12102	12103	12104	12105	12106	12107	12108	12109	12110	12111	12112	12113	12114	12115	12116	12117	12118	12119	12120	12121	12122	12123	12124	12125	12126	12127	12128	12129	12130	12131	12132	12133	12134	12135	12136	12137	12138	12139	12140	12141	12142	12143	12144	12145	12146	12147	12148	12149	12150	12151	12152	12153	12154	12155	12156	12157	12158	12159	12160	12161	12162	12163	12164	12165	12166	12167	12168	12169	12170	12171	12172	12173	12174	12175	12176	12177	12178	12179	12180	12181	12182	12183	12184	12185	12186	12187	12188	12189	12190	12191	12192	12193	12194	12195	12196	12197	12198	12199	12200	12201	12202	12203	12204	12205	12206	12207	12208	12209	12210	12211	12212	12213	12214	12215	12216	12217	12218	12219	12220	12221	12222	12223	12224	12225	12226	12227	12228	12229	12230	12231	12232	12233	12234	12235	12236	12237	12238	12239	12240	12241	12242	12243	12244	12245	12246	12247	12248	12249	12250	12251	12252	12253	12254	12255	12256	12257	12258	12259	12260	12261	12262	12263	12264	12265	12266	12267	12268	12269	12270	12271	12272	12273	12274	12275	12276	12277	12278	12279	12280	12281	12282	12283	12284	12285	12286	12287	12288	12289	12290	12291	12292	12293	12294	12295	12296	12297	12298	12299	12300	12301	12302	12303	12304	12305	12306	12307	12308	12309	12310	12311	12312	12313	12314	12315	12316	12317	12318	12319	12320	12321	12322	12323	12324	12325	12326	12327	12328	12329	12330	12331	12332	12333	12334	12335	12336	12337	12338	12339	12340	12341	12342	12343	12344	12345	12346	12347	12348	12349	12350	12351	12352	12353	12354	12355	12356	12357	12358	12359	12360	12361	12362	12363	12364	12365	12366	12367	12368	12369	12370	12371	12372	12373	12374	12375	12376	12377	12378	12379	12380	12381	12382	12383	12384	12385	12386	12387	12388	12389	12390	12391	12392	12393	12394	12395	12396	12397	12398	12399	12400	12401	12402	12403	12404	12405	12406	12407	12408	12409	12410	12411	12412	12413	12414	12415	12416	12417	12418	12419	12420	12421	12422	12423	12424	12425	12426	12427	12428	12429	12430	12431	12432	12433	12434	12435	12436	12437	12438	12439	12440	12441	12442	12443	12444	12445	12446	12447	12448	12449	12450	12451	12452	12453	12454	12455	12456	12457	12458	12459	12460	12461	12462	12463	12464	12465	12466	12467	12468	12469	12470	12471	12472	12473	12474	12475	12476	12477	12478	12479	12480	12481	12482	12483	12484	12485	12486	12487	12488	12489	12490	12491	12492	12493	12494	12495	12496	12497	12498	12499	12500	12501	12502	12503	12504	12505	12506	12507	12508	12509	12510	12511	12512	12513	12514	12515	12516	12517	12518	12519	12520	12521	12522	12523	12524	12525	12526	12527	12528	12529	12530	12531	12532	12533	12534	12535	12536	12537	12538	12539	12540	12541	12542	12543	12544	12545	12546	12547	12548	12549	12550	12551	12552	12553	12554	12555	12556	12557	12558	12559	12560	12561	12562	12563	12564	12565	12566	12567	12568	12569	12570	12571	12572	12573	12574	12575	12576	12577	12578	12579	12580	12581	12582	12583	12584	12585	12586	12587	12588	12589	12590	12591	12592	12593	12594	12595	12596	12597	12598	12599	12600	12601	12602	12603	12604	12605	12606	12607	12608	12609	12610	12611	12612	12613	12614	12615	12616	12617	12618	12619	12620	12621	12622	12623	12624	12625	12626	12627	12628	12629	12630	12631	12632	12633	12634	12635	12636	12637	12638	12639	12640	12641	12642	12643	12644	12645	12646	12647	12648	12649	12650	12651	12652	12653	12654	12655	12656	12657	12658	12659	12660	12661	12662	12663	12664	12665	12666	12667	12668	12669	12670	12671	12672	12673	12674	12675	12676	12677	12678	12679	12680	12681	12682	12683	12684	12685	12686	12687	12688	12689	12690	12691	12692	