

İnternet Tabanlı İmge Arama Sonuçlarının Histogram Tabanlı Baskın Kümeler ile Gruplandırılması

Evren Ferhat EMEKDAŞ^{1,3}

Ziya TELATAR²

^{1,2}Elektronik Mühendisliği Bölümü, Ankara Üniversitesi, Ankara

³Elektrik-Elektronik Mühendisliği Bölümü, Başkent Üniversitesi, Ankara

¹e-posta: eemekdas@baskent.edu.tr

²e-posta: Ziya.Telatar@eng.ankara.edu.tr

Özetçe

Gün geçtikçe artan çoklu ortam verilerinin kullanımı, beraberinde muazzam miktarlardaki görsel içerik içinden istenilen verilere erişim sorununu da beraberinde getirmektedir. Bu sorunu çözmek için kullanıcılar, internet tabanlı arama motorlarına yönelmişlerdir. Ne var ki sıklıkla kullanılan bu popüler arama motorlarının hiç birisi içerik tabanlı bir arama çözümü sunmamaktadır. Bu çalışmada, video bölütleme çalışmalarında başarıyla uygulanan baskın kümeler yönteminin içerik tabanlı bir ayıklama işlemi için kullanılabileceği ön görüşünden hareketle, internet tabanlı imge arama motorlarından elde edilen sonuçların gruplandırması gerçekleştirilmiştir.

1. Giriş

Çoklu ortam verilerinin kullanımı, masrafları azalan sayısal ortam gereçlerinin aşırı kullanımıyla beraber gün geçtikçe artmaktadır [1]. Günümüzde görsel verilerin saklanması veya paylaşılması bir sorun olmaktan çıkmış ancak beraberinde bu muazzam büyüklükteki içerik içinden istenilen verilere erişim sorununu getirmiştir. Bu uygulama tabanlı engelleri aşmak için yeni eğilim görüntüleri gruplandırma ve indisleme amaçlı algoritmalar geliştirilmesi yönündedir.

Bu çalışmanın amacı, Google™ ve benzeri arama motoru sonuçları üzerinde çizge tabanlı bir gruplandırma çalışması yapmaktır. Bu nedenle, bu çalışmanın ilgi alanı yeni imgeler aramak yerine arama motorlarının daha önceden bulabildiği imge veritabanları üzerinde çalışmalar yapmak olarak tanımlanabilir. Başka bir deyişle, veri tabanının yeniden oluşturulmasından ziyade, daha önceden oluşturulmuş veri tabanlarının içindeki imgelerin gruplandırılması optimizasyonu bu çalışmanın vizyonunu oluşturmaktadır.

Bu çalışmanın amaçlarından birisi de hesaplama yükünü azaltmak olduğu için daha az verinin bulunduğu bir yöntemin kullanılması daha akıllıca olacaktır. Bir imgenin histogramı büyük bir veri kaybına rağmen halen imge hakkında önemli bilgiler içermekte ve ihtiyaç duyulan daha küçük veri yığınlarını sağlayabilmektedir. Ayrıca veri miktarının azalması, problemin kötü koşullu (ill-conditioned) olma olasılığını azaltacak ve hesaplama hatasının getireceği sorunlar azalacaktır.

Baskın kümeler, imge bölütleme problemleri gibi parçalı (düz) gruplandırma ile ilişkisi ispatlanmış yeni bir çizge kuramsal anlayıştır. Bununla birlikte, pek çok bilgisayar tabanlı görüntü uygulamasında, örneğin bir imge

veritabanının düzenlenmesinde, verinin hiyerarşik bir düzenleme ile gruplandırılması önemlidir ve baskın küme kapsamında bu düzenlemenin nasıl yapılacağı çok açık değildir [2].

Düzenleyici parametre sıfıra eşitlendiğinde yerel çözümlerin baskın kümelerle birebir ilişki içinde olduğu bilinmektedir. Fakat, parametre pozitif olduğunda ortaya ilginç bir görüntü çıkmaktadır. Düzenleyici parametre için sınırlara önceden belirlenmiş bir eşik değerinden daha küçük büyüklüklere sahip grupları içeren yerel çözümlerin kümesini çıkarılmasını sağlayacak şekilde karar verilir. Bu durum gruplandırma süreci sırasında düzenleyici parametrenin uygun bir şekilde değiştirilmesi fikri üzerine kurulmuş olan yeni (parçalayıcı) hiyerarşik bir yaklaşımın kullanılmasını ortaya çıkarır. Literatürde, üç farklı benzerlik matrisi (ve veritabanı) ile yapılan deneyler olduğu görülmüş ve elde edilen sonuçlar, bu yaklaşımın etkinliğini doğrulamıştır.

