

Sanel Kandić

**GRUPISANJE I SEGMENTACIJA PODATAKA I SLIKA
KROZ PRIMJENU GRAFOVSKIH ELEMENATA**

- master rad -

Podgorica, 2024.

UNIVERZITET CRNE GORE
ELEKTROTEHNIČKI FAKULTET

Sanel Kandić

**GRUPISANJE I SEGMENTACIJA PODATAKA I SLIKA
KROZ PRIMJENU GRAFOVSKIH ELEMENATA**

- master rad -

Podgorica, 2024.

PODACI I INFORMACIJE O STUDENTU

Ime i prezime:

Sanel Kandić

Datum i mjesto rođenja:

19.03.2000. godine, Plav

Naziv završenog osnovnog studijskog programa i godina završetka studija:

Studijski program Primijenjenog računarstva, Elektrotehnički fakultet, Univerzitet Crne Gore, 180 ECTS kredita, 2021. godine.

INFORMACIJE O MASTER RADU

Naziv master studija:

Master studije primijenjenog računarstva

Naslov rada:

Grupisanje i segmentacija podataka i slika kroz primjenu grafovskih elemenata

Fakultet na kojem je rad odbranjen:

Elektrotehnički fakultet

UDK, OCJENA I ODBRANA MASTER RADA

Datum prijave magistarskog rada:

15.05.2024.

Datum sjednice Vijeća na kojoj je prihvaćena tema:

16.09.2024.

Komisija za ocjenu/odbranu rada:

Prof. dr Miloš Daković, ETF Podgorica, predsjednik
Prof. dr Ljubiša Stanković ETF Podgorica, mentor
Doc. dr Miloš Brajović, ETF Podgorica, član

Mentor:

Prof. dr Ljubiša Stanković

Datum odbrane:

27.12.2024

Ime i prezime autora: Sanel Kandić, BApp

ETIČKA IZJAVA

U skladu sa članom 22 Zakona o akademskom integritetu i članom 18 Pravila studiranja na master studijama, pod krivičnom i materijalnom odgovornošću, izjavljujem da je master rad pod naslovom

"GRUPISANJE I SEGMENTACIJA PODATAKA I SLIKA KROZ PRIMJENU GRAFOVSKIH ELEMENATA"

moje originalno djelo.

Podnosilac izjave,

Kandić Sanel

Sanel Kandić, BApp

U Podgorici, dana 27.12.2024. godine

Izvod teze

U radu se istražuje primjena grafovskih algoritama u svrhe grupisanja i segmentacije podataka i slika, s posebnim naglaskom na različite pristupe i algoritme. Kroz analizu performansi algoritama, kao što su Spektralno grupisanje i MST za grupisanje podataka, te Random Walks i Felzenszwalb-Huttenlocher za segmentaciju slika, rad pruža uvid u efikasnost i pouzdanost ovih tehnika u realnim scenarijima. Eksperimentalni dio istražuje otpornost algoritama na šum i razne vrste ometajućih faktora, što omogućava cjelovitu evaluaciju njihove robusnosti i preciznosti u segmentaciji i grupisanju složenih podataka

Ključne riječi: grupisanje podataka, segmentacija slike, algoritmi, robusnosti

Abstract

The thesis explores the application of graph-based algorithms for data clustering and image segmentation, with a particular emphasis on different approaches and algorithms. Through performance analysis of algorithms such as Spectral clustering and MST for data clustering, and Random Walks and Felzenszwalb-Huttenlocher for image segmentation, the thesis provides insights into the efficiency and reliability of these techniques in real-world scenarios. The experimental part examines the algorithms' resistance to noise and various types of disruptive factors, enabling a comprehensive evaluation of their robustness and accuracy in segmenting and clustering complex data.

Keywords: data clustering, image segmentation, robustness

Sadržaj

Izvod teze	1
Abstract	2
Slike	5
Tabele	7
Uvod	9
1. GRUPISANJE PODATAKA	14
2. SEGMENTACIJA SLIKE	18
2.1 Tehnike segmentacije	22
3. ALGORITMU ZA GRUPISANJE PODATAKA	23
3.1 Spektralno grupisanje	23
3.2 MST algoritam	27
4. ALGORITMI ZA SEGMENTACIJU SLIKA	34
4.1 Random Walks	34
4.2 Felzenszwalb-Huttenlocher algoritam	38
5. EKSPERIMENTALNI DIO	42
5.1 Datasetovi	43
5.1.1 Iris dataset	43
5.1.2 Wine dataset	44
5.1.3 Breast Cancer dataset	45
5.1.4 BSDS500 (Berkeley Segmentation Dataset 500)	46
5.1.5 Datasetovi sa medicinskim slikama	46
5.2 Metrike/metode	47
5.2.1 Silhouette Score	47
5.2.2 Calinski-Harabasz Index	48
5.2.3 Dice koeficijent	50
5.2.4 Jaccard-ov indeks	51

5.3 Spektralno grupisanje	52
5.3.1 Iris dataset	52
5.3.2 Wine dataset	54
5.3.3 Breast Cancer dataset	56
5.4 MST algoritam	58
5.4.1 Iris dataset	58
5.4.2 Wine dataset	60
5.4.3 Breast Cancer dataset	62
5.5 Random Walks algoritam	64
5.5.1 BSDS500 (Berkeley Segmentation Dataset 500)	64
5.5.2 Dataset sa medicinskim slikama	67
5.6 Felzenszwalb-Huttenlocher algoritam	70
5.6.1 BSDS500 (Berkeley Segmentation Dataset 500)	70
5.6.2 Dataset sa medicinskim slikama	73
Zaključak	76
Literatura	77

Slike

Slika 1. Grupisanje zasnovano na particiji.....	15
Slika 2. Hijerarhijsko grupisanje.....	16
Slika 3. Grupisanje zasnovano na gustini	17
Slika 4. Primjeri segmentacije slike	18
Slika 5. Primjer instance segmentacije	20
Slika 6. Primjer semantičke segmentacije	21
Slika 7. Primjer panoptičke segmentacije.....	21
Slika 8. Neorijentisani graf	25
Slika 9. Rezultati grupisanja	27
Slika 10. Graf sa početnim čvorovima i ivicama	29
Slika 11. Dodavanje ivice 7-6 bez formiranja ciklusa.....	29
Slika 12. Dodavanje ivice 8-2 bez formiranja ciklusa.....	30
Slika 13. Dodavanje ivice 6-5 bez formiranja ciklusa.....	30
Slika 14. Dodavanje ivice 0-1 bez formiranja ciklusa.....	30
Slika 15. Dodavanje ivice 2-5 bez formiranja ciklusa.....	31
Slika 16. Dodavanje ivice 2-3 bez formiranja ciklusa.....	31
Slika 17. Dodavanje ivice 0-7 bez formiranja ciklusa.....	32
Slika 18. Dodavanje ivice 3-4 bez formiranja ciklusa.....	32
Slika 19. Formiranje klastera uz prag od 7 ivica.....	33
Slika 20. Originalna slika sa dodanom bukom za testiranje segmentacije	35
Slika 21. Marker koji predstavljaju sjemena za objekat i pozadinu	35
Slika 22. Rezultat segmentacije	36
Slika 23. Originalna slika (lijevo), označeni markeri, koji predstavljaju sjemena za objekat i pozadinu (sredina) i segmentirana slika (desno)	36
Slika 24. Originalna slika (lijevo), označeni markeri uz pomoć Otsuovog metoda, koji predstavljaju sjemena za objekat i pozadinu (sredina) i segmentirana slika (desno).....	37
Slika 25. Prikaz piksela slike s intenzitetima	39
Slika 26. Segmentacija piksela s N4 susjednim povezivanjem	39
Slika 27. Povezivanje piksela sa težinama ivica prema razlikama intenziteta	40
Slika 28. Podjela piksela u segmente	40
Slika 29. Originalna i segmentirana slika uz pomoć Felzenszwalb-Huttenlocher algoritma	41

Slika 30. Prikaz koda za učitavanje i pregled Iris podataka	43
Slika 31. Prikaz koda za učitavanje i pregled Wine podataka	44
Slika 32. Prikaz koda za učitavanje i pregled Breast Cancer podataka.....	45
Slika 33. Ilustracija Dice koeficijenta.....	50
Slika 34. Ilustracija Jaccard-ovog indeksa.....	51
Slika 35. Spektralno grupisanje na Iris skupu podataka	52
Slika 36. Spektralno grupisanje na Iris skupu podataka sa bukom	53
Slika 37. Spektralno grupisanje na Wine skupu podataka.....	54
Slika 38. Spektralno grupisanje na Wine skupu podataka sa bukom	55
Slika 39. Spektralno grupisanje na Breast Cancer skupu podataka	56
Slika 40. Spektralno grupisanje na Breast Cancer skupu podataka sa bukom.....	57
Slika 41. MST grupisanje na Iris skupu podataka.....	58
Slika 42. MST grupisanje na Iris skupu podataka sa bukom.....	59
Slika 43. MST grupisanje na Wine skupu podataka	60
Slika 44. MST grupisanje na Wine skupu podataka sa bukom.....	61
Slika 45. MST grupisanje na Breast Cancer skupu podataka	62
Slika 46. MST grupisanje na Breast Cancer skupu podataka sa bukom	63
Slika 47. Prikaz originalnih slika i segmentiranih verzija korišćenjem Random Walks algoritma na dataset-u	65
Slika 48. Prikaz originalnih slika i segmentiranih verzija korišćenjem Random Walks algoritma na dataset-u sa bukom.....	66
Slika 49. Prikaz originalnih slika i segmentiranih verzija korišćenjem Random Walks algoritma na dataset-u	68
Slika 50. Prikaz originalnih slika i segmentiranih verzija korišćenjem Random Walks algoritma na dataset-u sa bukom.....	69
Slika 51. Prikaz originalnih slika i segmentiranih verzija korišćenjem Felzenszwalb-Huttenlocher algoritma na dataset-u.....	71
Slika 52. Prikaz originalnih slika i segmentiranih verzija korišćenjem Felzenszwalb-Huttenlocher algoritma na dataset-u sa bukom.....	72
Slika 53. Prikaz originalnih slika i segmentiranih verzija korišćenjem Felzenszwalb-Huttenlocher algoritma na dataset-u.....	74
Slika 54. Prikaz originalnih slika i segmentiranih verzija korišćenjem Felzenszwalb-Huttenlocher algoritma na dataset-u sa bukom.....	75

Tabele

Tabela 1. Matrica sličnosti (W) koja prikazuje povezanost između gradova kroz direktne veze, gde su sličnosti postavljene na 1 za povezane gradove i na 0 za nepovezane.....	25
Tabela 2. Matrica stepena (D) koja prikazuje ukupne vrijednosti sličnosti svakog grada sa svim ostalim gradovima, sa dijagonalnim vrijednostima koje predstavljaju stepen svakog grada	26
Tabela 3. Laplasijan matrica (L) dobijena oduzimanjem matrice sličnosti (W) od matrice stepena (D), korišćena za dalju analizu u spektralnom grupisanju.....	26
Tabela 4. Sortirani rezultati.....	29
Tabela 5. Evaluacija kvaliteta grupisanja Iris skupa podataka primjenom Spektralnog grupisanja	52
Tabela 6. Evaluacija kvaliteta grupisanja Iris skupa podataka primjenom Spektralnog grupisanja sa bukom.....	53
Tabela 7. Evaluacija kvaliteta grupisanja Wine skupa podataka primjenom Spektralnog grupisanja	54
Tabela 8. Evaluacija kvaliteta grupisanja Wine skupa podataka primjenom Spektralnog grupisanja sa bukom.....	55
Tabela 9. Evaluacija kvaliteta grupisanja Breast Cancer skupa podataka primjenom Spektralnog grupisanja.....	56
Tabela 10. Evaluacija kvaliteta grupisanja Breast Cancer skupa podataka primjenom Spektralno grupisanje algoritma sa bukom	57
Tabela 11. Evaluacija kvaliteta grupisanja Iris skupa podataka primjenom MST algoritma	58
Tabela 12. Evaluacija kvaliteta grupisanja Iris skupa podataka primjenom MST algoritma sa bukom.....	59
Tabela 13. Evaluacija kvaliteta grupisanja Wine skupa podataka primjenom MST algoritma	60
Tabela 14. Evaluacija kvaliteta grupisanja Wine skupa podataka primjenom MST algoritma sa bukom.....	61
Tabela 15. Evaluacija kvaliteta grupisanja Breast Cancer skupa podataka primjenom MST algoritma	62
Tabela 16. Evaluacija kvaliteta grupisanja Breast Cancer skupa podataka primjenom MST algoritma sa bukom	63
Tabela 17. Evaluacija kvaliteta segmentacije primjenom Random Walks algoritma na dataset-u	64

Tabela 18. Evaluacija kvaliteta segmentacije primjenom Random Walks algoritma na dataset-u sa bukom	66
Tabela 19. Evaluacija kvaliteta segmentacije primjenom Random Walks algoritma na dataset-u sa medicinskim slikama	67
Tabela 20. Evaluacija kvaliteta segmentacije primjenom Random Walks algoritma na dataset-u sa medicinskim slikama sa bukom	69
Tabela 21. Evaluacija kvaliteta segmentacije primjenom Felzenszwalb-Huttenlocher algoritma na dataset-u	70
Tabela 22. Evaluacija kvaliteta segmentacije primjenom Felzenszwalb-Huttenlocher algoritma na dataset-u sa bukom	72
Tabela 23. Evaluacija kvaliteta segmentacije primjenom Felzenszwalb-Huttenlocher algoritma na dataset-u sa medicinskim slikama	73
Tabela 24. Evaluacija kvaliteta segmentacije primjenom Felzenszwalb-Huttenlocher algoritma na dataset-u sa medicinskim slikama sa bukom	75

Uvod

Naslov rada "**Grupisanje i segmentacija podataka i slika kroz primjenu grafovskih elemenata**" odražava ključne teme koje su obrađene u ovom radu.

U današnjem digitalnom dobu, količina podataka rapidno raste u različitim sektorima i oblastima rada kao što su medicina, ekonomija, telekomunikacije, obrazovanje, naučna istraživanja i druge. Ovaj razvoj se u velikom procentu dešava usljed napretka tehnologije. Tako veliku količinu podataka je potrebno grupisati prema srodnosti u manje potkategorije radi lakše manipulacije i rada nad njima.

Grupisanje podataka igra ključnu ulogu u organizaciji i analizi obimnih skupova informacija, dok se segmentacijom olakšava analiza sličnih informacija i prilagođavaju se radne strategije prema karakteristikama. U radu je analizirano grupisanje i segmentacija podataka, proučavajući upotrebu grafova u tom kontekstu. Primjenom grafova je istražena vizualizacija i razumijevanje veza između različitih podataka, analize i otkrivanja klastera, pomažući pri identifikaciji zajedničkih karakteristika.

S obzirom na sveprisutnost i rastuće potrebe za podacima, istraživanjem ove teme nad određenim setom podataka istražena je efektivnost različitih algoritama za grupisanje i segmentaciju podataka.

Predmet istraživanja

Predmet istraživanja teme je primjena različitih algoritama koji se koriste za grupisanje podataka i segmentaciju slika. Analizom dostupnih podataka poput performansi, tehničkih karakteristika i slično, utvrđena je efikasnost istih u segmentaciji i grupisanju. Odrađena je uporedna analiza alogritama za segmentaciju podataka i slika kao što su: Spektralno grupisanje, MST algoritam, Random Walks i Felzenszwalb-Huttenlocher (Efficient Graph-Based Image Segmentation). Navedenim algoritmima su istražene preformanse koje direktno utiču na kvalitet grupisanja i segmentacije u obradi podataka i slika.

Motivi i ciljevi istraživanja

U savremenom digitalnom okruženju, gdje količina dostupnih podataka rapidno raste, istraživanje efikasnih metoda grupisanja i segmentacije postaje neophodno za izdvajanje ključnih informacija iz obimnih skupova podataka. Analizom performansi algoritama u stvarnom svijetu,

istraživanje pruža praktične smjernice za suočavanje s izazovima analize obimnih i različitih podataka.

Svrha istraživanja:

- Utvrđivanje načina na koji strukturne metode doprinose unapređenju grupisanja podataka i segmentacije slika.
- Analizirajući uticaj različitih algoritama na efikasnost i performanse identifikovane su najbolje prakse i strategije za integraciju strukturnih metoda u već postojeće algoritme grupisanja i segmentacije.

Ciljevi istraživanja:

Istraživanje ima za cilj proučavanje doprinosa grafova u efikasnosti grupisanja i segmentacije podataka, s fokusom na ponašanju ključnih algoritama poput Spektralno grupisanje, MST algoritam, Random Walks, Felzenszwalb-Huttenlocher. Analizirane su i performanse algoritama u različitim vrstama podataka i slika, s naglaskom na njihovu praktičnu primjenu, te su identifikovane konkretne primjene rezultata u stvarnim scenarijima.

Pregled dosadašnjih istraživanja

Autori rada [1] opisuje primjenu K-means algoritma za segmentaciju slika kroz subtractive clustering algoritam za generisanje početnih centroida, dok se istovremeno primjenjuje djelimično istežanje kontrasta radi poboljšanja kvaliteta originalne slike. Takođe, u radu se koristi i median filter za poboljšanje segmentirane slike.

Autori rada [12] proučavaju primjenu superpixela u segmentaciji slika, ističući smanjenje računalnog vremena i otpornost na šum. Autori zaključuju da superpixeli doprinose kompaktnosti i regularnosti u segmentaciji slika, podstičući daljnja istraživanja u ovom području.

Autori rada [5] predstavljaju Visual Clustering, brzi algoritam klasterovanja zasnovan na obučenom modelu segmentacije slika. Visual Clustering je inspirisan načinom na koji ljudi klasteruju podatke: plotujući skupove podataka u 2D i identifikujući grupe sličnih tačaka.

Autori naučnog istraživanja [14] istražuju poboljšanja segmentacije slika primjenom genetskog algoritma, K-means i detekcije kontura Sobelovim filterom. Proces uključuje hibrid K-means i Sobelovog filtra za klasifikaciju piksela i detekciju rubova regija uključujući primjenu

genetskog algoritma za poboljšanje kvalitete središta klasa. Ispitivanje učinkovitosti metode sprovodi se na različitim slikama, upoređujući je s K-means algoritmom.

Autori istraživanja [2] istražuju različite strategije klasterovanja i na taj način prepoznaje da svaka od njih ima svoj sistem i proces za obavljanje specifičnih zadataka. Autori ističu da se K-means klasterovanje ističe kao efikasan sistem koji dobro funkcioniše čak i na velikim skupovima podataka, ali može proizvesti neefikasne rezultate u prisustvu buke. U radu je utvrđeno da ova slabost može biti prevaziđena korišćenjem fuzzy C-means klasterovanja i drugih tehnika. Autori sugeriše na to da treba pažljivo birati strategiju klasterovanja, jer neprikladan izbor može dovesti do prekomjerne ili nedovoljne segmentacije.

Autori rada [8] ističu potrebu za organizacijom podataka u grupe u mnogim naučnim oblastima, ukazujući na izazove u traženju optimalnih pristupa klasterovanju. Sa raznovrsnim potrebama korisnika i aplikacija, autori podsjećaju zajednice mašinskog učenja i prepoznavanja oblika na ključne izazove u unapređenju razumijevanja klaster analize podataka.

U naučnom istraživanju [3] autori istražuju primjenu segmentacije za klasifikaciju područja na vazдушnim slikama, koristeći K-means algoritam. Oni prepoznaju nedostatak povećanog broja računanja udaljenosti u ovom algoritmu i predstavljaju membership K-MEANS kao predloženo rješenje koje smanjuje vrijeme izvršavanja i poboljšava tačnost.

Autor [10] opisuje pregled medicinske segmentacije slika, naglašavajući prednosti i nedostatke tehnika poput pragova, regijskih metoda, klusterskih tehnika, umjetnih neuronskih mreža i deformabilnih modela. Ističe potrebu za daljnjim istraživanjem radi razvoja univerzalnog segmentacijskog pristupa, pri čemu se deformabilni modeli ističu kao najrobustnija i preciznija tehnika, posebno pri rješavanju komplikovanih medicinskih problema.

Autori istraživačkog rada [7] zaključuju da je klaster analiza i dalje aktivno područje razvoja, s primjetnim radom u statistici, računarskim naukama, prepoznavanju oblika i kvantizaciji vektora.

Autori [13] u ovom radu predstavljaju rješenje za dva problema vezana uz primjenu Fuzzy C-Means (FCM) algoritma za segmentaciju slika. Prvi problem odnosi se na detekciju broja klastera, a autor predlaže dvije metode: DNCH koja koristi informacije histograma i DNCN koja koristi neuronske mreže.

U ovom naučnom istraživanju [6] autori razvijaju i evaluiraju grafovski metod za segmentaciju slika. Koristeći grafove bliskosti, slike se dijele u segmente susjednih piksela sa sličnim intenzitetima ili bojama. Metod je jednostavan i vrlo brz, a zbog iskorišćavanja piksela u jezgrima segmenata koji su generalno stabilni i pouzdani, metoda je otporna na šum.