12693	12694	12695	12696	12697	12698	12699	12700	12701	12702	12703	12704	12705	12706	12707	12708	12709	12710	12711	12712	12713	12714	12715	12716	12717	12718	12719	12720	12721	12722	12723	12724	12725	12726	12727	12728	12729	12730	12731	12732	12733	12734	12735	12736	12737	12738	12739	12740	12741	12742	12743	12744	12745	12746	12747	12748	12749	12750	12751	12752	12753	12754	12755	12756	12757	12758	12759	12760	12761	12762	12763	12764	12765	12766	12767	12768	12769	12770	12771	12772	12773	12774	12775	12776	12777	12778	12779	12780	12781	12782	12783	12784	12785	12786	12787	12788	12789	12790	12791	12792	12793	12794	12795	12796	12797	12798	12799	12800	12801	12802	12803	12804	12805	12806	12807	12808	12809	12810	12811	12812	12813	12814	12815	12816	12817	12818	12819	12820	12821	12822	12823	12824	12825	12826	12827	12828	12829	12830	12831	12832	12833	12834	12835	12836	12837	12838	12839	12840	12841	12842	12843	12844	12845	12846	12847	12848	12849	12850	12851	12852	12853	12854	12855	12856	12857	12858	12859	12860	12861	12862	12863	12864	12865	12866	12867	12868	12869	12870	12871	12872	12873	12874	12875	12876	12877	12878	12879	12880	12881	12882	12883	12884	12885	12886	12887	12888	12889	12890	12891	12892	12893	12894	12895	12896	12897	12898	12899	12900	12901	12902	12903	12904	12905	12906	12907	12908	12909	12910	12911	12912	12913	12914	12915	12916	12917	12918	12919	12920	12921	12922	12923	12924	12925	12926	12927	12928	12929	12930	12931	12932	12933	12934	12935	12936	12937	12938	12939	12940	12941	12942	12943	12944	12945	12946	12947	12948	12949	12950	12951	12952	12953	12954	12955	12956	12957	12958	12959	12960	12961	12962	12963	12964	12965	12966	12967	12968	12969	12970	12971	12972	12973	12974	12975	12976	12977	12978	12979	12980	12981	12982	12983	12984	12985	12986	12987	12988	12989	12990	12991	12992	12993	12994	12995	12996	12997	12998	12999	13000	13001	13002	13003	13004	13005	13006	13007	13008	13009	13010	13011	13012	13013	13014	13015	13016	13017	13018	13019	13020	13021	13022	13023	13024	13025	13026	13027	13028	13029	13030	13031	13032	13033	13034	13035	13036	13037	13038	13039	13040	13041	13042	13043	13044	13045	13046	13047	13048	13049	13050	13051	13052	13053	13054	13055	13056	13057	13058	13059	13060	13061	13062	13063	13064	13065	13066	13067	13068	13069	13070	13071	13072	13073	13074	13075	13076	13077	13078	13079	13080	13081	13082	13083	13084	13085	13086	13087	13088	13089	13090	13091	13092	13093	13094	13095	13096	13097	13098	13099	13100	13101	13102	13103	13104	13105	13106	13107	13108	13109	13110	13111	13112	13113	13114	13115	13116	13117	13118	13119	13120	13121	13122	13123	13124	13125	13126	13127	13128	13129	13130	13131	13132	13133	13134	13135	13136	13137	13138	13139	13140	13141	13142	13143	13144	13145	13146	13147	13148	13149	13150	13151	13152	13153	13154	13155	13156	13157	13158	13159	13160	13161	13162	13163	13164	13165	13166	13167	13168	13169	13170	13171	13172	13173	13174	13175	13176	13177	13178	13179	13180	13181	13182	13183	13184	13185	13186	13187	13188	13189	13190	13191	13192	13193	13194	13195	13196	13197	13198	13199	13200	13201	13202	13203	13204	13205	13206	13207	13208	13209	13210	13211	13212	13213	13214	13215	13216	13217	13218	13219	13220	13221	13222	13223	13224	13225	13226	13227	13228	13229	13230	13231	13232	13233	13234	13235	13236	13237	13238	13239	13240	13241	13242	13243	13244	13245	13246	13247	13248	13249	13250	13251	13252	13253	13254	13255	13256	13257	13258	13259	13260	13261	13262	13263	13264	13265	13266	13267	13268	13269	13270	13271	1
-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	-------	---