Verinin (ya da gruplandırmanın) katkısız parçalanması bilgisayar tabanlı görüntü araştırmalarında yaygınlaşan bir problemdir ve yakın zamanlarda gruplandırılacak verinin (pikseller, kenar elemanları, vs.) kenarlarının komşuluk ilişkilerini gösterdiği ve ağırlıkların veriler arasındaki benzerliği yansıttığı bir benzerlik (kenar ağırlıklı) çizgesinin tepe noktaları olarak belirlendiği çizge tabanlı yaklaşımların [3,4,5,6,7] ilgi alanına yeniden girdiği gözlenmektedir. Çizge kuramsal gruplandırma algoritmaları temel olarak en az kapsama ağacı [8] veya en az kesim [6,7,9] gibi benzerlik çizgesi içindeki kombinasyonel kesim yapılarının aranmasını içerir ve bu yöntemler arasında klasik bir yaklaşım (tam-bağlantı algoritması [10,11,12]) *clique* adındaki bir tüm alt çizgenin aranmasına indirger. “Grup” kavramının kesin ve üzerinde çalışılabilir bir tanımı olmadığından, bir parçayı elde etmek için tek bir “en iyi” kriter yoktur [12]. Bazı yazarlar [13,14] *maximal clique* kavramının grup olgusunun en katı tanımı olduğunu iddia etmektedir [2,15].

Bu çalışmanın ikinci bölümünde kuramsal temeller ele alınmıştır. Kuramsal temellerde ilk önce imge arama olgusu tartışılmış ve yapılacak çalışmada histogram kullanımının nedenleri açıklanmıştır. Daha sonraki alt bölümlerde sınıflandırma ve gruplandırma kavramlarının temelleri açıklanmıştır. Gruplandırma hakkında verilen bilgilerden sonra, çizge kuramı ve bir alt başlık olarak baskın kümeler açıklanmıştır. Bu bölümde son olarak performans değerlendirmesi için kullanılan yöntem belirtilmiştir. Üçüncü bölümde, açıklanan kuramsal bilgiler ışında yapılan uygulama ve sonuçlar hakkında bilgi verilmiştir. Son

bölümde ise önerilen yöntemle ulaşılan sonuçlar tartışılmıştır.

2. Kuramsal Temeller

2.1. İmge Arama

Sayısal ortam üzerinde artan veri miktarıyla beraber, istenen veriye ulaşılabilmesi için arama motorlarının kullanımı gerekli olmaktadır. Gerek ağ tabanlı, gerekse masaüstü arama motorları arama işlemini yazı tabanlı gerçekleştirmektedir. Pek çok ağ tabanlı arama motoru (Altavista™, Google™, Yahoo™, vb.) imge arama üzerinde farklı çalışmalar gerçekleştirmektedir. Günümüzde, Google™ arama motoru bu arama yöntemleri arasında en başarılılarından birisi olarak geçmektedir [16].

Google™, resime yakın olan metinleri, resim başlıklarını ve resim bilgilerini (HTML tabanlı meta verisi gibi), boyut ve çözünürlük bilgisini, erişim kolaylığını, erişim sayısını tarif edebilecek düzinelerce faktörü analiz eder. Ayrıca Google™, çiftleri elemek ve en kaliteli resimleri ilk olarak sunmak için, karmaşık algoritmalar kullanır. Google™'ın Resim Aramasını kullanarak, web üzerinde 250 milyon resimden fazlası aranabilmektedir. Bununla birlikte internet üzerinde henüz Google™'ın indeksine eklenmediği daha pek çok resim mevcuttur. [17]

Google™ grafik arama motoru, her ne kadar diğer arama motorlarına göre başarılı sonuçlar vermekte ise de arama sonucunda başarımlar her zaman istenildiği gibi olmamakta ve arama sonucunda ulaşılmak istenen imgelerden farklı sonuçlara da ulaşılmaktadır. Buna bir örnek oluşturmak için Google™ grafik arama motoru kullanılarak “ocean” kelimesi girilerek okyanus imgelerine ulaşılmaya çalışılmıştır. Arama sonucunda okyanus imgeleriyle beraber, “ocean” olarak isimlendirilmiş imgeler de verilmektedir. Arama işleminin sonucunda elde edilen bazı imgeler Şekil 1’de gösterilmiştir.