U naučnom istraživanju [11] autori istražuju korisne metode segmentacije slika poput pragova, rubova, regija, klastera i boje. Ova studija pruža korisne informacije za istraživače koji

se bave segmentacijom slika, naglašavajući važnost ovog procesa unutar modela obrade slika. Autori zaključuju da je segmentacija slika ključna komponenta u modelima obrade slika.

Autori rada [9] analiziraju metode klasterovanja, osvrćući se na njihove prednosti i mane. Navode da im je cilj razviti novu hibridnu tehniku kombinovanjem evolutivnih metoda i optimizacije, poput Genetskih Algoritama i Optimizacije Kolonijom Mrava, kako bi poboljšali efikasnost klasterisanja podataka. Takođe, istražuju praktičnu primjenu ovih algoritama, fokusirajući se na izazove u performansama, skalabilnosti i obradi velikih dimenzija.

Autori istraživačkog rada [4] predstavljaju novi pristup segmentaciji slika zasnovan na DP algoritmu klasterovanja. Ovaj metod segmentacije sposoban je direktno određivati broj klastera i centre na osnovu odlučujućeg grafa, koji se sastoji od gustine ρ i udaljenosti δ .

Naučne metode

Za ispitivanje istraživačkih pitanja vezanih za datu temu, primijenjene su različite metode, uključujući:

- Istraživanje - Istraživanje algoritama za grupisanje i segmentaciju podataka i slika.
- Testiranje - Testiranje performansi svakog od navedenih algoritama
- Uporedna analiza - Uporedna analiza performansi svakog od algoritama

Očekivani rezultati

Istraživanjem su postignuti rezultati u optimizaciji algoritama za grupisanje i segmentaciju podataka i slika, s naglaskom na poboljšanju preciznosti u segmentaciji kompleksnih struktura. Doprinos rada ogleda se u proširenju znanja o efikasnim tehnikama grupisanja i segmentacije podataka, s osvrtom na grafovske elemente. Kroz ovo istraživanje, istraženi su praktični alati za analitičke procese. Istraživanjem je prikazana perspektiva primjene razvijenih algoritama u analizi obimnih skupova podataka, čime se omogućava dublje razumijevanje strukture i uzoraka u velikim skupovima podataka, doprinoseći unapređenju analitičkih metoda u domenima informacionih tehnologija i društvenih interakcija.

Struktura master rada

1. Uvod

Objašnjenje značaja i konteksta teme grupisanja i segmentacije podataka i slika primjenom grafovskih pristupa. Postavljanje istraživačkih pitanja i uvid u relevantnost teme kroz različite domene.

2. Osnovni koncepti grafova

Opis osnovnih pojmova i tehnika u radu s grafovima, s naglaskom na različite vrste grafova i njihove karakteristike koje su relevantne za grupisanje i segmentaciju podataka.

3. Algoritmi za grupisanje

Analiza primijenjenih metoda grupisanja podataka, s posebnim fokusom na Spektralno grupisanje i MST algoritam. Prikaz prednosti, ograničenja i usporedna analiza ovih metoda na različitim skupovima podataka.

4. Algoritmi za segmentaciju:

Detaljno proučavanje algoritama za segmentaciju slike, uključujući Random Walks i Felzenszwalb-Huttenlocher algoritam. Prikaz primjera primjene ovih tehnika na različitim vrstama slika i analiza performansi.

5. Eksperimentalni dio

Evaluacija performansi navedenih metoda u kontekstu otpornosti na šum i tačnosti segmentacije. Analiza rezultata na različitim dataset-ovima, uključujući slike s medicinskom primjenom i standardne skupove podataka.

6. Zaključak

Rezime ključnih nalaza rada, doprinosi u oblasti grupisanja i segmentacije podataka i slika, te prijedlozi za daljnje istraživanje i poboljšanja.

1. GRUPISANJE PODATAKA

Grupisanje podataka definiše se kao proces raspoređivanja podataka na osnovu inherentnih sličnosti bez unaprijed definisanih kategorija. Ovaj proces donosi brojne prednosti, uključujući pojednostavljivanje složenih podataka, otkrivanje skrivenih struktura i pružanje podrške u donošenju odluka. Iako se intuitivno može pretpostaviti da grupisanje podataka znači samo njihovo raspoređivanje u različite grupe, važno je razmotriti širu upotrebu ovog koncepta, posebno u kontekstu njegove primjene u realnim situacijama.

Grupisanje podataka je neophodno jer omogućava analizu na osnovu opštih sličnosti između entiteta ili stavki, umjesto na osnovu njihovih individualnih vrijednosti. Na primjer, platforme za strimovanje, kao što je Netflix, koriste ovu tehniku da svrstaju filmove i serije u kategorije poput "triler", "animacija", "dokumentarci" i "drama", čime olakšavaju korisnicima pronalaženje sadržaja koji ih zanima. Ovaj pristup omogućava personalizovano iskustvo korisnicima, jer im se, na osnovu njihovih preferenci, automatski preporučuje novi sadržaj, stvarajući obostranu korist za korisnika i platformu [15].

Na primjeru maloprodajne kompanije, može se analizirati segmentacija baze korisnika kako bi se optimizovale marketinške kampanje. Analizom kupovnih obrazaca, kompanije mogu kreirati personalizovane ponude – na primjer, korisniku koji često kupuje luksuznu odjeću mogu se ponuditi popusti na te artikle, dok bi se drugom, koji preferira elektroniku, mogli ponuditi popusti na elektronske uređaje. Ovaj pristup ne samo da može rezultirati povećanom prodajom, već i većim zadovoljstvom kupaca.

Grupisanje podataka koristi se i za kompresiju informacija i izgradnju modela. Identifikovanjem sličnosti među podacima, moguće je predstaviti slične podatke sa manje simbola ili atributa, čime se pojednostavljuje složenost podataka. Pored toga, ako se identifikuju grupe podataka, iz tih grupisanja se može izgraditi model problema koji pruža temelj za analizu i predikciju [16].

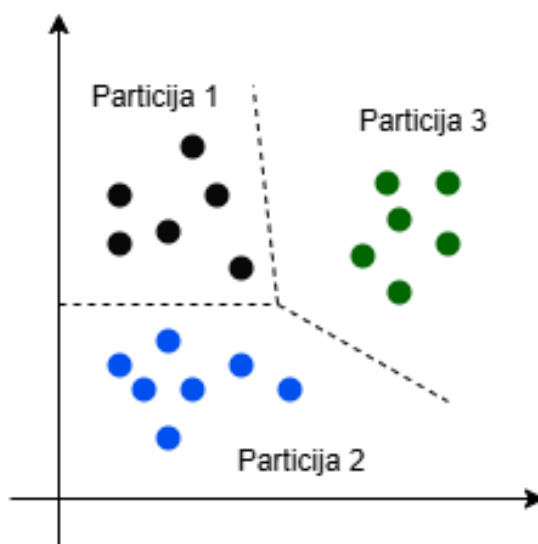
Još jedan važan aspekt grupisanja je otkrivanje relevantnih informacija unutar podataka. Na primjer, Francisco Azuaje i saradnici [17] su razvili sistem zaključivanja na osnovu slučajeva (CBR) koristeći model strukture rastuće ćelije (GCS). U ovom sistemu, podaci se skladište u bazi slučajeva, organizovani po kategorijama. Struktura rastuće ćelije omogućava dinamično dodavanje i uklanjanje podataka iz baze, a kada se predstavi novi upit, sistem identifikuje najrelevantnije slučajeve na osnovu njihove sličnosti sa upitom.

Tehnike grupisanja podataka mogu se podijeliti u tri glavne kategorije:

- **grupisanje zasnovano na particiji,**
- **hijerarhijsko grupisanje i**
- **grupisanje zasnovano na gustini.**

Grupisanje zasnovano na particiji

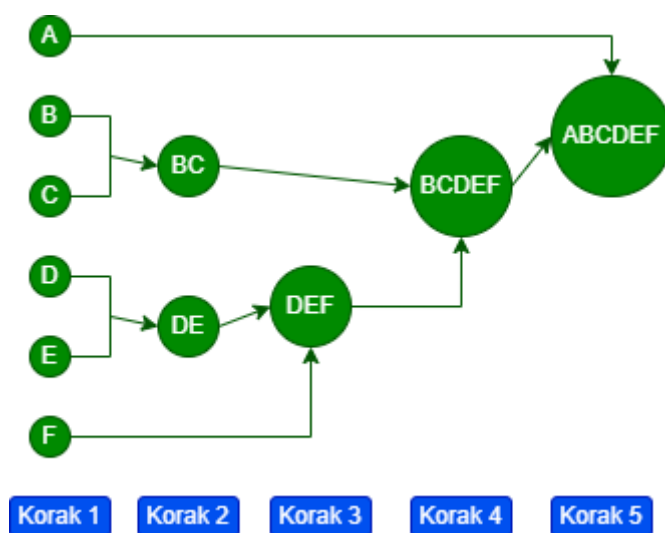
Grupisanje zasnovano na particiji (slika 1), kao što je K-means algoritam, zahtijeva da svaka tačka podataka bude dodijeljena samo jednom klasteru, pri čemu je broj klastera unaprijed definisan. Ova metoda je posebno efikasna kada je cilj brzo i jasno podijeliti podatke na razdvojene grupe. Koristi se u raznim aplikacijama, uključujući kompresiju slika, kategorizaciju dokumenata i segmentaciju korisnika, zbog svoje jednostavnosti i efikasnosti. Jedna od ključnih karakteristika ove tehnike je to što omogućava lako tumačenje rezultata, budući da svaki podatak pripada isključivo jednoj grupi, što olakšava analizu i donošenje odluka na osnovu identifikovanih klastera. Takođe, ova metoda je pogodna za slučajeve u kojima je moguće unaprijed odrediti broj grupa u podacima, čime se unaprijed postavljeni ciljevi grupisanja mogu precizno ispuniti. Ipak, jedna od potencijalnih slabosti ove tehnike je njena osjetljivost na inicijalne postavke i konfiguracije, zbog čega se često moraju testirati različiti parametri kako bi se postigao optimalan rezultat.



Slika 1. Grupisanje zasnovano na particiji

Hijerarhijsko grupisanje

Hijerarhijsko grupisanje (slika 2) gradi strukturu u obliku stabla unutar podataka, gdje se klasteri mogu vizualizovati na različitim nivoima, omogućavajući pregled podataka na različitim granicama grupisanja. Za razliku od metoda zasnovanih na particiji, hijerarhijsko grupisanje ne zahtijeva unaprijed definisan broj klastera, što ga čini fleksibilnijim i pogodnijim za primjene kao što su segmentacija slika, kategorizacija dokumenata i genetičko grupisanje. Ova metoda omogućava dinamično odlučivanje o broju klastera na osnovu strukture podataka, čime se mogu analizirati klasteri na različitim nivoima granularnosti. Dodatna prednost hijerarhijskog grupisanja je mogućnost uvida u odnose između klastera i njihovih podklastera, što olakšava identifikovanje hijerarhijskih veza unutar podataka. Ipak, treba napomenuti da ova tehnika može postati računarski zahtjevna kada se radi sa velikim skupovima podataka, jer je potrebno više memorije i vremena za izgradnju i analizu dendrograma.

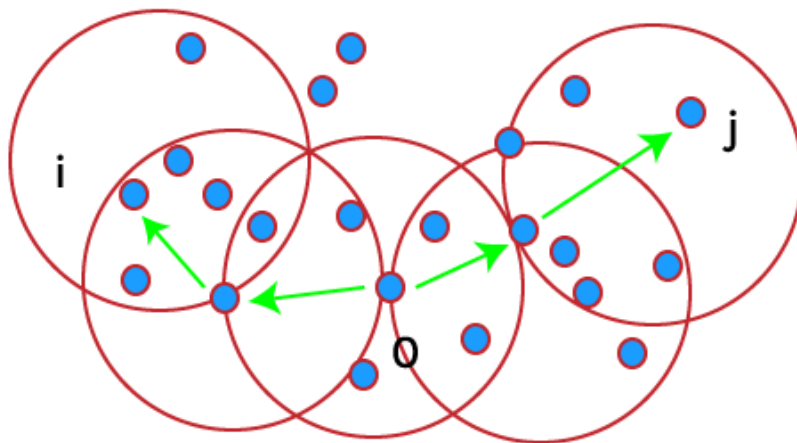


Slika 2. Hijerarhijsko grupisanje

Grupisanje zasnovano na gustini

Grupisanje zasnovano na gustini (slika 3) identifikuje klasterne na osnovu gustine tačaka u prostoru podataka, čime omogućava formiranje klastera proizvoljnih oblika, uključujući i neregularne i nepravilne forme. Ova tehnika pokazuje veliku otpornost na šum u podacima, jer tačke koje nijesu dio nijednog klastera tretiraju se kao šum i isključuju iz daljih analiza. Grupisanje zasnovano na gustini posebno je pogodno za rad sa velikim i kompleksnim skupovima podataka, gdje klasteri nisu jasno razdvojeni ili imaju nepravilne oblike, što ga čini efikasnim u situacijama

gdje tradicionalni algoritmi, poput onih zasnovanih na particiji, ne daju dobre rezultate. Ova tehnika se često koristi u geografskim informacionim sistemima (GIS) za analizu prostornih podataka i u detekciji anomalija, kao što su identifikovanje nepravilnosti u mrežnom saobraćaju ili otkrivanje potencijalnih prijetnji u sigurnosnim sistemima. Iako je robusna i fleksibilna, metoda zahtijeva pažljivo podešavanje parametara gustine kako bi se dobili optimalni rezultati, te može postati manje efikasna u slučajevima sa vrlo različitim gustinama podataka unutar istog skupa [18] [19].



Slika 3. Grupisanje zasnovano na gustini

2. SEGMENTACIJA SLIKE

U digitalnoj obradi slika i računarskoj viziji, segmentacija slika je osnovna tehnika koja podrazumijeva razdvajanje digitalne slike na više segmenata (regija ili objekata), poznatih kao slikovne regije ili objekti (skupovi piksela), s ciljem pojednostavljenja i analize slike (slika 4). Glavna svrha segmentacije je pretvaranje slike u oblik koji je lakše razumljiv i pogodniji za analizu, što čini proces obrade slike efikasnijim fokusiranjem na specifične interesne regije. Segmentacija omogućava prepoznavanje i izdvajanje važnih objekata unutar slike, čime se olakšava obrada, klasifikacija i donošenje odluka u različitim aplikacijama.



Slika 4. Primjeri segmentacije slike

Tipičan zadatak segmentacije prolazi kroz nekoliko koraka:

- Grupisanje piksela u slici na osnovu zajedničkih karakteristika, kao što su boja, intenzitet ili tekstura.
- Dodjeljivanje oznake svakom pikselu, koja označava pripadnost određenom segmentu ili objektu [20].

Rezultat segmentacije je set segmenata koji zajedno prekrivaju cijelu sliku, ili set kontura izdvojenih iz slike. Segmentisana slika se često prikazuje kao maska ili preklapanje koje naglašava različite segmente. Pikseli unutar jedne regije dijele slične karakteristike, dok se susjedne regije jasno razlikuju u tim osobinama.

Primjena segmentacije slika je široka i uključuje mnoge industrije. Na primjer, u medicinskim aplikacijama, segmentacija se koristi za identifikaciju tumora ili drugih abnormalnosti u slikama dobijenim magnetnom rezonancom (MRI) ili kompjuterskom tomografijom (CT). U autonomnim vozilima, segmentacija je ključna za prepoznavanje pješaka, vozila i drugih prepreka na putu, što omogućava sigurno kretanje kroz saobraćaj. Takođe, u poljoprivredi se koristi za analizu satelitskih slika i praćenje zdravlja usjeva, dok je u geološkim istraživanjima neophodna za modeliranje terena i identifikaciju ključnih struktura [21].

Segmentacija slika može se podijeliti na sljedeće kategorije:

- **instanca segmentacije,**
- **semantička segmentacija i**
- **panoptička segmentacija.**

Instanca segmentacija

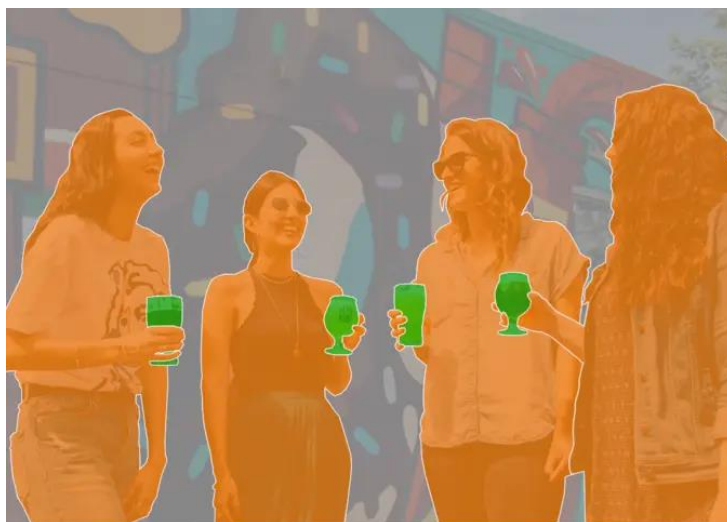
Instanca segmentacije (slika 5) je pristup koji identifikuje, za svaki piksel, kojoj tačno instanci objekta pripada, omogućavajući precizno odvajanje svakog pojedinačnog objekta u slici. Ona detektuje svaki jedinstveni objekat od interesa, čak i ako objekti pripadaju istoj kategoriji. Na primjer, u slici sa više osoba, instanca segmentacije će prepoznati i segmentisati svaku osobu kao poseban objekat, omogućavajući razlikovanje između pojedinaca po njihovim specifičnim karakteristikama, poput visine, položaja ili drugih atributa [22].



Slika 5. Primjer instance segmentacije

Semantička segmentacija

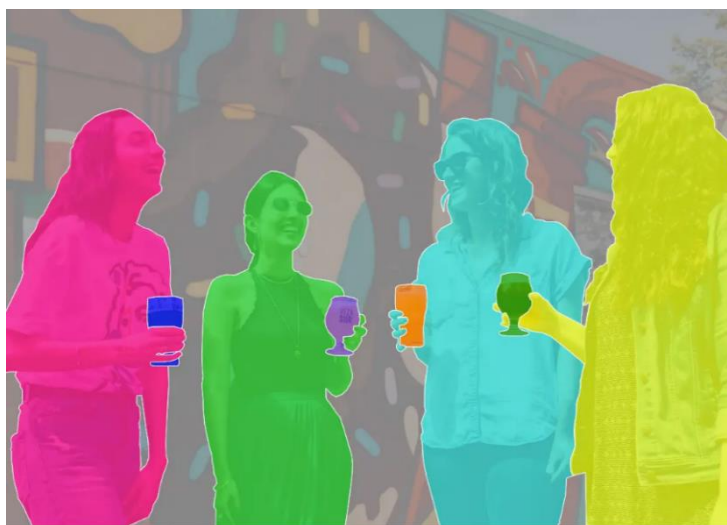
Semantička segmentacija (slika 6) je pristup koji detektuje klasu kojoj svaki piksel pripada, bez obzira na to da li su objekti unutar te klase pojedinačni ili ponavljajući. Na primjer, u slici sa više osoba, svi pikseli koji pripadaju ljudima biće označeni istom klasom, bez obzira na to što predstavljaju različite pojedince. Isto važi i za pozadinske elemente, koji će biti klasifikovani kao jedna cjelina, poput pozadine ili okruženja. Ova tehnika je korisna kada je cilj identifikacija širih kategorija objekata, ali ne omogućava prepoznavanje ili razlikovanje između različitih instanci unutar te klase. Semantička segmentacija se često koristi u aplikacijama gdje je važno razumjeti opštu strukturu slike, kao što su prepoznavanje urbanih scena, satelitske analize ili analiza zemljišta [23].



Slika 6. Primjer semantičke segmentacije

Panoptička segmentacija

Panoptička segmentacija (slika 7) kombinuje semantičku i instanca segmentaciju. Kao kod semantičke segmentacije, panoptička segmentacija identifikuje kojoj klasi pripada svaki piksel na slici. Međutim, ona ide korak dalje, jer poput instanca segmentacije, takođe razlikuje različite instance unutar iste klase. Na primjer, u slici sa više osoba, panoptička segmentacija ne samo da klasifikuje sve ljude u klasu "osoba", već prepoznaje i segmentiše svaku osobu kao poseban objekat. Ovaj pristup omogućava detaljniju analizu slike, pružajući i semantičku informaciju i precizno razlikovanje između pojedinačnih objekata unutar iste klase [24].