BÖLÜM 6. SONUÇ

Sürekli artan veri yoğunluğu verilerin analizini çok önemli hale getirmiştir. Giderek artan maliyet ve zaman kısıtlılığı, haritalama-indirgeme aşamalarında oluşan maliyetlere çözümler bulmayı gerektirmiştir. Burada analiz yapılırken kullanılacak işlemlerin daha az maliyetli olması amaçlanmıştır. Bu amaçtan dolayı homojen veri yığınları filtrelenmesi ve az maliyetli verilere öncelik verilmiştir. Bu geliştirme ve üzerinde yapılan çalışmada analiz işlemi esnasında, maliyetin minimize edilmesi hedeflenmiştir. Uygulamada gerçekleştirilecek işlemlerin daha az maliyetli olması için homojen veri yığınlarının filtrelenmesi amaçlanmıştır. Şekil 6.1 fitreme sonuçlarının farkları görsel olarak sunulmuştur.



Şekil 6.1. Veri azaltım grafiği görseli.

Geliştirilen algoritma ve program uygulaması sonucunda ortaya çıkan verileri birleşiminde büyük veri analizinde haritalamaindirgeme işlevsel bir uygulamadır analizine varırız. Haritalama aşaması indirgeme aşamasını etkilemektedir. Bu sebepten dolayı haritalama aşaması sırasında indirgeme planlanmalıdır. Haritalama aşamasına başlamada sınır değerler var ise sınıra uymayan değerler dışlanmalıdır ve buradan alınacak örneklem sonuçlanma için yeterli ise örneklem sınırları belirlenmelidir. Haritalama aşaması sonrası hızlı sıralama ile verilerin içinden az kaynağa ve daha çok veriye ulaşmamızı sağlar. Bu sınırlara verilerin en uygun kapladığı alan miktarının sonucu bize getirir. Hızlı sıralama hem az alana ihtiyaç duyduğu hemde sıralamalarda veriyi n adet parçaya bölerek daha az zaman tüketimi, bellek tüketimi, işlemci tüketimi ve disk tüketimi gibi konularda maliyet azaltıcı etki gösterecektir. Büyük veride yaklaşım **sınırlama->haritalama->sıralama->değerlendirme->indirgeme** olarak seçilmeli ve geliştirilmelidir. Her çalışmada sıralama ve sınırlama gerekmebilir, bu gibi durumlarda çalışmaların verimliği için metotların uygun olarak dönüştürülmesi (Ezilmesi) kullanıcı verimini arttıracaktır. Verilerin tekrar etme sıklıkları önem düzeyini çıkarmada bir sınır önemi sağlayabilir ve bu sebepten dolayı verilerin işlenmesinde, tekrar sınırlaması dikkate alınarak düzenlenmelidir.