Slika 7. Primjer panoptičke segmentacije

2.1 Tehnike segmentacije

Tradicionalne tehnike segmentacije slika, koje su postavile temelje za savremene metode segmentacije bazirane na algoritmima dubokog učenja, koriste metode kao što su pragiranje, detekcija ivica, segmentacija zasnovana na regijama, algoritmi klasterovanja i segmentacija metodom vododjelnica. Ove tehnike se uglavnom oslanjaju na osnovne principe obrade slika, matematičke operacije i heuristiku, kako bi se slika podijelila na značajne regije. Sljedeće metode su među najznačajnijima u tradicionalnoj segmentaciji slika:

- **Pragiranje:** Ova metoda zasniva se na definisanju praga vrijednosti piksela, gdje se pikseli slike klasifikuju u dvije glavne grupe – prednji plan i pozadinu – na osnovu njihovih intenziteta. Pragiranje se često koristi za jednostavno odvajanje objekata od pozadine u binarnim slikama.
- **Detekcija ivica:** Detekcija ivica identifikuje nagle promjene u intenzitetu slike ili diskontinuitete u teksturi. Ova metoda koristi algoritme poput Sobel, Canny i Laplaceovih detektora, koji naglašavaju ivice objekata, olakšavajući dalju obradu slike i identifikaciju objekata.
- **Segmentacija zasnovana na regijama:** Ova tehnika segmentacije dijeli sliku na manje regije, koje se zatim iterativno spajaju na osnovu zajedničkih atributa kao što su boja, intenzitet ili tekstura. Segmentacija zasnovana na regijama je posebno efikasna u prisustvu šuma i nepravilnosti, jer omogućava postepeno poboljšavanje granica regija.
- **Algoritmi klasterovanja:** Ova metoda koristi algoritme kao što su K-means i Gausovi modeli kako bi grupisala piksele sličnih karakteristika, poput boje ili texture, u klastere. Algoritmi klasterovanja su efikasni za segmentaciju slika gdje je potrebno identifikovati više objekata sa sličnim osobinama, a posebno su korisni u složenim scenarijima.
- **Segmentacija metodom vododjelnica:** Ova tehnika posmatra sliku kao topografski model, gdje se pikseli tretiraju kao visine terena. Linije vododjelnica predstavljaju granice između različitih regija na slici, identifikujući ih na osnovu intenziteta i povezanosti piksela. Vododjelnice se često koriste za segmentaciju složenih slika sa više preklapajućih objekata [25].

3. ALGORITMU ZA GRUPISANJE PODATAKA

3.1 Spektralno grupisanje

Spektralno grupisanje (Spectral clustering) je tehnika za particionisanje podataka, zasnovana na teoriji grafova i svojstvenim vektorima matrica sličnosti. Umjesto da direktno radi sa podacima u njihovom izvornom obliku, spektralno grupisanje transformiše podatke u novi prostor gdje se klasteri lakše identifikuju. Ključna prednost spektralnog grupisanja je to što može identifikovati klasterne nepravilnog oblika.

U Spektralnog grupisanju, podaci se predstavljaju kao grafovi, gdje su čvorovi tačke podataka, a ivice označavaju sličnosti između tačaka.

Graf je matematička struktura koja se sastoji od skupa čvorova (tačaka) i ivica (linija) koje povezuju te čvorove. U kontekstu grupisanja podataka, čvorovi predstavljaju pojedinačne tačke podataka, dok ivice predstavljaju sličnosti između tih tačaka. Veća sličnost između tačaka se izražava težim ivicama, dok manja sličnost rezultira manjim težinama ili čak odsustvom veze između čvorova.

Sličnost između tačaka može se izračunati različitim mjerama, ali se najčešće koristi Gaussian kernel:

$$W_{i,j} = \exp\left(-\frac{\|x_i - x_j\|^2}{2\sigma^2}\right)$$

gdje je $W_{i,j}$ sličnost između tačaka x_i i x_j , a σ kontroliše širinu Gaussovog kernela.

Da bi se odredili klasteri, koriste se svojstvene vrijednosti i svojstveni vektori Laplasijan matrice grafa. Laplasijan je matematički objekat koji opisuje globalne osobine grafa. Definicija Laplasijana je:

$$L = D - W$$

gdje je D dijagonalna matrica stepena čvorova koja se izračunava kao:

$$D_i = \sum_{j=1}^n W_{i,j},$$

a W matrica sličnosti.

Cilj spektralne analize je da se izračunaju sopstvene vrijednosti i dekompozicije vektora Laplasijana. Ovi svojstveni vektori definišu novi prostor u kojem je moguće lakše identifikovati klastere. Proces se bazira na rješavanju sljedeće jednačine:

$$Lv = \lambda v$$

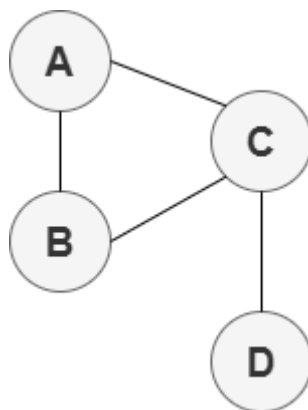
Izračunavanjem nekoliko najmanjih svojstvenih vrijednosti (osim najmanje koja je uvijek nula) i njihovih odgovarajućih svojstvenih vektora, dobijamo novu reprezentaciju podataka u tzv. spektralnom prostoru. Svaka tačka podataka sada je predstavljena novim skupom koordinata, koje se koriste za grupisanje.

Nakon što se podaci transformišu u spektralni prostor, koriste se algoritmi za grupisanje, kao što je **K-means**, kako bi se pronašli klasteri. U ovom novom prostoru, podaci koji pripadaju istom klasteru su bliži jedan drugom, dok su podaci iz različitih klastera međusobno udaljeni [26].

K-means

K-means je popularan algoritam za grupisanje podataka, koji nastoji da podijeli podatke u k klastera na osnovu njihove sličnosti. Proces započinje nasumičnim postavljanjem k centroida, koji predstavljaju centre klastera. Svaki podatak se zatim pridružuje najbližem centru na osnovu udaljenosti (najčešće euklidske), a potom se centroidi ažuriraju tako da predstavljaju sredinu svojih pripadajućih podataka. Ovaj proces se ponavlja sve dok promjene u centroidima postanu minimalne, čime se postiže stabilno grupisanje podataka [27].

U nastavku je dat primjer kako bi se lakše objasnio funkcionisanje algoritma. Na slici 8 je prikazano četiri grada: A, B, C i D. Cilj je da se grupišu gradovi na osnovu njihove blizine. Gradovi koji su bliže jedni drugima trebaju biti u istom klasteru, dok bi oni koji su daleko trebali pripadati različitim klasterima.



Slika 8. Neorijentisani graf

1. Kreiranje matrice sličnosti W

Matrica sličnosti W prikazuje povezanost između gradova. Ako postoji direktna veza između dva grada (gradovi su blizu), sličnost je visoka (1), a ako nema direktne veze, sličnost je 0. Na tabeli 1 se vidi da su gradovi A, B i C povezani jedan sa drugim, pa njihove međusobne sličnosti imaju vrijednost 1, dok je grad D povezan samo sa gradom C, pa su njegove sličnosti sa ostalim gradovima postavljene na 0.

Tabela 1. Matrica sličnosti (W) koja prikazuje povezanost između gradova kroz direktne veze, gde su sličnosti postavljene na 1 za povezane gradove i na 0 za nepovezane.

Gradovi	A	B	C	D
A	0	1	1	0
B	1	0	1	0
C	1	1	0	1
D	0	0	1	0

2. Kreiranje stepen matrice D

Matrica stepena D sadrži zbir sličnosti svakog grada sa svim ostalim gradovima. Zbir za svaki grad stavlja se na dijagonalne elemente. Na tabeli 2 se vidi da grad A ima stepen 2 jer je povezan sa gradovima B i C, dok C ima stepen 3 jer je povezan sa gradovima A, B, i D.

Tabela 2. Matrica stepena (D) koja prikazuje ukupne vrijednosti sličnosti svakog grada sa svim ostalim gradovima, sa dijagonalnim vrijednostima koje predstavljaju stepen svakog grada

Gradovi	A	B	C	D
A	2	0	0	0
B	0	2	0	0
C	0	0	3	0
D	0	0	0	1

3. Izračunavanje Laplasijana matrice

Laplasijan matrica se kreira oduzimanjem matrice sličnosti W od stepena matrice D kao što je prikazano u tabeli 3.

Tabela 3. Laplasijan matrica (L) dobijena oduzimanjem matrice sličnosti (W) od matrice stepena (D), korišćena za dalju analizu u spektralnom grupisanju

Gradovi	A	B	C	D
A	2	-1	-1	0
B	-1	2	-1	0
C	-1	-1	3	-1
D	0	0	-1	1

4. Svojstvena dekompozicija

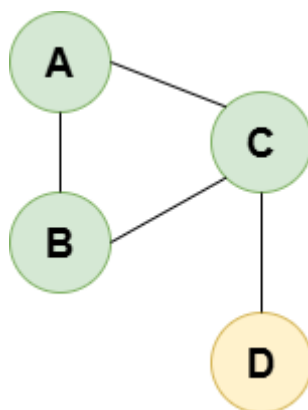
Sada se primjenjuje svojstvena dekompozicija na Laplasijan matrici L, čime se dobijaju svojstvene vrijednosti i svojstveni vektori. Ovi svojstveni vektori omogućavaju da se podaci transformišu u novi prostor, gdje se može lakše izvršiti grupisanje gradova.

5. Grupisanje u spektralnom prostoru

Nakon što su izračunati svojstveni vektori, koristi se K-means algoritam kako bi se konačno izvršilo grupisanje gradova. Gradovi koji imaju slične vrijednosti u spektralnom prostoru grupišu se zajedno.

Rezultat grupisanja (slika 9):

- Klaster 1: Gradovi A, B, i C su grupisani zajedno jer su povezani direktnim vezama.
- Klaster 2: Grad D formira poseban klaster jer je povezan samo sa gradom C [28].



Slika 9. Rezultati grupisanja

3.2 MST algoritam

Razapinjuće stablo je definisano kao podgraf povezanog, neusmjerenog grafa koji ima strukturu stabla i koji uključuje sve čvorove grafa. Jednostavnim jezikom rečeno, to je podskup ivica grafa koji formira stablo (graf bez ciklusa) u kojem je svaki čvor dio tog stabla. Drugim riječima, razapinjuće stablo povezuje sve čvorove grafa koristeći najmanji mogući broj veza (ivica) i bez stvaranja povratnih putanja ili ciklusa. U bilo kojem povezanom grafu, može postojati više različitih razapinjućih stabala, ali svako od njih sadrži isti broj čvorova i $V - 1$ ivica, gdje je V broj čvorova u grafu.

Minimalno razapinjuće stablo (**MST**) ima sve osobine razapinjućeg stabla, uz dodatni uslov da ima najmanju moguću sumu težina među svim mogućim razapinjućim stablima. Kao i kod razapinjućeg stabla, može postojati više mogućih minimalnih razapinjućih stabala za dati graf, u zavisnosti od težina ivica i strukture grafa. MST algoritam koristi se za efikasno povezivanje

tačaka u grafu na način koji minimizira troškove ili resurse potrebne za tu povezanost. MST algoritmi, poput Kruskalovog i Primovog, funkcionišu tako što postepeno dodaju najkraće moguće ivice, izbjegavajući stvaranje ciklusa, sve dok se ne povežu svi čvorovi u grafu.

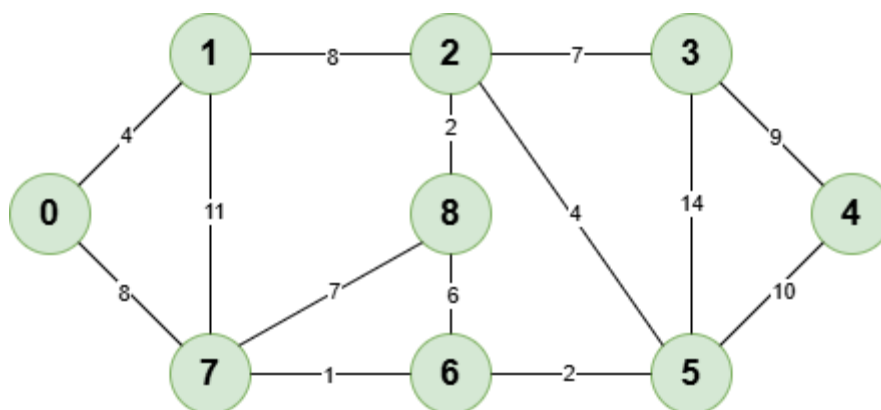
Razapinjuće stablo ima sljedeća ključna svojstva:

1. **Broj čvorova (V)** u grafu i razapinjućem stablu je isti.
2. **Broj ivica** u razapinjućem stablu je uvijek jedan manje od broja tjemena u grafu, dakle $E = V - 1$, gdje je E broj ivica, a V broj čvorova.
3. **Razapinjuće stablo mora biti povezano**
4. **Razapinjuće stablo mora biti aciklično:** To znači da ne smije biti nikakvih ciklusa unutar stabla (nema povratnih putanja između čvorova).
5. **Ukupni trošak (ili težina)** razapinjućeg stabla je definisan kao zbir težina svih ivica koje pripadaju tom stablu.
6. **Više mogućih razapinjućih stabala:** Za dati graf može postojati više razapinjućih stabala [29] [30].

Postoji nekoliko MST algoritama koji se mogu koristiti za pronalaženje minimalnog razapinjućeg stabla iz datog grafa, a fokus će biti stavljen na Kruskalov algoritam za minimalno razapinjuće stablo.

Kod Kruskalovog algoritma, sve ivice datog grafa se sortiraju u rastućem redosljedu. Potom se postepeno dodaju nove ivice i čvorovi u MST, pod uslovom da dodata ivica ne formira ciklus. Najprije se bira ivica sa najmanjom težinom, a na kraju ona sa najvećom težinom. Tokom ovog procesa, prag može biti postavljen kako bi se ograničile ivice koje prelaze određenu težinu, čime se utiče na formiranje klastera i optimizaciju strukture stabla koje se kasnije koristi za grupisanje podataka.

U nastavku je dat primjer kako bi se lakše objasnilo funkcionisanje algoritma. Graf sadrži 9 čvorova i 14 ivica (slika 10). Dakle, minimalno razapinjuće stablo koje se formira imaće $9 - 1 = 8$ ivica. Nakon sortiranja dobijaju se rezultati iz tabele 4.

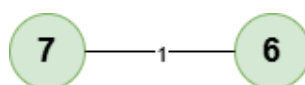


Slika 10. Graf sa početnim čvorovima i ivicama

Tabela 4. Sortirani rezultati

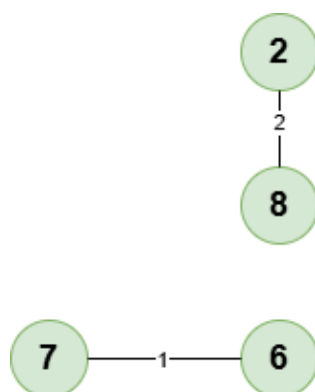
Težina	Izvor	Destinacija	Težina	Izvor	Destinacija
1	7	6	8	7	8
2	8	2	9	0	7
3	6	5	10	1	2
4	0	1	11	3	4
5	2	5	12	5	4
6	8	6	13	1	7
7	2	3	14	3	5

Biraju se sve ivice jedna po jedna iz sortirane liste. Bira se ivica 7-6 i gleda da li je ciklus napravljen, u ovom slučaju nije, pa se ova ivica uključuje (slika 11).



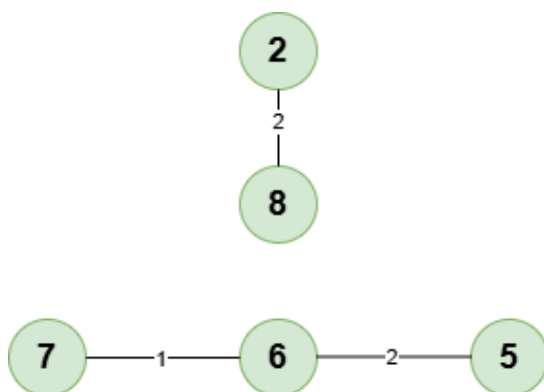
Slika 11. Dodavanje ivice 7-6 bez formiranja ciklusa

Bira se ivica $8 - 2$, ciklus nije formiran, pa se uključuje (slika 12).



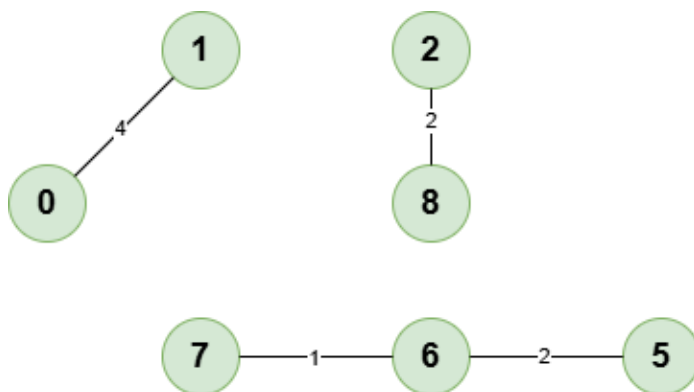
Slika 12. Dodavanje ivice 8-2 bez formiranja ciklusa

Bira se ivica $6 - 5$, ciklus nije formiran, pa se uključuje (slika 13).



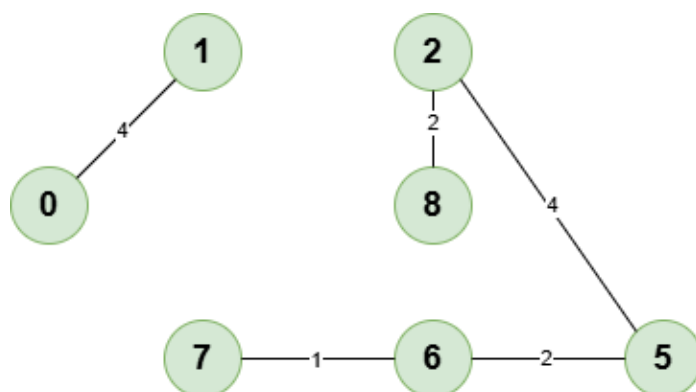
Slika 13. Dodavanje ivice 6-5 bez formiranja ciklusa

Bira se ivica $0 - 1$, ciklus nije formiran, pa se uključuje (slika 14).



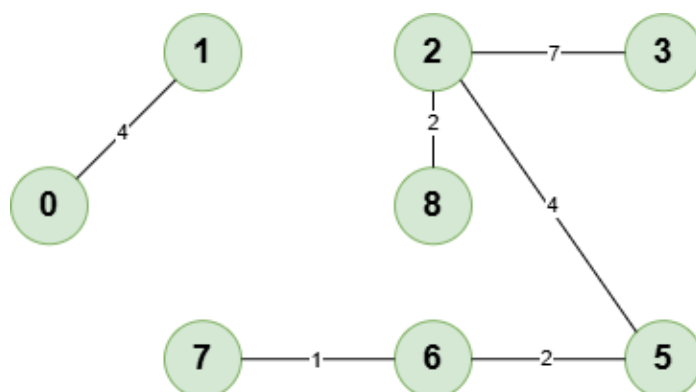
Slika 14. Dodavanje ivice 0-1 bez formiranja ciklusa

Bira se ivica 2 -5, ciklus nije formiran, pa se uključuje (slika 15).



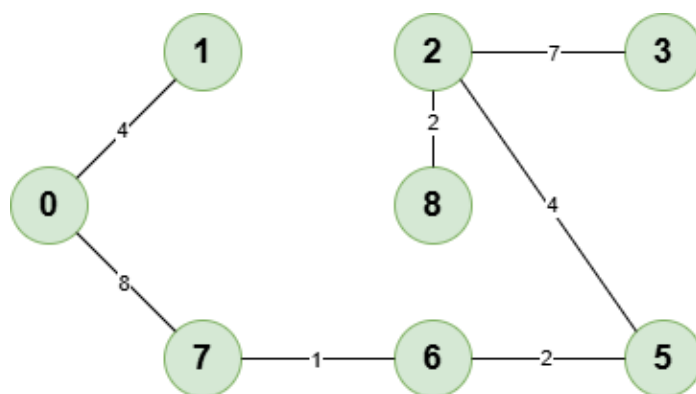
Slika 15. Dodavanje ivice 2-5 bez formiranja ciklusa

Bira se ivica 8 – 6, uključivanje rezultira formiranjem ciklusa, pa se odbacuje. Bira se ivica 2 -3, ciklus nije formiran, pa se uključuje (slika 16).



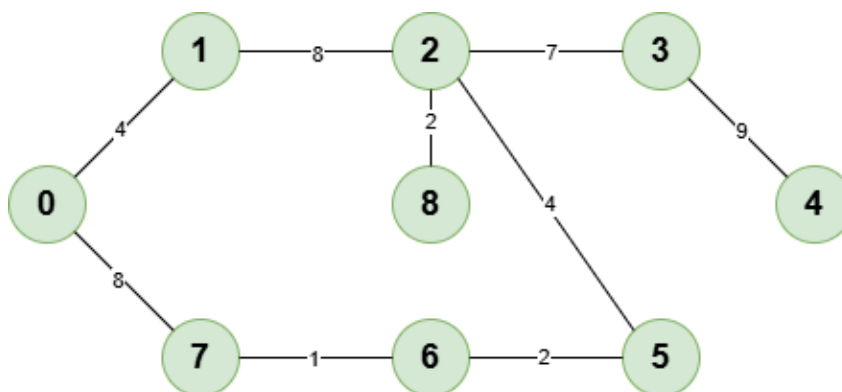
Slika 16. Dodavanje ivice 2-3 bez formiranja ciklusa

Bira se ivica 7 – 8, uključivanje rezultira formiranjem ciklusa, pa se odbacuje. Bira se ivica 0 -7, ciklus nije formiran, pa se uključuje (slika 17).