Haritalama öncesi aynı değerler ile analiz içinden çıkarılmış maliyet kazanımı elde edilmiş ve aynı zamanda homojenlik arttırılmıştır. Bu değerlerin indekslenerek maliyet azaltılmıştır. Tüm veri yığnında belli bir alan seçimine girilmesi veya girilmemesi durumuna bakılmaksızın az maliyetli ve tekrar frekansı yüksek olan veriler öncelikli hale gelmiştir. Böylece işlemler için indeksleme ile maliyeti oluşmuş fakat verilerin işleme esnasında oluşacak maliyet minimuma indirilmiştir. Sıralama algoritmalarında hızlı sıralama (böl yönet) algoritmaları verilerin sıralamasında zaman maliyetini azaltmıştır ve ayrıca hızlı sıralamada üç parçalı yapı kullanılarak daha fazla maliyet azatımı sağlanmıştır. Maliyet kazanımı için ortalama tekrar frekansı ve bayt uzunluğu belirlenmiş bu değerler sıralamada kullanılmıştır. Bu analiz ile veriler yüzdelik puanlama ile değerlendirilmiştir. Değerli veriler listelemede öne alınmış ve böylece haritalama ve indirgeme çalışması ile veri analizinde gerekli kaynak ve

alışma süresi azaltılmıştır. Bu alışma ile uygulamada büyük verinin artan maliyeti düşürülmüştür.

Yapılan bu alışma sonucunda elde edilen tecrübeler dayalı olarak, harita indirgeme öncesinde veriler en sağlıklı puanlama için homojen olmalıdır. Veriler en sağlıklı puanlama içinde homojen olmalıdır. Bu kullanıcının belli oranda kontrolünde bir seçenektir ve bu sebepten homojenliği artıracak yöntemler üzerinde alışılmalıdır. Ayrıca farklı verilerin birbirlerini ait değeri etkileyebileceği durumlar için. Verilerin apraz puanlaması sağlanmalıdır. Bayt uzunluğu alanı içinde değersiz veri varsa azaltılan listeye eklenmeden önce değerlendirme dışına itilmelidir. Bu duruma örnek olarak, boşluk veya virgülden sonraki sayı basamak değeri verilebilir. Son olarak sıralama algoritması alışırken yoğun ram tüketimini azaltıcı donanımsal ve yazılımsal özüm geliştirilmelidir.

KAYNAKLAR

- [1] Dean J., Ghemawat S., Simplified Data Processing On Large Clusters, Communications of the ACM January, 2008.
- [2] Yang H.C., Parker D.S., Simplified Indexing on Large Map-Reduce-Merge Clusters, DASFAA 2009: Database Systems for Advanced Applications pp 308-322, 2009.
- [3] Zaharia M., Konwinski A., Joseph A.D., Katz R.H., Stoica I., Improving Map Reduce Performance in Heterogeneous Environments, Stoica - Osd, 2008 - static.usenix.org
- [4] Yang H., Dasdan A., Hsiao R.L., Parker D.S., Map-Reduce-Merge: Simplified Relational Data Processing on Large Clusters, SIGMOD '07: Proceedings of the 2007 ACM SIGMOD, International Conference on Management of Data June 2007 Pages 1029–104, 2007.
- [5] Kolb L., Thor A., Rahm E., Multi-Pass Sorted Neighborhood Blocking with Map Reduce, Computer Science – Research and Development, Vol. 27, pp. 45–63, 2012.
- [6] Verma A., Cho B., Zea N., Gupta I., Campbell R.H., Breaking the Map Reduce stage barrier, Published: 10 September 2011 Cluster Computing Vol. 16, pp.191–206, 2013.
- [7] Goodrich M.T., Sitchinava N., Zhang Q., Sorting, Searching, and Simulation in the Map Reduce Framework, ISAAC 2011: Algorithms and Computation, pp 374-383, 2011.
- [8] Herodotou H., Babu S., Profiling What-if Analysis, and Cost-based Optimization of Map Reduce Programs, in Proceedings of the VLDB Endowment 4(11):1111-1122 · August 2011, DOI: 10.14778/3402707.3402746
- [9] Holmes A., Hadoop in Practice, ISBN 9781617290237; 536 pages, October 2012.
- [10] Slagter K., Hsu C.H., Chung Y.C., Zhang D., An improved partitioning mechanism for optimizing massive data analysis using Map Reduce, The Journal of Supercomputing, Vol. 66, pages 539–555, 2013.