Slika 17. Dodavanje ivice 0-7 bez formiranja ciklusa

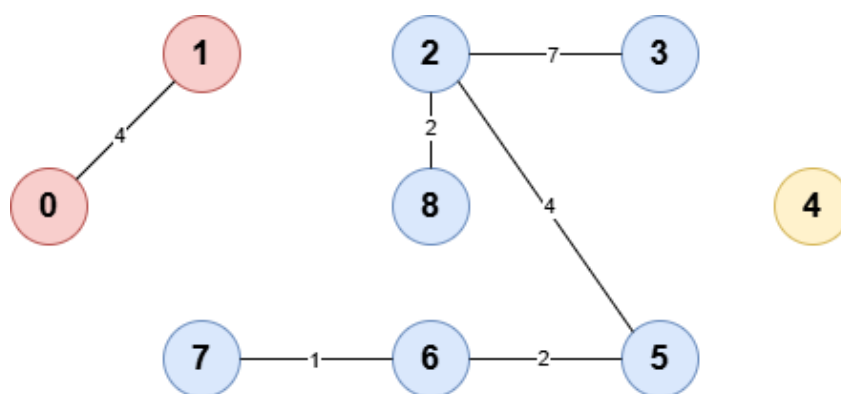
Bira se ivica 1 – 2, uključivanje rezultira formiranjem ciklusa, pa se odbacuje. Bira se ivica 3 -4, ciklus nije formiran, pa se uključuje (slika 18).



Slika 18. Dodavanje ivice 3-4 bez formiranja ciklusa

Slika 18.

Pošto broj ivica uključenih u MST iznosi $(V - 1)$, algoritam ovdje prestaje. U primjeru je odabran prag 7, što rezultira formiranjem tri klastera pri grupisanju podataka (slika 19).



Slika 19. Formiranje klastera uz prag od 7 ivica

4. ALGORITMI ZA SEGMENTACIJU SLIKA

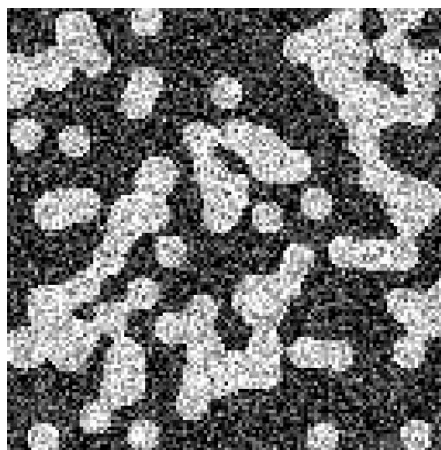
4.1 Random Walks

Algoritam za segmentaciju slike Random Walks koristi se za donošenje odluka o tome kojoj klasi pripada svaki piksel na slici. Slika se predstavlja kao graf, gdje je svaki piksel prikazan čvorom grafa, a susjedni pikseli su povezani ivicama. Težine ivica određuju sličnost između susjednih piksela – veće težine se dodjeljuju većoj sličnosti, dok se manje težine koriste za veće razlike.

Proces segmentacije započinje ručnim označavanjem malog broja piksela sa poznatim oznakama, koji se nazivaju sjemeni (npr. pikseli označeni kao "objekat" ili "pozadina"). Potom se izračunava kojoj oznaci će biti dodijeljeni svi ostali pikseli koji nisu označeni. Svaki piksel koji nije označen zamišlja se kao početna tačka slučajnog hodača – zamišlja se da se slučajni hodač "kreće" po grafu i prelazi na susjedne piksele. Vjerovatnoća da slučajni hodač, koji polazi sa određenog piksela, prvo stigne do jednog od sjemeni sa poznatom oznakom, računa se za sve piksele u slici. Na osnovu ovih vjerovatnoća, dodjeljuje se oznaka onoj klasi čijem sjemenu se hodač najvjerovatnije približi.

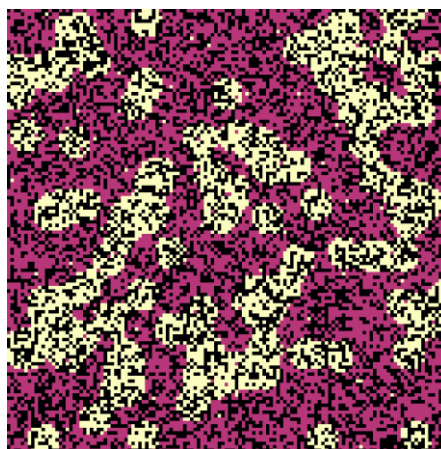
Matematički, ovaj proces se formuliše rješavanjem sistema linearnih jednačina, pri čemu se određuju vjerovatnoće za sve piksele. Ovim pristupom omogućava se integracija informacija o sličnosti između piksela (kroz težine ivica), čime se obezbjeđuje segmentacija koja je robusna na šum ili varijacije u slici [31].

U nastavku je dat primjer kako bi se lakše objasnilo funkcionisanje algoritma. Data je originalna slika (slika 20) koja sadrži razne nijanse sivih tonova. Na slici su prisutne tamne i svijetle oblasti koje je potrebno segmentisati, kako bi objekti bili odvojeni od pozadine. Buka je dodata slici kako bi se simulirali stvarni uslovi i otežao zadatak segmentacije.



Slika 20. Originalna slika sa dodanom bukom za testiranje segmentacije

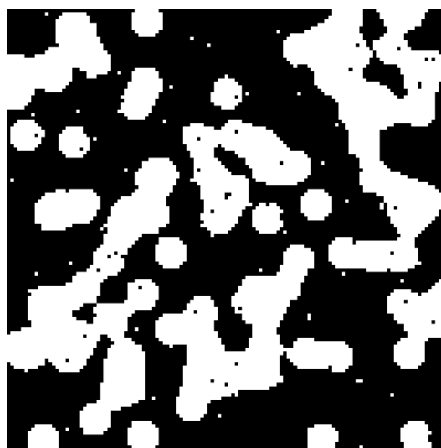
Prije pokretanja algoritma, određeni pikseli su označeni kao sjemena (markeri). Pikseli sa vrijednostima koje su veoma tamne (npr. ispod -0.5, na skali normalizovane slike) označeni su kao "objekat" (oznaka 1), dok su pikseli sa vrijednostima koje su veoma svijetle (npr. iznad 0.5) označeni kao "pozadina" (oznaka 2). Ovakvo postavljanje markera omogućava algoritmu da prepozna koje oblasti pripadaju objektu, a koje pozadini (slika 21).



Slika 21. Markeri koji predstavljaju sjemena za objekat i pozadinu

Algoritam Random Walks koristi informacije iz postavljenih markera kako bi izvršio segmentaciju. Zamišlja se da se sa svakog piksela oslobađa slučajni hodač, koji nasumično prelazi na susjedne piksele. Algoritam izračunava vjerovatnoću da će hodač, koji počinje sa određenog piksela, prvo stići do jednog od markera (bilo "objekta" ili "pozadine"). Na osnovu tih vjerovatnoća, algoritam odlučuje kojoj klasi će piksel pripadati – objektu ili pozadini.

Parametar koji se koristi u algoritmu određuje koliko će hodač biti skloniji prelasku na slične piksele. Na taj način, slični pikseli imaju veće šanse da budu klasifikovani kao ista klasa. Nakon završetka algoritma, svaki piksel je dodijeljen ili klasi "objekta" ili klasi "pozadine", na osnovu izračunate vjerovatnoće. Rezultat je segmentisana slika na kojoj su oblasti koje pripadaju objektu odvojene od onih koje pripadaju pozadini (slika 22) [32].



Slika 22. Rezultat segmentacije

Na slici 23 je prikazan primjer segmentacije na realističnim podacima.



Slika 23. Originalna slika (lijevo), označeni markeri, koji predstavljaju sjemena za objekat i pozadinu (sredina) i segmentirana slika (desno)

Algoritam Random Walks može se kombinovati s drugim metodama za odabir markera, poput Otsuove metode za automatsko određivanje praga, koja omogućava precizniji izbor markera na osnovu distribucije intenziteta piksela na slici. Ovakva kombinacija pomaže u postavljanju pouzdanijih početnih sjemena za segmentaciju, što rezultira boljim performansama algoritma. Na slici 24 je prikazan primjer segmentacije koristeći Random Walks i Otsuov metod gdje metod pomaže algoritmu da preciznije odabere sjemena, za početne uslove segmentacije [33].



Slika 24. Originalna slika (lijevo), označeni markeri uz pomoć Otsuovog metoda, koji predstavljaju sjemena za objekat i pozadinu (sredina) i segmentirana slika (desno)

Otsuov metod za automatsko određivanje praga

Otsuov metod je tehnika za automatsko određivanje optimalnog praga pri segmentaciji slike, posebno korisna kod slika sa dvostrukim pikovima u distribuciji intenziteta piksela. Cilj metode je pronaći takav prag koji minimizira varijansu unutar klasa, odnosno razlike u intenzitetima piksela unutar objekta i pozadine. Proces uključuje pretragu svih mogućih vrijednosti praga kako bi se izračunala i odabrala ona vrijednost koja najbolje odvaja piksele u dvije klase. [34]

4.2 Felzenszwalb-Huttenlocher algoritam

Felzenszwalb-Huttenlocher algoritam, poznat i kao Efikasna segmentacija slika zasnovana na grafu, koristi se za segmentaciju slike primjenom teorije grafova, pri čemu se slika predstavlja kao neusmjereni graf. Pikseli slike predstavljaju čvorove u grafu, dok su određeni susjedni pikseli povezani neusmjerenim ivicama, čije težine mjere nesličnost između tih piksela. Algoritam omogućava prilagodljivu segmentaciju koja prilagođava kriterijume na osnovu stepena varijabilnosti u susjednim regionima slike, za razliku od klasičnih metoda segmentacije koje koriste statične pristupe.

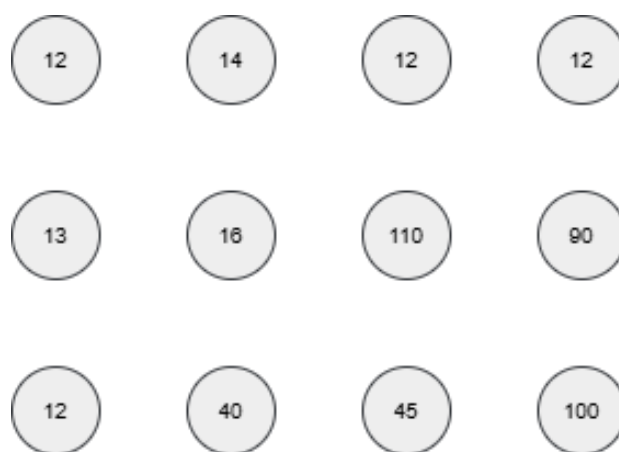
U okviru ovog pristupa, težine ivica izračunavaju se prema razlikama u intenzitetu boje ili svjetlosti između piksela, pri čemu manje razlike ukazuju na sličnije piksele. Nakon što se ivice sortiraju po težini, one se postepeno dodaju u graf, spajajući piksele u segmente na osnovu prilagodljivog praga, koji uzima u obzir unutrašnju varijaciju unutar segmenata i prilagođava se prema njihovim osobinama.

Cilj algoritma je obuhvatiti perceptivno važne regione koji odražavaju globalne karakteristike slike, a radi se na način koji omogućava efikasnu segmentaciju sa vremenskom složenošću od $O(n * \log(n))$. Adaptivni kriterijum za spajanje segmenata definisan je tako da postoji granica između dva susjedna regiona C_i i C_j kada je varijacija između tih regiona veća od varijacije unutar svakog od njih pojedinačno.

Ova strategija osigurava da se segmentacija uskladi s prirodnim granicama objekata u slici, prilagođavajući se globalnim svojstvima koja nisu odmah očigledna.

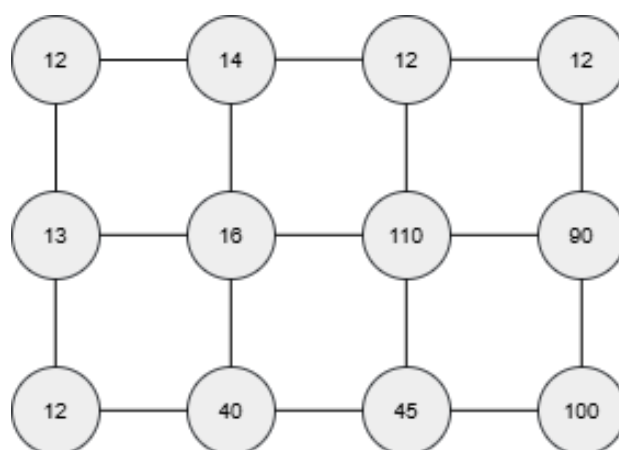
Kao završni korak, segmenti manji od minimalne veličine spajaju se sa susjednim većim segmentima kako bi se uklonile nepravilnosti i postigla uspješna segmentacija. Zbog efikasnosti i skalabilnosti, ovaj algoritam primjenjuje se u različitim aplikacijama računarskog vida, uključujući izdvajanje objekata, medicinsku analizu slika, nadzor i autonomnu vožnju.

U nastavku je dat primjer kako bi se lakše objasnilo funkcionisanje algoritma. Razmatra se slika koja se sastoji od 12 piksela sa intezitetima (slika 25).



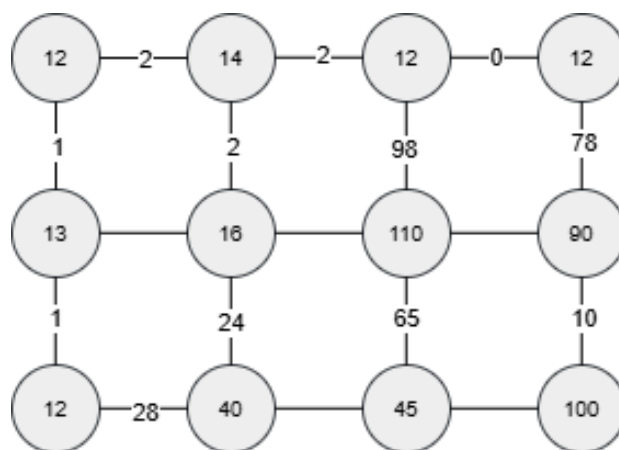
Slika 25. Prikaz piksela slike s intenzitetima

Slika se može predstaviti kao neusmjereni graf koristeći N4 sistem. Prema N4 sistemu, svaki čvor povezuje se sa četiri susjedna čvora (slika 26).



Slika 26. Segmentacija piksela s N4 susjednim povezivanjem

Težine ivica se dodjeljuju na osnovu razlike između intenziteta (slika 27).



Slika 27. Povezivanje piksela sa težinama ivica prema razlikama intenziteta

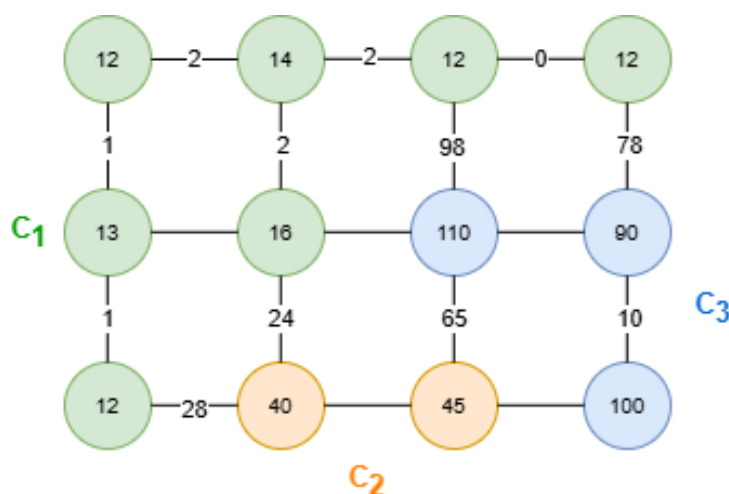
Problem segmentacije može biti formulisano kao particionisanje skupa čvorova V datog neusmjerenog grafa G na komponente $C_1, C_2, C_3, \dots$ tako da:

- ivice između dva čvora unutar istog segmenta C_i imaju niže težine,
- ivice između čvorova u različitim segmentima C_i i C_j imaju više težine.

Prema toj formulaciji, jedno moguće rješenje može biti predstavljeno skupovima C_1, C_2, C_3 pri čemu su:

- C_1 prikazani zelenom bojom,
- C_2 prikazani narandžastom bojom,
- C_3 prikazani plavom bojom.

Čvorovi su segmentirani na osnovu dodijeljenih težina ivica i sličnosti intenziteta piksela, pri čemu je podjela prilagođena tako da odražava razlike između različitih regiona slike (slika 28).



Slika 28. Podjela piksela u segmente

Na slici 29 je prikazan primjer segmentacije koristeći Felzenszwalb-Huttenlocher algoritam.



Slika 29. Originalna i segmentirana slika uz pomoc Felzenszwalb-Huttenlocher algoritma

5. EKSPERIMENTALNI DIO

U okviru eksperimentalnog dijela rada, primijenjena su dva algoritma za grupisanje podataka i dva algoritma za segmentaciju slike kako bi se sveobuhvatno analizirala njihova efikasnost, stabilnost i otpornost na šumove u podacima i slikama. Ovaj dio istraživanja pruža uvid u to kako svaki algoritam odgovara na različite strukture podataka i na promenljive uslove u segmentaciji slika.

Za grupisanje podataka korišćeni su algoritmi Spektralno grupisanje i MST algoritam. Ovi algoritmi su odabrani zbog svojih različitih pristupa i mogućnosti prilagođavanja strukturi podataka. Spektralno grupisanje koristi spektralne karakteristike podataka kako bi identifikovao prirodne grupe, dok MST formira hijerarhijsku strukturu podataka pomoću minimalnog rastućeg stabla. Upoređivanjem ovih algoritama, cilj je da se istraži kako različiti matematički pristupi grupisanju funkcionišu na različitim tipovima podataka. Za evaluaciju performansi algoritama za grupisanje odabrana su tri skupa podataka: Iris, Wine i Breast Cancer. Za ocjenu kvaliteta grupisanja korišćene su sljedeće metrike: Silhouette Score i Calinski-Harabasz Index.

Za segmentaciju slike primijenjeni su algoritmi Random Walks i Efficient Graph-Based Image Segmentation (Felzenszwalb-Huttenlocher algoritam). Svaki algoritam je primijenjen na dva skupa slika: jedan set sadrži 500 standardnih slika različitog sadržaja, dok drugi set obuhvata 300 medicinskih slika. Cilj je bio procijeniti sposobnost algoritama da izvrše preciznu segmentaciju u različitim uslovima i sa različitim vrstama slika. Random Walks algoritam koristi teoriju slučajnog hoda za povezivanje piksela u segmente na osnovu vjerovatnoće, dok Felzenszwalb-Huttenlocher algoritam omogućava brzu i efikasnu segmentaciju, zasnovanu na intenzitetima susjednih piksela kroz grafovsko modelovanje. Evaluacija uspješnosti segmentacije izvršena je pomoću Dice koeficijenta i Jaccardovog indeksa, čime su analizirane preciznost i dosljednost algoritama u izdvajanju relevantnih segmenata iz slika različitog tipa.

Tokom eksperimenta se analizira i otpornost algoritama na šum i buku. Kako bi se dodatno testirala stabilnost i robusnost algoritama, slike su obogaćene Gausovim šumom. Cilj ove analize je da se ispita kako različiti algoritmi za segmentaciju reaguju na promjene u kvalitetu slika i koliko su sposobni da zadrže koherentnost segmentacije i pored prisustva ometajućih faktora.

Ovakav eksperimentalni pristup omogućava cjelovit prikaz performansi svakog algoritma u realnim uslovima, sa naglaskom na njihove prednosti, ograničenja i otpornost na promjene u strukturi i kvalitetu podataka.

5.1 Dataset-ovi

5.1.1 Iris dataset

Iris dataset je jedan od najpoznatijih i najčešće korišćenih skupova podataka u mašinskom učenju, posebno u zadacima klasifikacije i klasteringa. Prvobitno ga je predstavio britanski statističar i biolog Ronald Fišer 1936. godine. Ovaj skup podataka postao je klasičan primjer u statistikama i oblastima mašinskog učenja zbog svoje strukture i jednostavnosti.