- [11] Gopalani S., Arora R., Comparing Apache Spark and Map Reduce with Performance Analysis using K-Means, *International Journal of Computer Applications* (0975 – 8887) Vol. 113 – No. 1, March 2015
- [12] McCreddie R., Macdonald C., Ounis I., Map Reduce Indexing Strategies: Studying Scalability and Efficiency, *Information Processing & Management* Vol. 48, Issue 5, Pages 873-888, September 2012.
- [13] Bakratsas, M., Basaras, P., Katsaros, D., Tassioulas, L., Hadoop MapReduce Performance on SSDS for Analyzing Social Networks, *Big Data Research*, Vol. 11, pp.1-10, March 2018.
- [14] Auradkar, P., Prashanth, T., Aralihaili, S., Kumar, S.P., Sitaram, D., Performance Tuning Analysis of Spatial Operations on Spatial Hadoop Cluster with SSD, *Procedia Computer Science*, V. 167, pp. 2253-2266, 2020.
- [15] Lu, L., Feng, Q.Y., An Optimal Solution of Storing and Processing Small Image Files on Hadoop, *Procedia Computer Science*, Vol. 154, pp. 581-587, 2019.
- [16] Babar, M., Arif, F., Jan, M.A., Tan, Z., Khan, F., Urban Data Management System: Towards Big Data Analytics for Internet of Things Based Smart Urban Environment Using Costimized Hadoop, *Future Generation Computer Systems*, Vol. 96, pp. 398-409, July 2019.
- [17] Tallada, P., Carretero, J., Casals, J., Acosta-Silva, C., Tonello, N., COSMOHUB: Interactive Exploration and Distribution of Astronomical Data on Hadoop, *Astronomy and Computing*, Vol. 32, Article 100391, July 2020.
- [18] Ibrahim, S., Phan, T.D., Carpen-Amarie, A., Chhoub, H.E., Antoniu, G., Governing Energy Consumption in Hadoop Through CPU Frequency Scaling: An Analysis, *Future Generation Computer Systems*, Vol. 54, pp. 219-232, Jan. 2016.
- [19] Rasooli, A., Down, D.G., COSHH: A Classification and Optimization Based Scheduler for Heterogeneous Hadoop Systems, *Future Generation Computer Systems*, Vol. 36, pp. 1-15, July 2014.
- [20] Wu, W.T., Lin, W.W., Hsu, Ching-Hsien, H., LiGang, O., Energy-efficient Hadoop for Big Data Analytics and Computing: A Systematic Review and Research Insights, *Future Generation Computer Systems*, Vol. 86, pp. 1351-1367, September 2018.
- [21] Pradeep, D., Sundar, C., QAOC: Novel Query Analysis and Ontology Based Clustering for Data Management in Hadoop, *Future Generation Computer Systems*, Vol. 108, pp. 849-860, July 2020.