Iris dataset sadrži ukupno 150 primjera koji predstavljaju uzorke tri različite vrste (ili klase) cvijeta irisa: Iris Setosa, Iris Versicolor i Iris Virginica. Svaka klasa cvijeta ima po 50 primjera, a svaki primjer je opisan kroz četiri numeričke karakteristike koje su mjerene na cvijetu. Ove karakteristike uključuju dužinu čašičnog lista (Sepal Length) i širinu čašičnog lista (Sepal Width), dužinu laticе (Petal Length) i širinu laticе (Petal Width), sve izražene u centimetrima.

Iris dataset je posebno značajan jer sadrži tri klase koje se mogu prirodno razlikovati po karakteristikama laticе i čašičnih listova. On pruža jasan, ali i jednostavan primjer problema klasifikacije i grupisanja, gde su dvije klase (Iris Setosa i Iris Virginica) lako separabilne, dok su dvije klase (Iris Versicolor i Iris Virginica) djelimično preklapajuće, što ga čini izazovnim za algoritme klasteringa i klasifikacije. Na slici 30 je prikazan kod za implementaciju učitavanja Iris dataset-a u Python-u [35].

```
from sklearn.datasets import load_iris

iris = load_iris()
X = iris.data
y = iris.target

# Prikaz prvih nekoliko redova podataka
print("Karakteristike (features):")
print(X[:5])
print("\nOznake klase (labels):")
print(y[:5])
```

Slika 30. Prikaz koda za učitavanje i pregled Iris podataka

5.1.2 Wine dataset

Wine dataset je skup podataka koji se koristi u mašinskom učenju i statistici za klasifikaciju i klastering. Ovaj skup sadrži podatke o hemijskim karakteristikama tri vrste vina iz Italije, porijeklom iz istočne regije Pijemont. Podaci su korišćeni kako bi se klasifikovale sorte vina na osnovu analize hemijskih svojstava, što je izazovno zbog velikog broja karakteristika i njihove kompleksnosti. Wine dataset se sastoji od 178 uzoraka podijeljenih u tri klase (tri vrste vina), a svaki uzorak je opisan kroz 13 numeričkih karakteristika koje se odnose na različite hemijske osobine vina.

Wine dataset se koristi za testiranje algoritama klasifikacije i klasteringa jer obuhvata visokodimenzionalne podatke (13 karakteristika) i pruža izazov u razdvajanju klasa. Pošto su tri klase vina definisane na osnovu hemijskog sastava, ovaj dataset je pogodan za analize u kojima algoritmi treba da prepoznaju suptilne razlike između grupa. Na slici 31 je prikazan kod za implementaciju učitavanja Wine dataset-a u Python-u [36].

```
from sklearn.datasets import load_wine

wine = load_wine()
X = wine.data
y = wine.target

# Prikaz prvih nekoliko redova podataka
print("Karakteristike (features):")
print(X[:5])
print("\nOznake klasa (labels):")
print(y[:5])
```

Slika 31. Prikaz koda za učitavanje i pregled Wine podataka

5.1.3 Breast Cancer dataset

Breast Cancer dataset je skup podataka koji se koristi za klasifikaciju tumora dojke kao benignih ili malignih na osnovu njihovih fizičkih i hemijskih karakteristika. Ovaj dataset je značajan u medicinskoj analizi, jer omogućava procjenu algoritama za detekciju tumora, i daje mogućnost za testiranje algoritama u binarnim klasifikacionim zadacima.

Breast Cancer dataset sadrži:

- **569 uzoraka:** Svaki uzorak predstavlja jedan tumor.
- **2 klase:** Benigni i maligni tumori.
- **30 karakteristika:** Svaki uzorak je opisan kroz 30 numeričkih karakteristika kao što su veličina tumora, tekstura, oblik, simetrija itd.

Ove karakteristike su ključne za otkrivanje razlika između benignih i malignih tumora i često se koriste u algoritmima za medicinsku klasifikaciju i klastering. Na slici 33 je prikazan kod za implementaciju učitavanja Breast Cancer dataset-a u Python-u [37].

```
from sklearn.datasets import load_breast_cancer

breast_cancer = load_breast_cancer()
X = breast_cancer.data
y = breast_cancer.target

# Prikaz prvih nekoliko redova podataka
print("Karakteristike (features):")
print(X[:5])
print("\nOznake klase (labels):")
print(y[:5])
```

Slika 32. Prikaz koda za učitavanje i pregled Breast Cancer podataka

5.1.4 BSDS500 (Berkeley Segmentation Dataset 500)

BSDS500 (Berkeley Segmentation Dataset 500) se sastoji od 500 slika različitih prirodnih scena, koje se koriste za evaluaciju algoritama za segmentaciju slike. Svaka slika u ovom dataset-u ima pripadajuće ground-truth oznake, koje predstavljaju referentne segmente ručno kreirane od strane ljudi. Ground-truth označava pravu segmentaciju slike koju algoritmi pokušavaju da predvide i koristi se kao standard za poređenje tačnosti predikcija segmentacije.

5.1.5 Datasetovi sa medicinskim slikama

Dataset sa medicinskim slikama sastoji se od ukupno 300 slika, koje obuhvataju različite anatomske strukture, poput pluća, mozga, ćelija i slično. Ovaj dataset je kreiran spajanjem slika iz nekoliko poznatih medicinskih dataset-ova kako bi se omogućilo sveobuhvatno istraživanje segmentacije na različitim tipovima medicinskih slika. Ova raznovrsnost omogućava procjenu performansi algoritama segmentacije u različitim medicinskim kontekstima.

5.2 Metrike/metode

5.2.1 Silhouette Score

Silhouette Score je alat koji se koristi za procjenu prikladnosti rezultata grupisanja, pružajući kvantitativnu mjeru o tome koliko su klasteri dobro definisani i različiti. Silhouette Score kvantifikuje koliko se dobro tačka uklapa u dodijeljeni klaster i koliko se razlikuje od drugih klastera. Ova mjera se zasniva na koheziji i separaciji podataka unutar klastera i koristi se za utvrđivanje da li su klasteri jasno odvojeni i unutrašnje homogeni.

Silhouette Score predstavlja metriku za procjenu performansi grupisanja. Procjena kvaliteta grupisanja je ključna za određivanje efikasnosti i pouzdanosti algoritama grupisanja.

Za računanje Silhouette Score za skup podataka potrebno je za svaku tačku i izračunati sljedeće vrijednosti:

- a_i - Prosječna udaljenost tačke i do svih drugih tačaka unutar istog klastera (intra-klasterska udaljenost):

$$a_i = \frac{1}{|C| - 1} \sum_{j \in C, j \neq i} d(i, j),$$

gdje je $d(i, j)$ udaljenost između tačaka i i j , a $|C|$ broj tačaka u klasteru C ;

- b_i - Prosječna udaljenost tačke i do svih tačaka u najbližem klasteru (inter-klasterska udaljenost):

$$b_i = \min_{D \neq C} \frac{1}{|D|} \sum_{j \in D} d(i, j),$$

gdje je $|D|$ broj tačaka u klasteru D , a $d(i, j)$ udaljenost između tačaka i i j .

Vrijednosti a_i i b_i se koriste za izračunavanje SC za svaku tačku i pomoću formule:

$$SC(i) = \frac{b_i - a_i}{\max(a_i, b_i)}$$

Ukupni SC za rezultat grupisanja dobija se tako što se izračuna SC za sve tačke a zatim izračuna njihov prosjek:

$$SC = \frac{\sum_{i=1}^n SC(i)}{n}$$

Silhouette Score može imati vrijednosti između -1 i 1, gde vrijednosti bliže 1 označavaju dobro definisane i jasno odvojene klasterne, dok vrijednosti bliže 0 ili negativne ukazuju na preklapanje ili pogrešnu dodjelu klastera [38].

5.2.2 Calinski-Harabasz Index

Calinski-Harabasz Index, poznat i kao Variance Ratio Criterion, koristi se za ocjenu kvaliteta grupisanja na osnovu varijanse između klastera i unutar klastera. Ovaj indeks mjeri koliko su klasteri međusobno udaljeni i koliko su unutar sebe kompaktni. Što su grupe međusobno dalje i unutar sebe zbijenije, to će vrijednost Calinski-Harabasz Index-a biti veća, što ukazuje na bolje definisano grupisanje.

Za skup podataka sa n uzoraka podeljenih u k klastera, Calinski-Harabasz Index se definiše kao:

$$CH = \frac{S_b/(k-1)}{S_w/(n-k)},$$

gdje:

- S_b označava varijansu između klastera, odnosno razliku između centra svakog klastera i ukupnog centra podataka,
- S_w označava varijansu unutar klastera, odnosno razliku između tačaka unutar svakog klastera i centra tog klastera

Varijansa između klastera se izračunava kao:

$$S_b = \sum_{j=1}^b |C_j| * \|\mu_j - \mu\|^2$$

gdje:

$|C_j|$ označava broj tačaka u klasteru j ,

μ_j označava centar klastera j ,

μ označava ukupan centar svih podataka

$\|\mu_j - \mu\|^2$ označava kvadrat euklidske udaljenosti između centra klastera j i ukupnog centra

Varijansa unutar klastera se izračunava kao:

$$S_w = \sum_{j=1}^k \sum_{x \in C_j} \|x - \mu_j\|^2$$

gdje:

x označava tačku unutar klastera j .

$\|x - \mu_j\|^2$ označava kvadrat euklidske udaljenosti između tačke x i centra njenog klastera

μ_j

Veća vrijednost Calinski-Harabasz Index-a ukazuje na bolje definisane i jasnije razdvojene klastere [39].

5.2.3 Dice koeficijent

Dice koeficijent, F-mjera, takođe poznata kao F-skor, je jedna od najrasprostranjenijih metrika za mjerenje performansi u kompjuterskom vidu i u MIS-u (Medicinska Segmentacija Slika).

Dice koeficijent se izračunava na osnovu preciznosti i odziva predikcije. Ocjenjuje preklapanje između predviđene segmentacije i ground-truth segmentacije. Takođe, penalizuje lažno pozitivne slučajeve, što je čest faktor kod dataset-ova sa velikim neravnotežama u klasama.

Dice koeficijent predstavlja harmonijsku sredinu preciznosti i odziva. Formula za Dice koeficijent je:

$$Dice = \frac{2TP}{2TP + FP + FN}$$

gdje:

TP predstavlja broj tačno pozitivnih piksela (true positives).

FP predstavlja broj lažno pozitivnih piksela (false positives).

FN predstavlja broj lažno negativnih piksela (false negatives).

Na slici 33 je prikazana ilustracija Dice koeficijenta [40].



Slika 33. Ilustracija Dice koeficijenta

5.2.4 Jaccard-ov indeks

Jaccard-ov indeks koristi se kao mjera sličnosti između dva asimetrična binarna vektora ili za procjenu sličnosti između dva skupa, te je često primjenjivana u segmentaciji slike za računanje uspješnosti segmentacije. Smatra se uobičajenom mjerom bliskosti za izračunavanje sličnosti između dva objekta, kao što su dva tekstualna dokumenta ili segmenti slike. Indeks se kreće u rasponu od 0 do 1, pri čemu vrijednosti bliže 1 označavaju veću sličnost između dva skupa podataka.

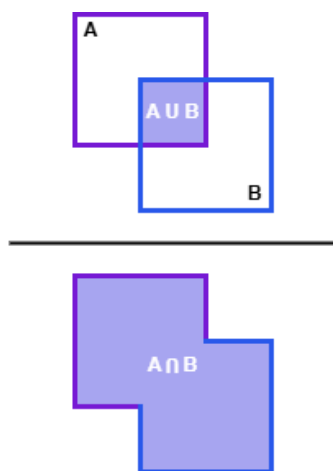
Izračunava se prema formuli:

$$J = \frac{\text{broj elemenata u oba skupa}}{\text{broj elemenata u bilo kojem skupu}}$$

ili matematički:

$$J(A, B) = \frac{|A \cap B|}{|A \cup B|}$$

Ako dva skupa podataka sadrže iste elemente, njihov Jaccard-ov indeks sličnosti će iznositi 1, dok će indeks biti 0 ako nemaju zajedničkih elemenata. Ovom mjerom sličnosti može se procijeniti koliko su osobine u dataset-u međusobno slične, što je korisno za evaluaciju kvaliteta segmentacije. Na slici 34 je prikazana ilustracija metrike [41].

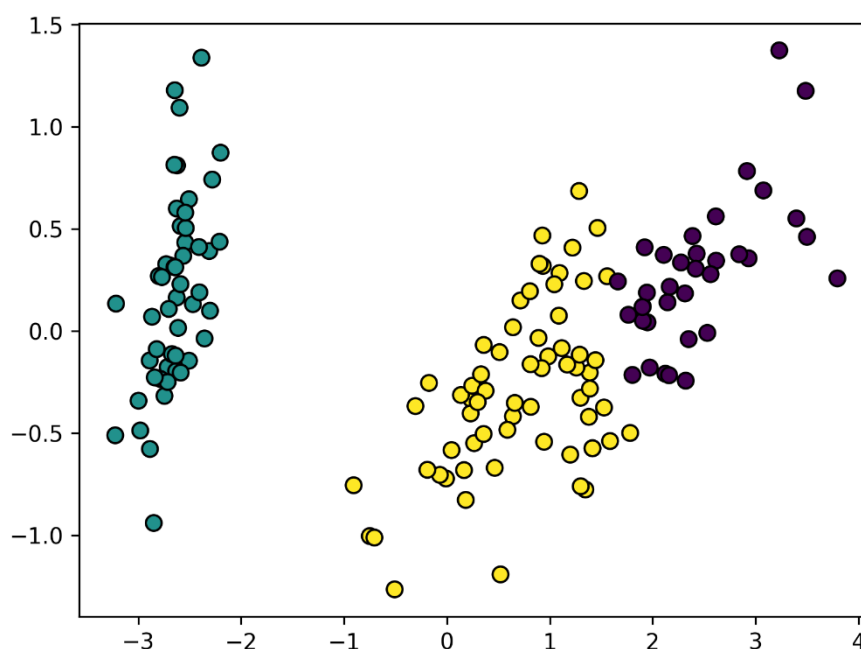


Slika 34. Ilustracija Jaccard-ovog indeksa

5.3 Spektralno grupisanje

5.3.1 Iris dataset

U nastavku je prikazana analiza eksperimentalnih rezultata postignutih primjenom Spektralnog grupisanja na Iris skupu podataka. Eksperimenti su izvedeni na originalnom skupu podataka, kao i na skupu podataka sa dodatkom Gausove buke, kako bi se procijenio uticaj šuma na kvalitet grupisanja.



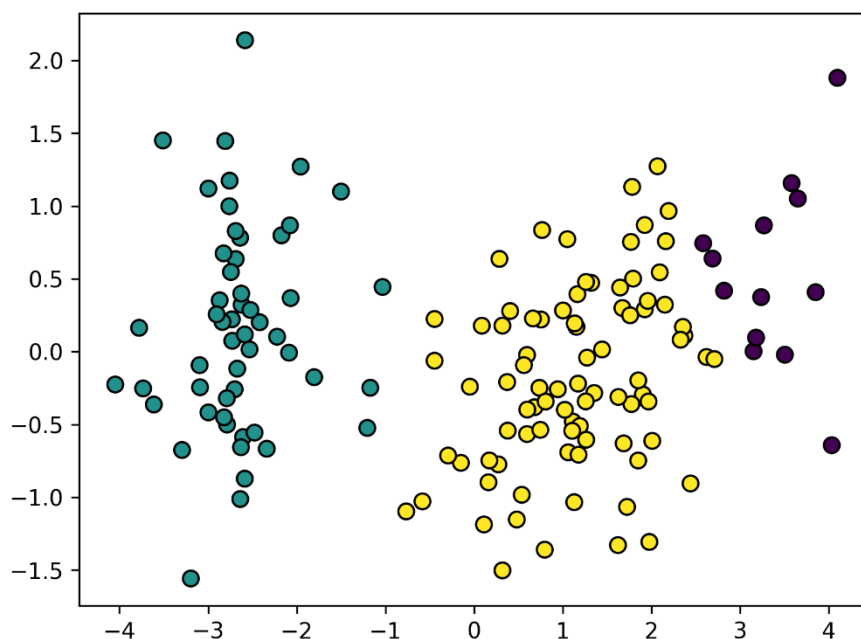
Slika 35. Spektralno grupisanje na Iris skupu podataka

Tabela 5. Evaluacija kvaliteta grupisanja Iris skupa podataka primjenom Spektralnog grupisanja

Metode	
Silhouette Score	0.56
Calinski-Harabasz Index	556.12

Na slici 35 prikazano je grupisanje na originalnom Iris dataset-u bez buke. Rezultati grupisanja evaluirani su pomoću metrika Silhouette Score i Calinski-Harabasz Index. Vrijednosti

ovih metrika, koje iznose 0.56 za Silhouette Score i 556.12 za Calinski-Harabasz Index, ukazuju na relativno dobro definisane klastere sa jasnom separacijom među grupama (tabela 5).



Slika 36. Spektralno grupisanje na Iris skupu podataka sa bukom

Tabela 6. Evaluacija kvaliteta grupisanja Iris skupa podataka primjenom Spektralnog grupisanja sa bukom

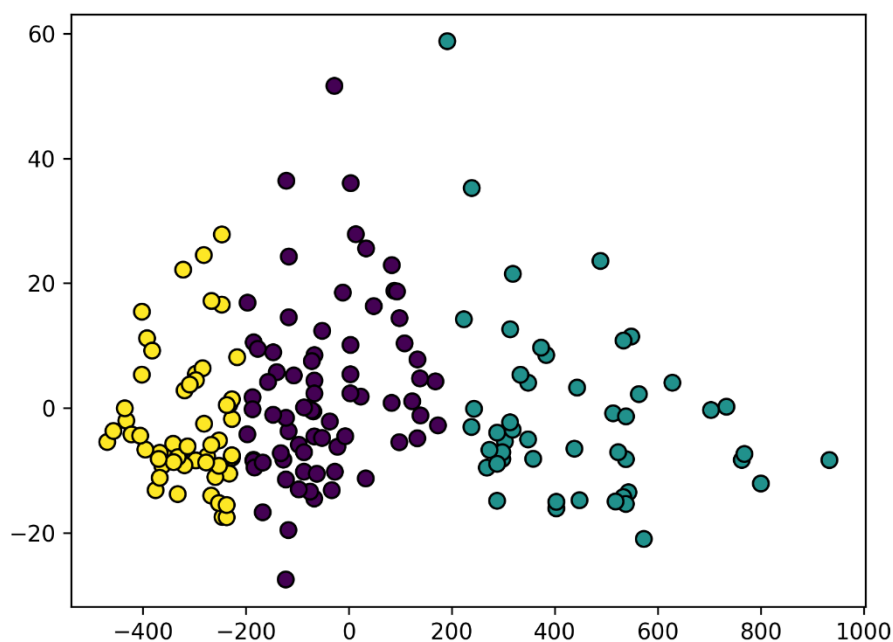
Metode	
Silhouette Score	0.44
Calinski-Harabasz Index	210.04

Na slici 36 prikazano je grupisanje na Iris dataset-u sa dodatkom Gausove buke. Dodavanjem šuma, vrijednosti metrika su se smanjile na 0.44 za Silhouette Score i 210.04 za Calinski-Harabasz Index (tabela 6). Smanjenje ovih vrijednosti ukazuje na slabiji kvalitet grupisanja, jer je struktura klastera postala manje jasna, što je posljedica povećane varijabilnosti unutar klastera i smanjene separacije između klastera.

Ove slike pružaju uvid u otpornost Spektralnog grupisanja na šum, pokazujući da prisustvo buke može značajno uticati na rezultate grupisanja, smanjujući koherenciju i razdvojenost klastera.

5.3.2 Wine dataset

U nastavku je prikazana analiza eksperimentalnih rezultata postignutih primjenom Spektralnog grupisanja na Wine skupu podataka, sa različitim intenzitetima dodate buke kako bi se procijenio uticaj šuma na kvalitet grupisanja.

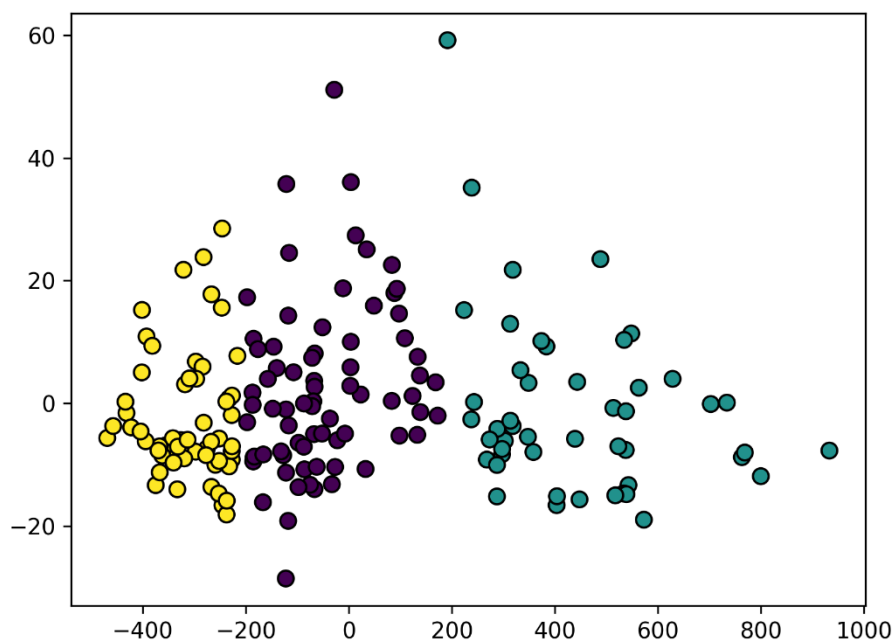


Slika 37. Spektralno grupisanje na Wine skupu podataka

Tabela 7. Evaluacija kvaliteta grupisanja Wine skupa podataka primjenom Spektralnog grupisanja

Metode	
Silhouette Score	0.56
Calinski-Harabasz Index	533.86

Na slici 37 prikazano je grupisanje na originalnom Wine dataset-u bez dodate buke. Rezultati grupisanja evaluirani su pomoću metrika Silhouette Score i Calinski-Harabasz Index. Vrijednosti ovih metrika, koje iznose 0.56 za Silhouette Score i 533.86 za Calinski-Harabasz Index, ukazuju na dobro definisane klastere sa jasno uočljivom separacijom između grupa (tabela 7).



Slika 38. Spektralno grupisanje na Wine skupu podataka sa bukom

Tabela 8. Evaluacija kvaliteta grupisanja Wine skupa podataka primjenom Spektralnog grupisanja sa bukom

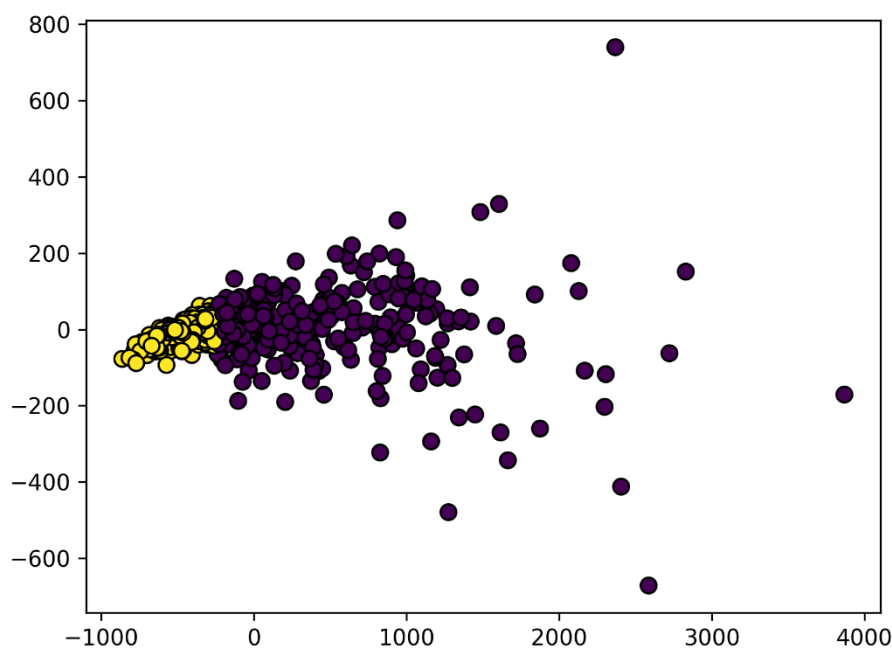
Metode	
Silhouette Score	0.56
Calinski-Harabasz Index	533.86

Na slici 38 prikazano je grupisanje na Wine dataset-u sa dodatnom Gausovom bukom gdje je intenzitet buke postavljen na 0.5. Iako je buka blago narušila strukturu podataka, vrijednosti metrika su ostale iste kao u slučaju bez buke, sa Silhouette Score od 0.56 i Calinski-Harabasz Index od 533.86. Ovo pokazuje da dodavanje blagog šuma nije imalo značajan uticaj na kvalitet grupisanja, a struktura klastera je ostala stabilna (tabela 8).

Važno je napomenuti da otpornost Spektralnog grupisanja na niže nivoe šuma može biti rezultat specifične strukture Wine skupa podataka, koja je možda manje podložna narušavanju u prisustvu umjerenog šuma. Ovi rezultati ukazuju na to da se Spektralno grupisanje pokazalo efikasnim u prisustvu manjeg šuma, što upućuje na njegovu otpornost, kao i na potencijalnu stabilnost samog skupa podataka.

5.3.3 Breast Cancer dataset

U nastavku je prikazana analiza eksperimentalnih rezultata postignutih primjenom Spektralnog grupisanja na Breast Cancer skupu podataka, sa i bez dodatka buke, kako bi se procijenio uticaj šuma na kvalitet grupisanja.

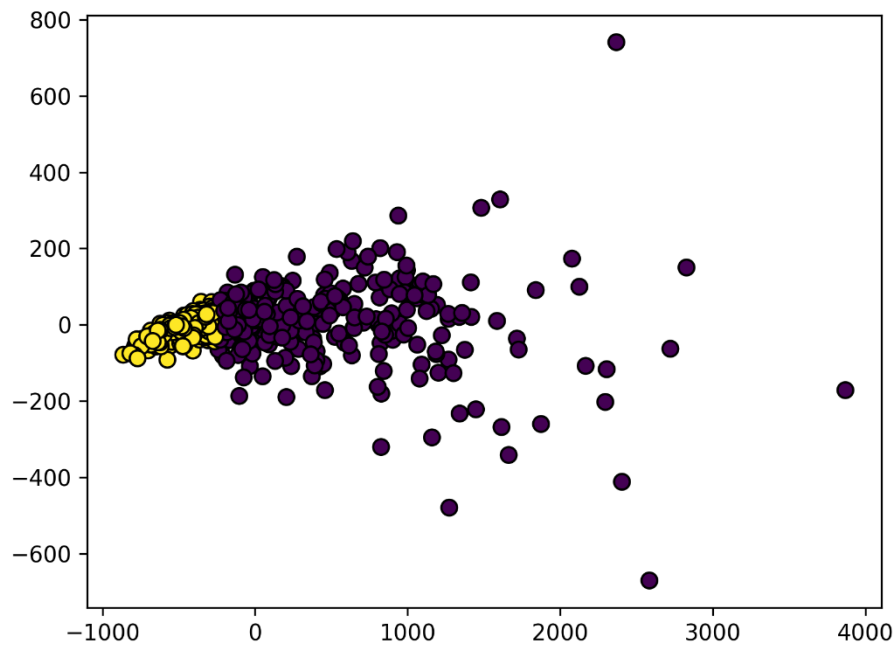


Slika 39. Spektralno grupisanje na Breast Cancer skupu podataka

Tabela 9. Evaluacija kvaliteta grupisanja Breast Cancer skupa podataka primjenom Spektralnog grupisanja

Metode	
Silhouette Score	0.41
Calinski-Harabasz Index	445.02

Na slici 39 prikazano je grupisanje na originalnom Breast Cancer dataset-u bez buke. Rezultati grupisanja evaluirani su pomoću metrika Silhouette Score i Calinski-Harabasz Index. Vrijednosti ovih metrika, koje iznose 0.41 za Silhouette Score i 445.02 za Calinski-Harabasz Index, ukazuju na umjereno definisane klastere, s obzirom na to da vrijednosti nijesu ekstremno visoke, ali ukazuju na određeni nivo separacije (tabela 9).



Slika 40. Spektralno grupisanje na Breast Cancer skupu podataka sa bukom

Tabela 10. Evaluacija kvaliteta grupisanja Breast Cancer skupa podataka primjenom Spektralno grupisanje algoritma sa bukom

Metode	
Silhouette Score	0.41
Calinski-Harabasz Index	442

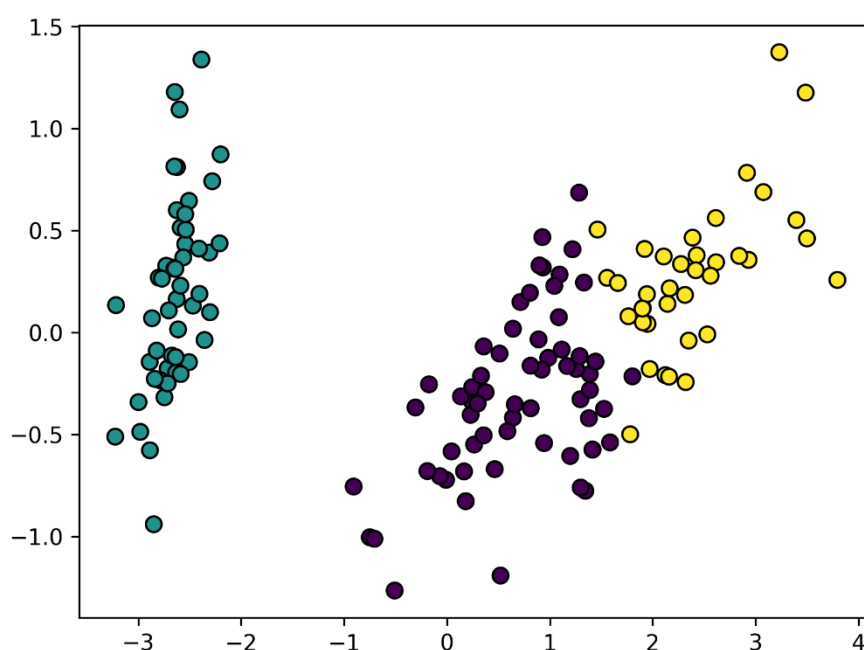
Na slici 40 prikazano je grupisanje na Breast Cancer dataset-u sa dodanom Gausovom bukom. Uprkos dodatku buke, vrijednosti metrika su se vrlo malo promijenile, sa Silhouette Score ostajući na 0.41 i Calinski-Harabasz Index koji je blago opao na 442. Ovaj minimalan pad sugerise da blagi dodaci šuma nisu značajno uticali na strukturu klastera (tabela 10).

Ovi rezultati ukazuju na otpornost Spektralno grupisanje algoritma prema niskom nivou šuma, što može biti posljedica same prirode Breast Cancer skupa podataka, gdje su klasteri možda čvršće strukturirani i manje podložni narušavanju u prisustvu umjerenog šuma.

5.4 MST algoritam

5.4.1 Iris dataset

U nastavku je prikazana analiza eksperimentalnih rezultata postignutih primjenom MST algoritma za grupisanje na Iris skupu podataka, sa i bez dodatka buke, kako bi se procijenio uticaj šuma na kvalitet grupisanja.

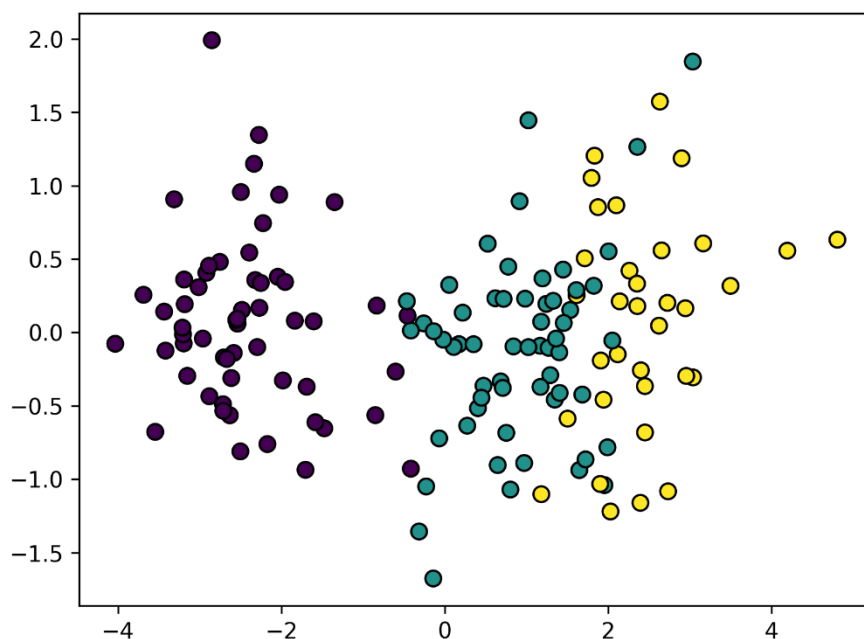


Slika 41. MST grupisanje na Iris skupu podataka

Tabela 11. Evaluacija kvaliteta grupisanja Iris skupa podataka primjenom MST algoritma

Metode	
Silhouette Score	0.55
Calinski-Harabasz Index	556.14

Na slici 41 prikazano je grupisanje na originalnom Iris dataset-u bez buke. Rezultati grupisanja evaluirani su pomoću metrika Silhouette Score i Calinski-Harabasz Index. Vrijednosti ovih metrika, koje iznose 0.55 za Silhouette Score i 556.14 za Calinski-Harabasz Index, ukazuju na relativno dobro definisane klastere sa jasnom separacijom među grupama (tabela 11).



Slika 42. MST grupisanje na Iris skupu podataka sa bukom

Tabela 12. Evaluacija kvaliteta grupisanja Iris skupa podataka primjenom MST algoritma sa bukom

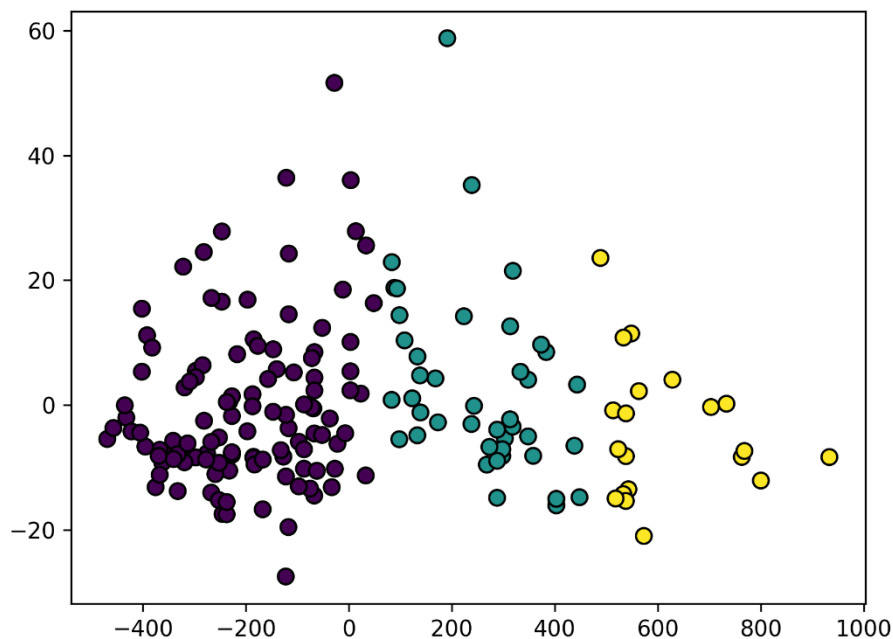
Metode	
Silhouette Score	0.33
Calinski-Harabasz Index	182.01

Na slici 42 prikazano je grupisanje na Iris dataset-u sa dodatkom Gausove buke. Nakon dodavanja šuma, vrijednosti metrika su se značajno smanjile, sa Silhouette Score na 0.33 i Calinski-Harabasz Index na 182.01. Ovaj značajan pad ukazuje na smanjenje kvaliteta grupisanja, jer je šum u velikoj mjeri uticao na koherentnost klastera i njihovu međusobnu separaciju (tabela 12).

Ovi rezultati pokazuju da MST algoritam za grupisanje nije toliko otporan na buku u Iris dataset-u, pri čemu čak i umjereni nivoi šuma značajno narušavaju strukturu klastera.

5.4.2 Wine dataset

U nastavku je prikazana analiza eksperimentalnih rezultata postignutih primjenom MST algoritma za grupisanje na Wine skupu podataka, sa i bez dodatka buke, kako bi se procijenio uticaj šuma na kvalitet grupisanja.

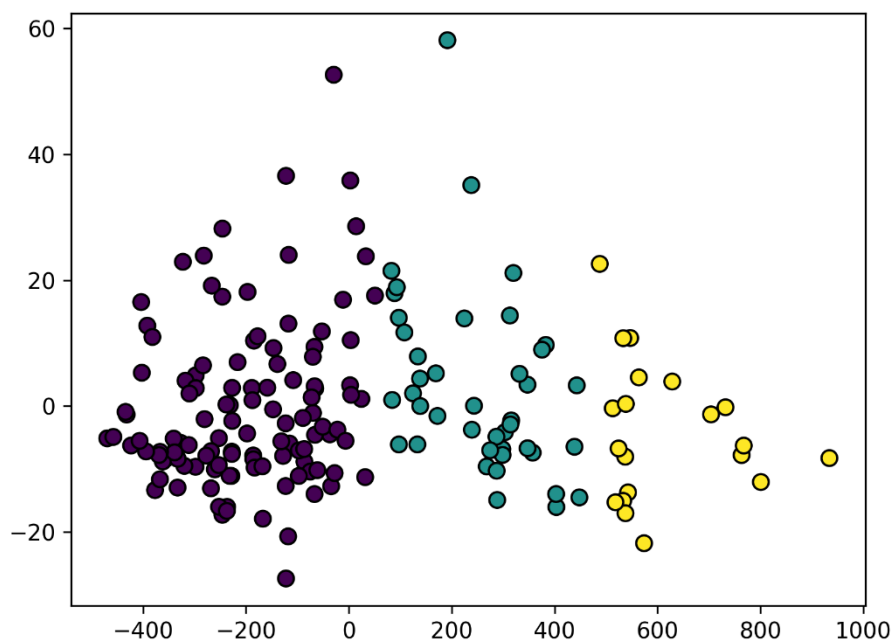


Slika 43. MST grupisanje na Wine skupu podataka

Tabela 13. Evaluacija kvaliteta grupisanja Wine skupa podataka primjenom MST algoritma

Metode	
Silhouette Score	0.60
Calinski-Harabasz Index	457.31

Na slici 43 prikazano je grupisanje na originalnom Wine dataset-u bez dodatka buke. Rezultati grupisanja evaluirani su pomoću metrika Silhouette Score i Calinski-Harabasz Index. Vrijednosti ovih metrika, koje iznose 0.60 za Silhouette Score i 457.31 za Calinski-Harabasz Index, ukazuju na dobro definisane klastere sa jasnom separacijom među grupama (tabela 13).



Slika 44. MST grupisanje na Wine skupu podataka sa bukom

Tabela 14. Evaluacija kvaliteta grupisanja Wine skupa podataka primjenom MST algoritma sa bukom

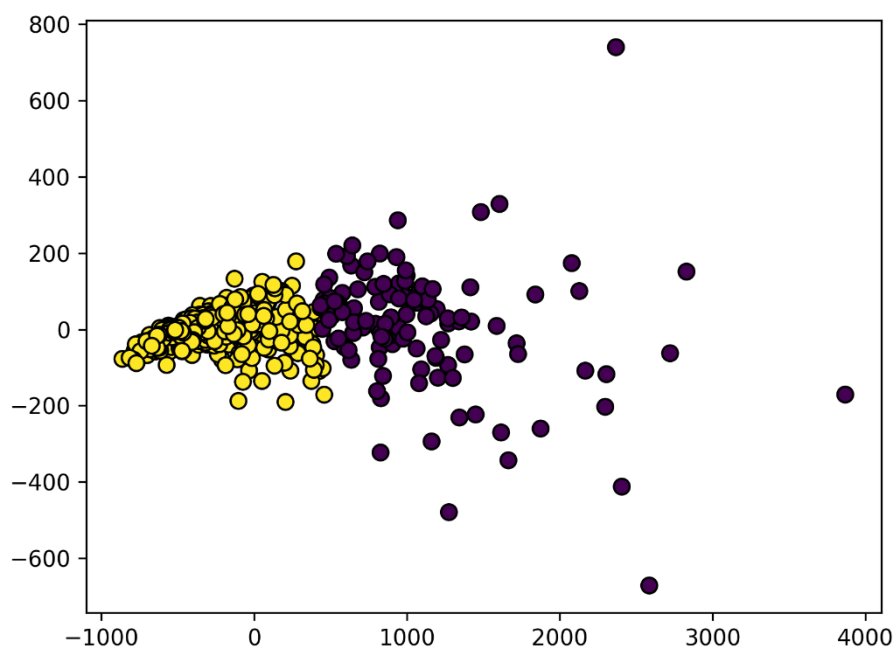
Metode	
Silhouette Score	0.60
Calinski-Harabasz Index	457.17

Na slici 44 prikazano je grupisanje na Wine dataset-u sa dodatkom Gausove buke. Nakon dodavanja šuma, vrijednosti metrika ostale su gotovo identične, sa Silhouette Score od 0.60 i Calinski-Harabasz Index blago opao na 457.17. Ove vrijednosti pokazuju da dodatak šuma nije imao značajan uticaj na strukturu klastera, što ukazuje na otpornost MST algoritma u ovom slučaju (tabela 14).

Ovi rezultati sugerišu da je MST algoritam za grupisanje na Wine dataset-u otporan na umjerene nivoe šuma, vjerovatno zbog stabilne strukture ovog skupa podataka, što je omogućilo algoritmu da održi visoku koherentnost klastera čak i u prisustvu buke.