- [22] Mazumdar, S., Scionti, A., Chapter One: Fast Execution of RFD Queries Using Apache Hadoop, *Advances in Computers*, Vol. 119, pp. 1-33, 2020
- [23] Usama, M., Liu, M., Chen, M., Hadoop environment: testing real-life schedulers using benchmark programs, *Digital Communications and Networks*, Vol. 3, Issue 4, pp. 260-273, November 2017
- [24] Kusumakumari, V., Sherigar, D., Chandran, R., Patil, N., Frequent Pattern Mining on Stream Data Using Hadoop Tree-GTree, *Procedia Computer Science*, Vol. 115, pp. 226-273, 2017.
- [25] Meena, K., Tayal, D.K., Castillo, O., Jain, A., Handling Data-skewness in Character Based String Similarity Join Using Hadoop, *Applied Computing and Informatics*, in pres, 16 Nov. 2018
- [26] Bende, S., Shedge, R., Dealing with Small Files Problem in Hadoop distributed File System, *Procedia Computer Science*, Vol. 79, pp. 1001-1012, 2016.
- [27] GOM-Hadoop: A Distributed Framework for Efficient Analytics on Ordered Datasets, *Journal of Parallel and Distributed Computing*, Vol. 83, pp. 58-69, Sept. 2015.
- [28] Lin, H.K., Harding, J.A., Chen, C.I., A Hyperconnected Manufacturing Collaboration System Using The Semantic Web and Hadoop Ecosystem System, *Procedia CIRP*, Vol. 52, pp. 18-23, 2016.
- [29] Lee, C.W., Hsieh, K.Y., Hsieh, S.Y., Hsiao, H.C., A Dynamic Data Placement Strategy for Hadoop in Heterogeneous Environments, *Big Data Research*, Vol. 1, p. 14-22, Aug. 2014.
- [30] Zhou, Y., Taneja, S., Dudeja, G., Qin, X., Alghamdi, M.I., Towards thermal-aware Hadoop clusters, *Future Generation Computer Systems*, Vol. 88, pp. 40-54, November 2018.
- [31] Zhou, Y., Taneja, S., Dudeja, G., Qin, X., Zhang, J., Jiang, M., Alghamdi, M.I., Towards Thermal-aware Hadoop Clusters, *Future Generation Computer Systems*, Vol. 88, pp. 40-54, November 2018.
- [32] Sowkuntla P., Prasad P.S.V.S, Mapreduce Based Improved Quick Reduct Algorithm With Granular Refinement Using Vertical Partitioning Scheme, *Knowledge-Based Systems*, Vol. 18915, Article 105104, February 2020
- [33] Yao, X., Mokbel, M.F., Alarabi, L., Eldawy, A., Zhu, D., Spatial Coding-based Approach for Partitioning Big Spatial Data in Hadoop, *Computers & Geosciences*, Vol. 106, pp. 60-67, Sep. 2017.

- [34] Roy, S., Gupta, S., Omkar, S.N., Case study on: Scalability of Preprocessing Procedure of Remote Sensing in Hadoop, *Procedia Computer Science*, Vol. 108, pp. 1672-1681, 2017.
- [35] Sharma, N., Bagga, S., Girdhar, A., Novel Approach for Denoising Using Hadoop Image Processing Interface, *Procedia Computer Science*, Vol. 132, pp.1327-1350, 2018.
- [36] Demchenko Y., Cees de L., Membrey P., Defining Architecture Components Of The Big Data Eco System. In *International Conference On Collaboration Technologies And Systems (CTS)* Pp. 104-112, 2014.
- [37] Uzunkaya, C., Ensari T., Hadoop Ecosystem and its Analysis on Tweets, *Procedia-Social and Behavioral Sciences*, Vol. 1953, pp.1890-1897, July 2015.
- [38] <https://hadoop.apache.org>.
- [39] https://hadoop.apache.org/docs/r1.2.1/mapred_tutorial.html, Erişim Tarihi: 15.12.2019.
- [40] <https://jar-download.com/>, Erişim Tarihi: 15.12.2019.
- [41] <https://ckan.coegss.hlr.de/>, Erişim Tarihi: 15.12.2019.

ÖZGEÇMİŞ

Suat ERDOĞAN, 07.05.1987’de İstanbul Eyüp’te doğdu. İlk, orta ve lise eğitimini İstanbul’da tamamladı. 2005 yılında Profilo Anadolu Meslek Lisesi’nden mezun oldu. 2007 yılında başladığı İstanbul Kültür Üniversitesi Bilgisayar Teknolojisi ve Programlamadan mezun oldu. Sonrasında 2010 yılında başladığı Anadolu Üniversitesi İşletme Fakültesi’nden mezun oldu, ayrıca 2013 yılında başladığı Karabük Üniversitesi Bilgisayar Mühendisliği Bölümü’nü 2016 yılında bitirerek lisans eğitimini tamamladı. 2018 yılında başladığı yüksek lisans eğitimine Sakarya Üniversitesi Mühendislik Yönetimi Bölümü’nde devam etmektedir.