5.4.3 Breast Cancer dataset

U nastavku je prikazana analiza eksperimentalnih rezultata postignutih primjenom MST algoritma za grupisanje na Breast Cancer skupu podataka, sa i bez dodatka buke, kako bi se procijenio uticaj šuma na kvalitet grupisanja.

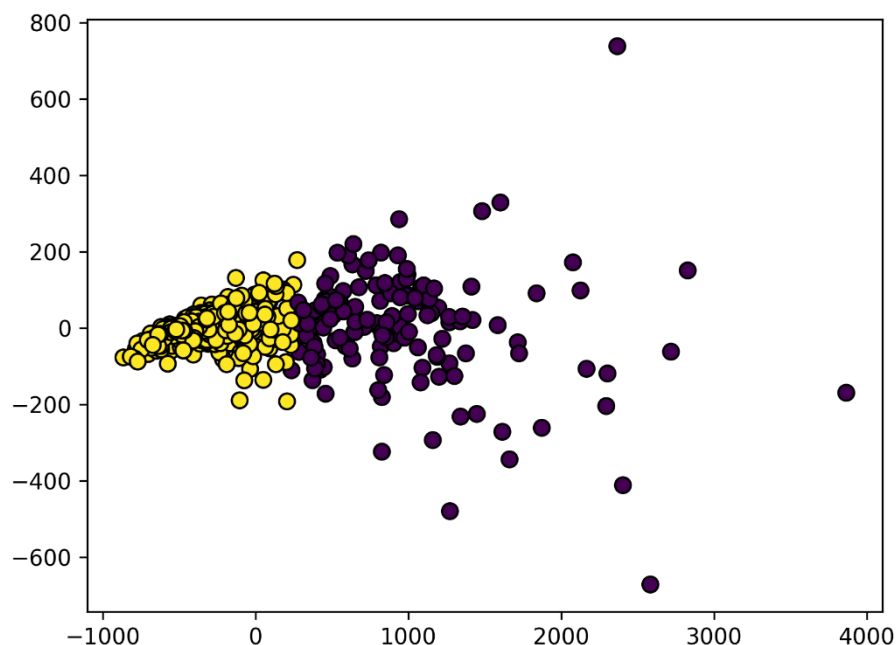


Slika 45. MST grupisanje na Breast Cancer skupu podataka

Tabela 15. Evaluacija kvaliteta grupisanja Breast Cancer skupa podataka primjenom MST algoritma

Metode	
Silhouette Score	0.60
Calinski-Harabasz Index	1282.02

Na slici 45 prikazano je grupisanje na originalnom Breast Cancer dataset-u bez buke. Rezultati grupisanja evaluirani su pomoću metrika Silhouette Score i Calinski-Harabasz Index. Vrijednosti ovih metrika, koje iznose 0.70 za Silhouette Score i 1282.02 za Calinski-Harabasz Index, ukazuju na dobro definisane klastere sa jasnom separacijom među grupama (tabela 15).



Slika 46. MST grupisanje na Breast Cancer skupu podataka sa bukom

Tabela 16. Evaluacija kvaliteta grupisanja Breast Cancer skupa podataka primjenom MST algoritma sa bukom

Metode	
Silhouette Score	0.69
Calinski-Harabasz Index	1276.12

Na slici 46 prikazano je grupisanje na Breast Cancer dataset-u sa dodatkom Gausove buke. Nakon dodavanja šuma, vrijednosti metrika su se blago smanjile, sa Silhouette Score padajući na 0.69 i Calinski-Harabasz Index na 1276.12. Ovaj minimalan pad vrijednosti ukazuje na to da dodatak buke nije značajno uticao na strukturu klastera, što potvrđuje otpornost MST algoritma u ovom slučaju (tabela 16).

Ovi rezultati pokazuju da MST algoritam za grupisanje na Breast Cancer skupu podataka zadržava stabilnost i visok kvalitet klastera čak i u prisustvu umjerene količine šuma, što može biti posljedica stabilne strukture samog skupa podataka.

5.5 Random Walks algoritam

5.5.1 BSDS500 (Berkeley Segmentation Dataset 500)

Na slici 47 prikazani su pojedinačni primjeri rezultata segmentacije dobijeni primjenom Random Walks algoritma na dataset od 500 slika bez dodatka šuma ili buke. Ovaj algoritam koristi pristup zasnovan na vjerojatnoći, gdje se segmentacija ostvaruje širenjem markera kroz sliku, što omogućava prepoznavanje objekata na osnovu njihovog intenziteta i povezivanja piksel po piksel. Prikazani primjeri ilustruju sposobnost algoritma da odvoji glavne objekte od pozadine, pokazujući efikasnost metode u jasnom odvajanju objekata. Na desnoj strani za svaku sliku prikazana je segmentacija koja obuhvata prepoznate oblike i konture, što je od ključnog značaja za zadatke u kojima se zahtijeva precizno odvajanje objekata od pozadine.

Važno je napomenuti da rezultati prikazani na ovim slikama predstavljaju samo nekoliko primjera i nisu reprezentativni za čitav dataset. Performanse algoritma mogu varirati u zavisnosti od složenosti slike, osvjetljenja, kontrasta i drugih faktora koji mogu uticati na preciznost segmentacije.

U tabeli 17 prikazane su prosječne vrijednosti evaluacionih metrika za cijeli dataset. Dice koeficijent iznosi 0.71, što ukazuje na dobro preklapanje između segmentiranih regiona i ground-truth oznaka. Jaccard-ov indeks, sa vrijednošću od 0.57, potvrđuje ovaj rezultat i ukazuje da su postignute zadovoljavajuće performanse algoritma na ovom datasetu u uslovima bez dodatnog šuma. Ove vrijednosti sugerišu da algoritam može pouzdano segmentisati većinu slika u datasetu, uz mogućnost varijacija u zavisnosti od specifičnosti svake slike.

Tabela 17. Evaluacija kvaliteta segmentacije primjenom Random Walks algoritma na dataset-u

Metode	
Dice koeficijent	0.71
Jaccard-ov indeks	0.57



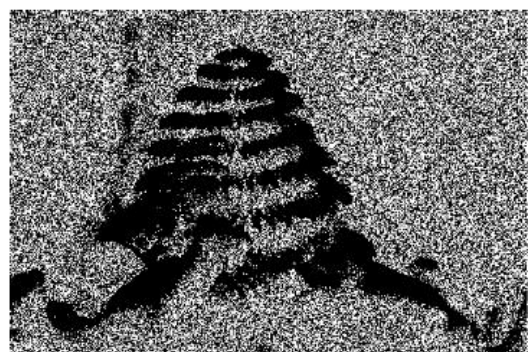
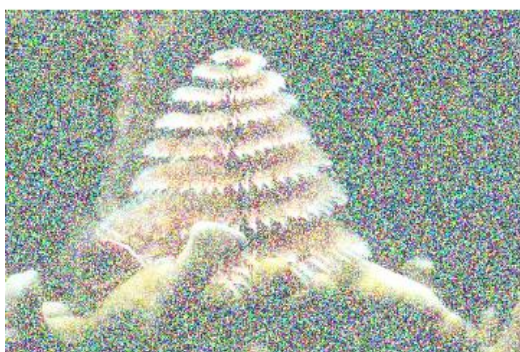
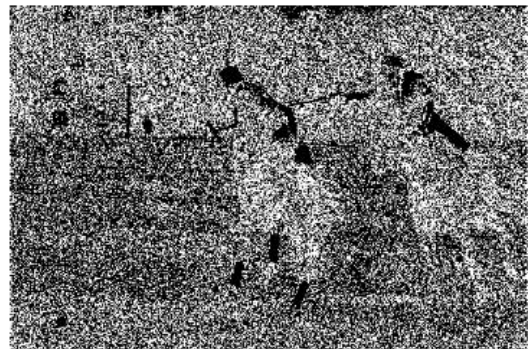
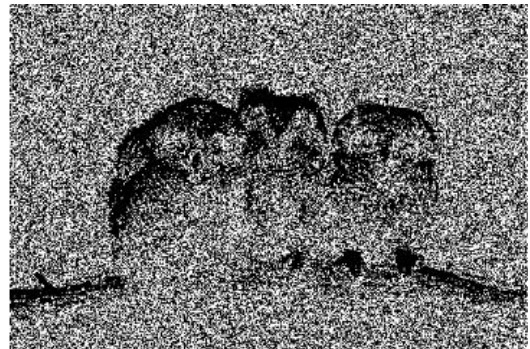
Slika 47. Prikaz originalnih slika i segmentiranih verzija korišćenjem Random Walks algoritma na dataset-u

Na slici 48 prikazani su pojedinačni primjeri rezultata segmentacije dobijeni primjenom Random Walks algoritma na isti dataset od 500 slika, ali sa dodatkom šuma. Gausovski šum je dodat sa varijansom od 0.01 kako bi se testirala robusnost algoritma u uslovima narušenog kvaliteta slike. Prikazani primjeri ilustruju uticaj šuma na preciznost segmentacije, pri čemu se jasno uočava kako dodatni šum otežava prepoznavanje objekata i njihovo odvajanje od pozadine.

U tabeli 18 prikazane su prosječne vrijednosti evaluacionih metrika za cijeli dataset sa dodatkom šuma. Dice koeficijent iznosi 0.53, što ukazuje na smanjenje preklapanja između segmentiranih regiona i ground-truth oznaka u poređenju sa originalnim dataset-om bez šuma. Jaccard-ov indeks, sa vrijednošću od 0.36, potvrđuje smanjenu preciznost algoritma u prisustvu šuma. Primjetno smanjenje ovih vrijednosti sugerise da dodatak šuma značajno otežava tačnost segmentacije i negativno utiče na performanse algoritma i smanjuje njegovu sposobnost.

Tabela 18. Evaluacija kvaliteta segmentacije primjenom Random Walks algoritma na dataset-u sa bukom

Metode	
Dice koeficijent	0.53
Jaccard-ov indeks	0.36



Slika 48. Prikaz originalnih slika i segmentiranih verzija korišćenjem Random Walks algoritma na dataset-u sa bukom

5.5.2 Dataset sa medicinskim slikama

Na slici 49 prikazani su pojedinačni primjeri rezultata segmentacije medicinskih slika dobijeni primjenom Random Walks algoritma na dataset koji se sastoji od 300 medicinskih slika. Ovaj dataset obuhvata razne tipove slika, uključujući snimke oka, mozga i pluća, što omogućava sveobuhvatnu procjenu algoritma u različitim medicinskim kontekstima. Random Walks algoritam koristi pristup zasnovan na vjerojatnoći, pri čemu se segmentacija ostvaruje širenjem markera kroz sliku, što omogućava prepoznavanje struktura na osnovu intenziteta i povezanosti piksel po piksel.

Prikazani primjeri na slici ilustriraju sposobnost algoritma da identifikuje i odvoji ključne anatomske strukture, pokazujući efikasnost metode u preciznom odvajanju objekata. Na desnoj strani za svaku sliku prikazana je segmentacija koja obuhvata prepoznate anatomske oblike, što je od suštinskog značaja za zadatke u medicinskoj dijagnostici i analizi.

U tabeli 19 prikazane su prosječne vrijednosti evaluacionih metrika za cijeli medicinski dataset. Dice koeficijent iznosi 0.59, što ukazuje na solidno preklapanje između segmentiranih regiona i ground-truth oznaka, ali nešto niže u poređenju sa nemedicinskim slikama. Jaccard-ov indeks, sa vrijednošću od 0.45, takođe pokazuje zadovoljavajuću, ali umjereniju preciznost segmentacije u ovom specifičnom kontekstu. Ovi rezultati sugerišu da algoritam može pouzdano segmentisati medicinske slike, uz mogućnost poboljšanja u zavisnosti od specifičnih karakteristika svake vrste slike.

Tabela 19. Evaluacija kvaliteta segmentacije primjenom Random Walks algoritma na dataset-u sa medicinskim slikama

Metode	
Dice koeficijent	0.59
Jaccard-ov indeks	0.45



Slika 49. Prikaz originalnih slika i segmentiranih verzija korišćenjem Random Walks algoritma na dataset-u

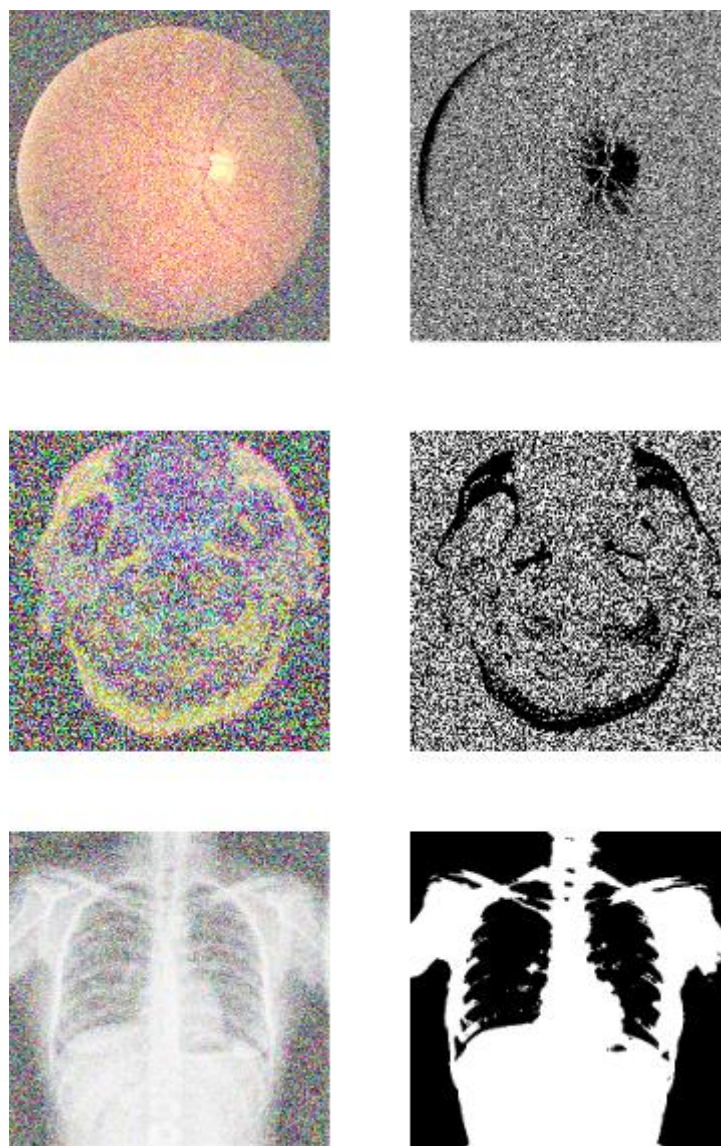
Na slici 50 prikazani su pojedinačni primjeri rezultata segmentacije medicinskih slika dobijeni primjenom Random Walks algoritma na dataset sa dodatkom gausovskog šuma (varijansa šuma: 0.01). Cilj dodavanja šuma je procjena robusnosti algoritma u uslovima narušenog kvaliteta slike, kako bi se utvrdilo koliko dobro algoritam može identifikovati i odvojiti strukture u prisustvu dodatnih smetnji. Prikazani primjeri ilustruju kako šum otežava prepoznavanje finih detalja, smanjujući jasnoću granica između anatomske strukture i pozadine.

U tabeli 20 prikazane su prosječne vrijednosti evaluacionih metrika za cijeli dataset sa dodatkom šuma. Dice koeficijent iznosi 0.56, što pokazuje smanjenje preklapanja između segmentiranih regiona i ground-truth oznaka u poređenju sa slikama bez šuma. Jaccard-ov indeks, sa vrijednošću od 0.40, dodatno potvrđuje nižu preciznost segmentacije u ovim uslovima. Ovi

rezultati sugerišu da šum negativno utiče na sposobnost algoritma da precizno identifikuje anatomske strukture, što je ključno za medicinsku analizu slika u praksi.

Tabela 20. Evaluacija kvaliteta segmentacije primjenom Random Walks algoritma na dataset-u sa medicinskim slikama sa bukom

Metode	
Dice koeficijent	0.56
Jaccard-ov indeks	0.40



Slika 50. Prikaz originalnih slika i segmentiranih verzija korišćenjem Random Walks algoritma na dataset-u sa bukom

5.6 Felzenszwalb-Huttenlocher algoritam

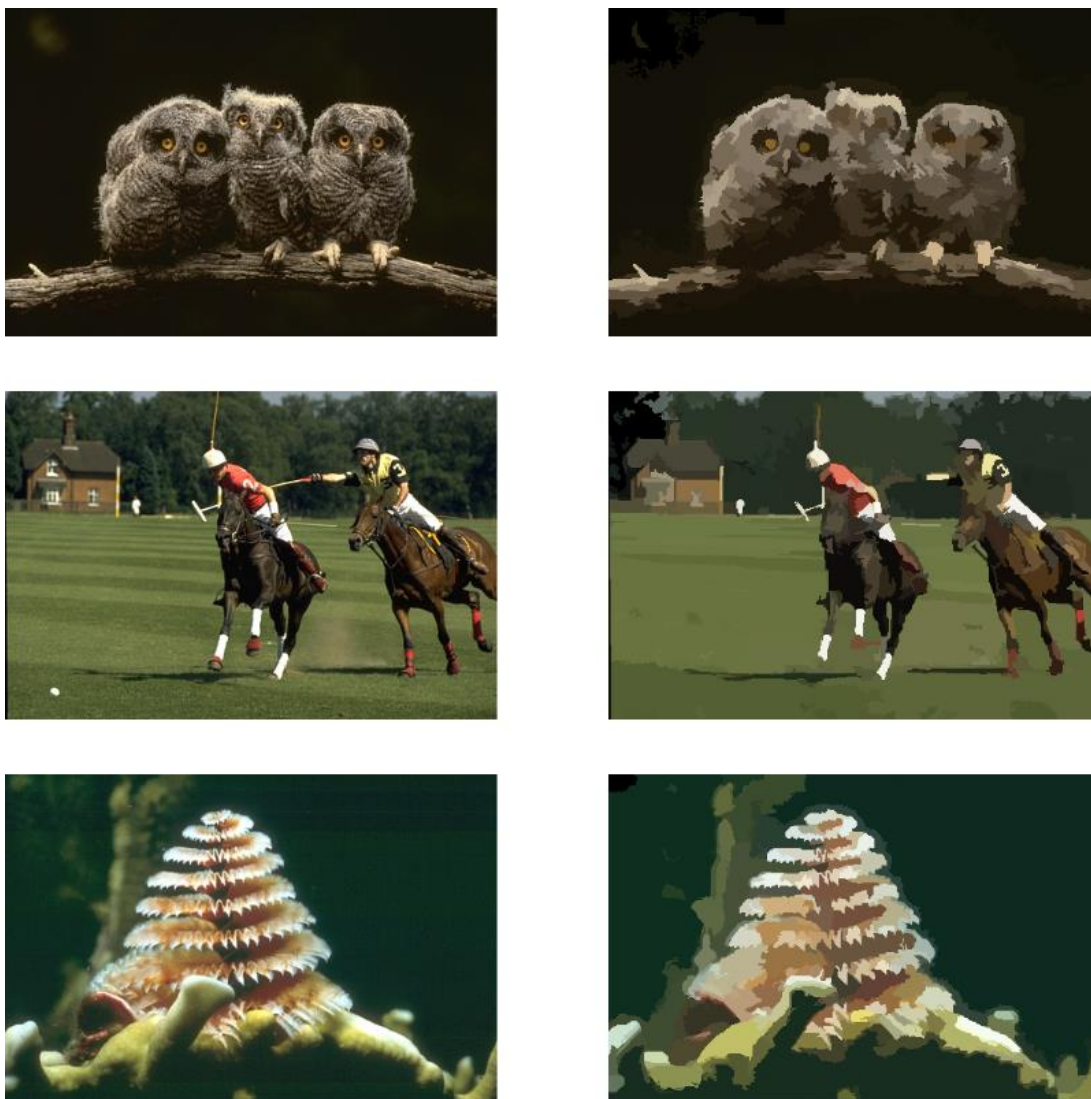
5.6.1 BSDS500 (Berkeley Segmentation Dataset 500)

Na slici 51 prikazani su pojedinačni primjeri rezultata segmentacije dobijeni primjenom Felzenszwalb-Huttenlocher algoritma na dataset od 500 slika, bez dodatka šuma. Ovaj algoritam koristi grafički pristup zasnovan na minimalnim razlikama između susjednih piksela, pri čemu se kreiraju regioni koji predstavljaju homogene segmente u slici. Prikazani primjeri ilustruju sposobnost algoritma da grupiše piksele u homogene regione, čuvajući osnovne oblike i konture objekata. Ova metoda omogućava segmentaciju slike na višem nivou, prikazujući objekte kao cjelovite entitete, ali može biti osjetljiva na sitne detalje unutar kompleksnih struktura.

U tabeli 21 prikazane su prosječne vrijednosti evaluacionih metrika za cijeli dataset. Dice koeficijent iznosi 0.44, što ukazuje na manje preklapanje između segmentiranih regiona i ground-truth oznaka u poređenju sa prethodnim algoritmom. Jaccard-ov indeks, sa vrijednošću od 0.31, potvrđuje nižu preciznost segmentacije pomoću ovog algoritma. Ove vrijednosti sugerišu da Felzenszwalb-Huttenlocher algoritam daje solidne, ali ne optimalne rezultate za ovaj tip dataset-a bez dodatnog šuma, što može biti rezultat specifične metode segmentacije zasnovane na boji i teksturi.

Tabela 21. Evaluacija kvaliteta segmentacije primjenom Felzenszwalb-Huttenlocher algoritma na dataset-u

Metode	
Dice koeficijent	0.44
Jaccard-ov indeks	0.31



Slika 51. Prikaz originalnih slika i segmentiranih verzija korišćenjem Felzenszwalb-Huttenlocher algoritma na dataset-u

Na slici 53 prikazani su pojedinačni primjeri rezultata segmentacije dobijeni primjenom Felzenszwalb-Huttenlocher algoritma na dataset od 500 slika sa dodatkom Gausovskog šuma (varijansa šuma: 0.01). Dodavanje šuma ima za cilj testiranje otpornosti algoritma na degradaciju kvaliteta slike i procjenu koliko će segmentacija ostati stabilna u uslovima narušenog signala. Prikazani primjeri ilustruju kako šum smanjuje jasnoću granica i ometa precizno odvajanje objekata od pozadine, što utiče na performanse algoritma.

U tabeli 22 prikazane su prosječne vrijednosti evaluacionih metrika za cijeli dataset sa dodatkom šuma. Dice koeficijent iznosi 0.42, što ukazuje na smanjenje preklapanja između segmentiranih regiona i ground-truth oznaka u poređenju sa datasetom bez šuma. Jaccard-ov indeks, sa vrijednošću od 0.30, takođe potvrđuje smanjenu preciznost segmentacije u prisustvu šuma. Ovi rezultati sugerišu da dodatak šuma otežava algoritmu prepoznavanje finih detalja i jasno odvajanje objekata, čime se smanjuje njegova pouzdanost u uslovima narušenih slika.

Tabela 22. Evaluacija kvaliteta segmentacije primjenom Felzenszwalb-Huttenlocher algoritma na dataset-u sa bukom

Metode	
Dice koeficijent	0.42
Jaccard-ov indeks	0.30



Slika 52. Prikaz originalnih slika i segmentiranih verzija korišćenjem Felzenszwalb-Huttenlocher algoritma na dataset-u sa bukom

5.6.2 Dataset sa medicinskim slikama

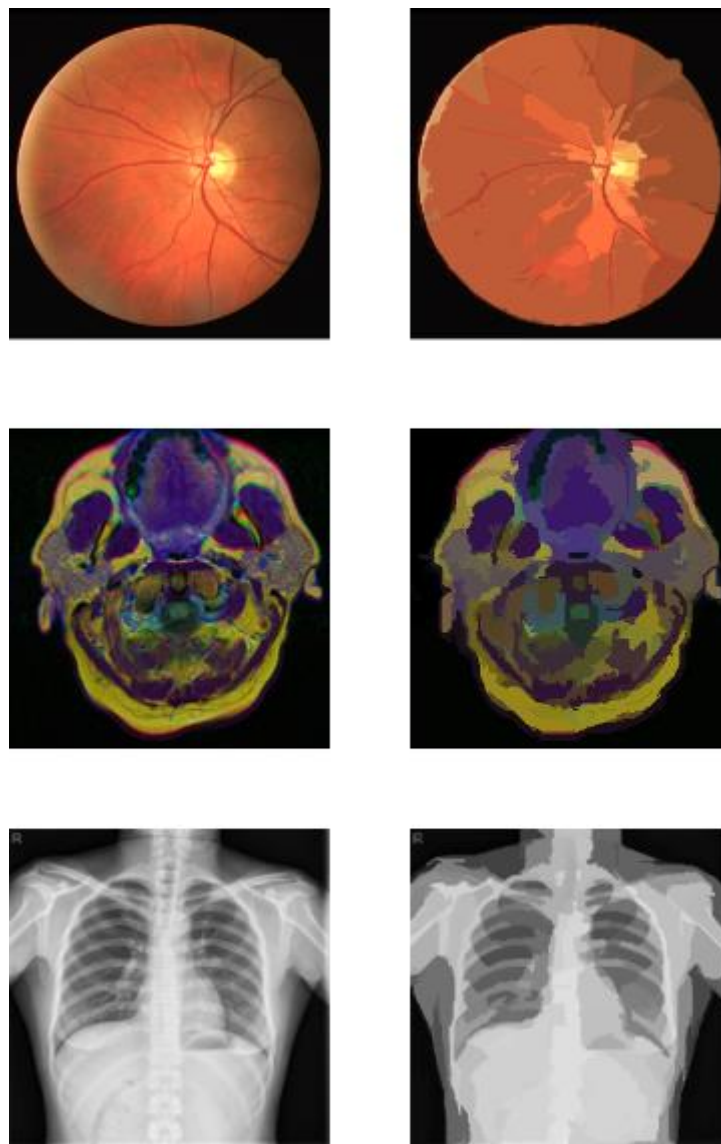
Na slici 54 prikazani su pojedinačni primjeri rezultata segmentacije medicinskih slika dobijeni primjenom Felzenszwalb-Huttenlocher algoritma na dataset koji se sastoji od 300 medicinskih slika. Ovaj dataset obuhvata razne tipove slika, uključujući snimke oka, mozga i pluća, čime se omogućava evaluacija algoritma u različitim medicinskim kontekstima.

Prikazani primjeri ilustruju sposobnost algoritma da zadrži osnovne konture i strukture objekata u različitim tipovima medicinskih slika. Ova metoda omogućava odvajanje objekata u cjelinu, ali može biti osjetljiva na suptilne promjene unutar složenih struktura, što može uticati na preciznost segmentacije.

U tabeli 23 prikazane su prosječne vrijednosti evaluacionih metrika za cijeli medicinski dataset. Dice koeficijent iznosi 0.70, što ukazuje na solidno preklapanje između segmetisanih regiona i ground-truth oznaka. Jaccard-ov indeks, sa vrijednošću od 0.54, takođe potvrđuje zadovoljavajuću preciznost algoritma u ovim uslovima. Ovi rezultati sugerišu da Felzenszwalb-Huttenlocher algoritam postiže dobre rezultate za segmentaciju medicinskih slika, omogućavajući jasnu segmentaciju ključnih anatomske struktura u većini slučajeva.

Tabela 23. Evaluacija kvaliteta segmentacije primjenom Felzenszwalb-Huttenlocher algoritma na dataset-u sa medicinskim slikama

Metode	
Dice koeficijent	0.70
Jaccard-ov indeks	0.54



Slika 53. Prikaz originalnih slika i segmentiranih verzija korišćenjem Felzenszwalb-Huttenlocher algoritma na dataset-u

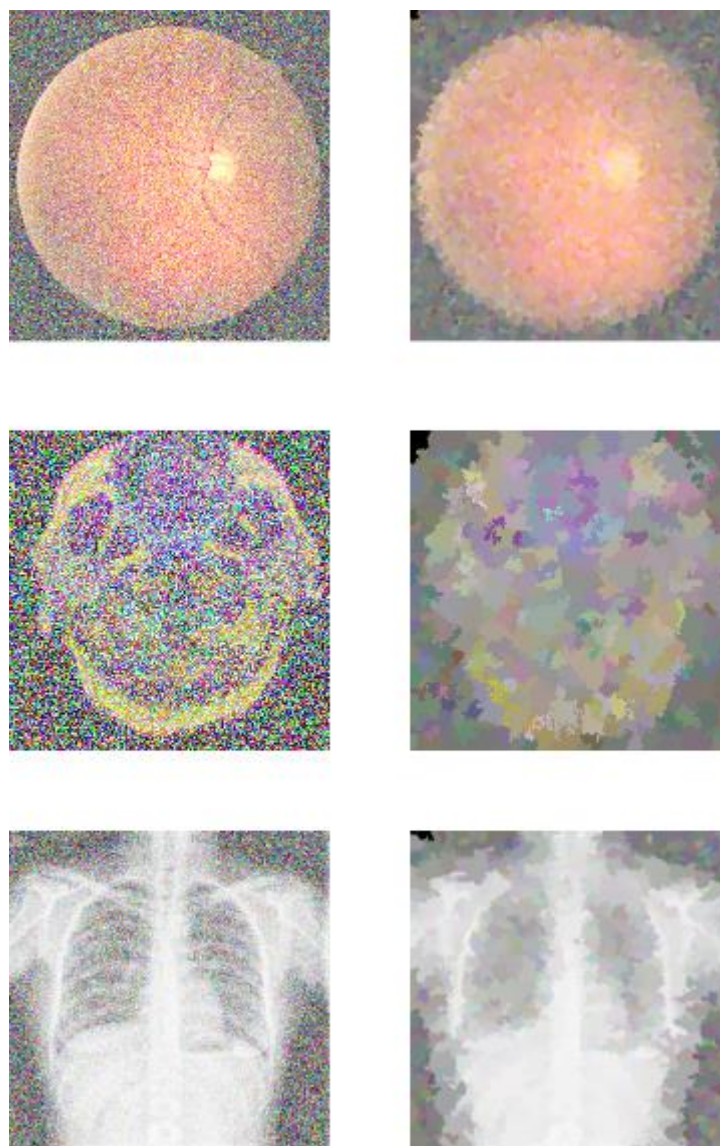
Na slici 55 prikazani su pojedinačni primjeri rezultata segmentacije medicinskih slika dobijeni primjenom Felzenszwalb-Huttenlocher algoritma na dataset sa dodatkom Gausovskog šuma (varijansa šuma: 0.01). Dodavanje šuma omogućava procjenu otpornosti algoritma na smetnje u slici, čime se testira njegova sposobnost da zadrži ključne anatomske strukture uprkos narušenom kvalitetu slike. Prikazani primjeri pokazuju kako šum utiče na jasnoću granica i izaziva smanjenje preciznosti u prepoznavanju i odvajanju objekata.

U tabeli 24 prikazane su prosječne vrijednosti evaluacionih metrika za cijeli medicinski dataset sa dodatkom šuma. Dice koeficijent iznosi 0.68, što ukazuje na nešto smanjen stepen preklapanja između segmentiranih regiona i ground-truth oznaka u poređenju sa slikama bez šuma. Jaccard-ov indeks, sa vrijednošću od 0.51, takođe pokazuje blago smanjenu preciznost

segmentacije. Ovi rezultati sugerišu da, iako šum utiče na performanse algoritma, Felzenszwalb-Huttenlocher algoritam i dalje pokazuje otpornost i sposobnost da identifikuje glavne anatomske strukture, uz umjeren pad tačnosti u prisustvu šuma.

Tabela 24. Evaluacija kvaliteta segmentacije primjenom Felzenszwalb-Huttenlocher algoritma na dataset-u sa medicinskim slikama sa bukom

Metode	
Dice koeficijent	0.68
Jaccard-ov indeks	0.51



Slika 54. Prikaz originalnih slika i segmentiranih verzija korišćenjem Felzenszwalb-Huttenlocher algoritma na dataset-u sa bukom

Zaključak

U radu se istražuju ključni koncepti grafovskog grupisanja podataka i segmentacije slika, s naglaskom na značaj i specifičnosti ovih tehnika u obradi kompleksnih informacija. Grupisanje podataka omogućava organizaciju podataka na osnovu sličnosti i zajedničkih karakteristika, što pojednostavljuje analizu i otkriva strukturu unutar velikih i neorganizovanih skupova podataka. Ovaj proces je posebno važan u oblastima gdje je potrebno brzo prepoznati obrasce, kao što su analiza ponašanja korisnika, medicinska istraživanja, te sistemi preporuka. Grupisanje u radu obuhvata tehnike koje koriste grafove za reprezentaciju odnosa među podacima, čime se omogućava jasnija vizuelizacija i analiza grupa koje se formiraju na osnovu prirodnih veza između pojedinačnih elemenata.

Segmentacija slika je druga ključna tema istraživanja, koja predstavlja proces dijeljenja slike na razumljive i koherentne regije kako bi se olakšala dalja analiza i prepoznavanje objekata. U domenu obrade slika, segmentacija je od suštinskog značaja, posebno u oblastima kao što su medicina, autonomna vožnja i geološka istraživanja, gdje je potrebno izdvojiti značajne strukture na slikama i omogućiti detaljnu analizu. Primjenom grafovskih tehnika u segmentaciji slika postiže se visoka preciznost u izdvajanju objekata, što je posebno korisno kod složenih scena i slika sa puno detalja.

Opisani su algoritmi Spektralno grupisanje i Minimum Spanning Tree (MST) za grupisanje podataka, kao i Random Walks i Felzenszwalb-Huttenlocher za segmentaciju slika. Pored samih algoritama, predstavljeni su korišćeni datasetovi, među kojima su Iris, Wine i Breast Cancer za grupisanje, dok su za segmentaciju korišćeni standardni setovi slika i medicinske slike. Evaluacija algoritama sprovedena je uz pomoć četiri metrike: Silhouette Score, Davies-Bouldin Index, Calinski-Harabasz Index za grupisanje, te Dice koeficijent i Jaccardov indeks za segmentaciju.

Eksperimentalni dio rada testira otpornost algoritama na šum i buku, što omogućava sveobuhvatno ocjenjivanje njihove stabilnosti u otežanim uslovima. Analiza rezultata pokazuje da su algoritmi zadržali visoke performanse čak i uz prisustvo ometajućih faktora, čime je potvrđena njihova primjenjivost u stvarnim scenarijima. Ovaj rad predstavlja vrijedan resurs za buduća istraživanja i unaprijeđenja u oblasti grupisanja i segmentacije, te doprinosi efikasnijem odabiru odgovarajućih algoritama za različite primjene.

Literatura

- [1] A.K. Jain, “Data Clustering: 50 Years Beyond KMeans”, *Pattern Recognition Letters*, Vol 31 Issue 8, pp. 651-666., 2010.
- [2] Noh, Y., Koo, D., Kang, Y., Park, D., & Lee, D. “Automatic Crack Detection on Concrete Images Using Segmentation via Fuzzy C-means Clustering”, pp. 877–880., 2017.
- [3] K. M. and Y. J. Chanu N. Dhanachandra, “Image segmentation using K-means clustering algorithm and subtractive clustering algorithm”, *ScienceDirect, Computer science procedia*, pp. 764–771, 2015.
- [4] N. R. Pal, S. K. Pal, “A review on image segmentation techniques”, *Pattern recognition* 26 (9) pp. 1277–1294., 1993.
- [5] Mathilde Caron, Piotr Bojanowski, Armand Joulin, and Matthijs Douze. “Deep clustering for unsupervised learning of visual features”, *Proceedings of the European Conference on Computer Vision (ECCV)*, pp. 132–149, 2018.
- [6] Le, T. V., Kulikowski, C. A., Muchnik, I. B, “Coring method for clustering a graph”, *Pattern Recognition, 19th International Conference on*. pp. 1-4, 2008.
- [7] L. Kaufman and P. J. Rousseeuw, “Finding Groups in Data: an Introduction to Cluster Analysis”, *John Wiley and Sons*. pp. 126-163, 1990.
- [8] Zhang, Tian, Ramakrishnan, Raghu, Livny, Miron. “BIRCH: An efficient data clustering method for very large databases”, *Internat. Conf. on Management of data*, vol. 25, pp. 103–114, 1996.
- [9] Y. Kim, K. Shim, M. S. Kim and S. Lee, “DBCURE-MR: an efficient density-based clustering algorithm for large data using MapReduce”, *Information Systems*, vol. 42, pp. 15-35, 2014.

- [10] Anamika Ahirwar, "Study of Techniques used for Medical Image Segmentation and Computation of Statistical Test for Region Classification of Brain MRI", *International Journal of Information Technology and Computer Science*, pp. 44- 53, 2013.
- [11] Rozy Kumari and Narinder Sharma, "A Study on the Different Image Segmentation Technique", *International Journal of Engineering and Innovative Technology (IJEIT)* Volume 4, Issue 1, July, pp. 284-289, 2014.
- [12] Lei, T.; Jia, X.; Zhang, Y.; Liu, S.; Meng, H.; Nandi, A.K. "Superpixel-based fast fuzzy C-means clustering for color image segmentation", *IEEE Trans. Fuzzy Syst.*, pp. 1753–1766, 2018.
- [13] Yang Y., "Image segmentation based on fuzzy clustering with neighbourhood information", *Opt. Appl*, pp. 135-147, 2009.
- [14] Khrissi L, El Akkad N, Satori H, Satori K, "Image segmentation based on k-means and genetic algorithms", *Embedded systems and artificial intelligence. Springer, Singapore*, pp 489–497, 2020.
- [15] Data Clustering: Methods & Applications, dostupno na: <https://encord.com>, pristupano 5. novembra 2024.
- [16] Jang, J.-S. R., Sun, C.-T., Mizutani, E. Neuro-Fuzzy and Soft Computing – A Computational Approach to Learning and Machine Intelligence, *Prentice Hall*, 1997.
- [17] Azuaje, F., Dubitzky, W., Black, N., Adamson, K., "Discovering Relevance Knowledge in Data: A Growing Cell Structures Approach", *IEEE Transactions on Systems, Man, and Cybernetics-Part B: Cybernetics*, Vol. 30, No. 3, pp. 448, 2002.
- [18] A few types of clustering algorithms, dostupno na: <https://computing4all.com/a-few-types-of-clustering-algorithms/>, pristupano 5. novembra 2024.
- [19] Data Clustering: Methods & Applications, dostupno na: <https://encord.com>, pristupano 5. novembra 2024.

- [20] Shapiro, L. G., Stockman, G. C. “Computer Vision”, *Prentice-Hall, New Jersey*, pp. 279–325, 2001.
- [21] Barghout, L., Lee, L. W., “Perceptual information processing system”, *Paravue Inc. U.S. Patent Application 10/618,543*, 2003.
- [22] Guo, D., Pei, Y., Zheng, K., Yu, H., Lu, Y., Wang, S., “Degraded Image Semantic Segmentation With Dense-Gram Networks”, *IEEE Transactions on Image Processing*, Vol. 29, pp. 782–795, 2020.
- [23] Yi, J., Wu, P., Jiang, M., Huang, Q., Hoeppner, D. J., Metaxas, D. N., “Attentive neural cell instance segmentation”, *Medical Image Analysis*, Vol. 55, str. 228–240, 2019.
- [24] Kirillov, A., He, K., Girshick, R., Rother, C., Dollár, P., “Panoptic Segmentation”, 2018
- [25] Explain Image Segmentation: Techniques and Applications, dostupno na: <https://www.geeksforgeeks.org>, pristupano 5. novembra 2024.
- [26] Educational Data Mining, dostupno na: <https://educationaldatamining.org>, pristupano 5. novembra 2024.
- [27] What is k-means clustering?, dostupno na: <https://www.ibm.com/topics/k-means-clustering>, pristupano 5. novembra 2024.\
- [28] Spectral clustering: Definition, Operation, Use, dostupno na: <https://datascientest.com>, pristupano 5. novembra 2024.
- [29] Mininum Spanning Tree dostupno na: <https://docs.scipy.org/doc/scipy/mtg.html>, pristupano 5. novembra 2024.’
- [30] What is MST, dostupno na: <https://www.geeksforgeeks.org/what-is-minimum-spanning-tree-mst/>, pristupano 5. novembra 2024.
- [31] Grady, L., “Random Walks for Image Segmentation”, *IEEE Transactions on Pattern Analysis and Machine Intelligence*, Vol. 28, No. 11, str. 1768–1783, 2006.

- [32] Random walker segmentation", dostupno na: https://random_walker_segmentation.html , pristupano 5. novembra 2024.
- [33] Research platform for scientists, dostupno na: <https://www.researchgate.net>, pristupano 5. novembra 2024.
- [34] N. Otsu (1979). "A threshold selection method from gray-level histograms", IEEE Transactions on Systems, Man, and Cybernetics, 1979.
- [35] R. A. Fisher, "The use of multiple measurements in taxonomic problems", *Annals of Eugenics*, 179–188, 1936.
- [36] M Yasser H, "Wine Quality Dataset", dostupno na: <https://www.kaggle.com/wine-quality-dataset>, pristupano 5. novembra 2024.
- [37] Breast Cancer Dataset, dostupno na: <https://www.kaggle.com/breast-cancer-wisconsin-data>, pristupano 5. novembra 2024.
- [38] What is Silhouette Score?, dostupno na: <https://www.educative.io/answers/what-is-silhouette-score>, pristupano 5. novembra 2024.
- [39] Caliński, Tadeusz; Harabasz, "A dendrite method for cluster analysis". Communications in Statistics, pp. 1–27, 1974
- [40] Understanding Evaluation Metrics in Medical Image Segmentation, dostupno na: <https://medium.com>, pristupano 5. novembra 2024.
- [41] Jaccard Similarity Made Simple: A Beginner's Guide to Data Comparison, dostupno na: <https://medium.com>, pristupano 5. novembra 2024.