Şekil 1: Google™ grafik arama motoruyla yapılan “ocean” sorgusu sonucunda çıkan bazı imgeler

Literatürde arama motorlarından gelen sonuçların bir yöntem çerçevesinde tekrar sıralanmasını öneren çok fazla çalışma bulunmadığı gibi sonuçların gruplandırılması üzerine çok daha az çalışma bulunmaktadır [16]. Ben Haim *et al.* [18]’in çalışması bu az sayıda sıralama çalışmalarından birisi olarak sayılabilir. Schroff *et al.* [19]’un yaptığı çalışma da bu alanda yapılan çalışmaların bir başka örneğidir. Schroff bu çalışmada, internet üzerinde belirli bir konu üzerinde

yapılan imge araştırmasının sonuçlarını toplayarak kullanılabilir bir veri tabanı oluşturmaktadır. Bunun için Google™’ın ağ tabanlı aramasını ve imge arama sonuçlarını kullanarak önem sırasına göre dizmekte ve önemli resimleri veri kümesine eklemektedir. Deselaers *et al.* [20]’ın yaptığı çalışma, bu alanda yapılan bir başka çalışma olarak gösterilebilir. Farklı bir bakış açısı olarak, Deselaers çalışmasında, benzer öznitelikler bulmak yerine değişmez özniteliklerin belirlenmesi ve bu öznitelikler kullanılarak gruplandırma çalışmalarının yapılmasını amaçlamıştır. Çalışmada Google™ imge arama sonuçları kullanılarak bir veri kümesi oluşturulmakta ve bu veri kümesindeki imgelerin “değişmez öznitelikleri” bulunarak bu imgeler gruplandırılmaktadır.

Daha önce de belirtildiği üzere, bu çalışmada da Google™ imge arama sonuçlarından elde edilecek veritabanlarının kullanılması amaçlanmaktadır. Burada kritik olan nokta, imgelerin hangi özniteliklerinden faydalanılarak gruplandırma işleminin yapılacağıdır. İmgelerin uzamsal düzlemdeki öznitelikleri kullanılabileceği gibi belirli dönüşüm uzaylarındaki öznitelikleri de kullanılabilir. Bu noktada unutmamak gerekir ki, dönüşüm uzayına geçiş sürecinde fazladan bir hesaplama yükü ile karşılaşmaktadır. Ne var ki başarımların artırımı amaçlanırken, aynı zamanda uygulanabilirliği sağlamak amacıyla hesaplama yükünün en az seviyede tutulması da gerekmektedir. Bu nedenle günümüz teknolojisi dahilinde dönüşüm uzayına geçişin kullanılması çok akıllıca gözükmemektedir. Uzamsal düzlemde yapılabilecek uygulamaların başında piksel tabanlı benzerliklere bakmanın geldiği söylenebilir. Fakat, bu uygulamada aynı içeriği içeren imgeler arasında imgenin elde edildiği ortam parametrelerinin (ışıklandırma, görüntü kaynağı, gürültü, nesnenin dokusundaki bozukluklar, vb.) farklılığından dolayı benzer içeriği farklı gruplara ayırmak mümkündür. Bu noktada içeriğin fiziksel özelliklerine bakmak bir başka çözüm yolu olabilir. Ne yazık ki bu yöntemin de belirgin dezavantajları vardır. Bunlar arasında çözünürlük farklılığı, imgenin elde edildiği kaynaktan dolayı oluşan renk ve doku farklılıkları sayılabilir. Bir başka yöntem de imge bölütlerinin benzerliklerine bakmaktır. Fakat imge içindeki nesne bir imgede farklı açıdan görüntülenirken, başka bir imgede çok daha farklı bir açıdan görüntülenmiş olabilir. Örneğin, bir insanın bir portre fotoğrafı ile bir profil fotoğrafı bölütlendiğinde birbirine benzer olmayacaktır.

Bu noktaya kadar tartışılan yöntemlerin bir başka kötü yanı da büyük veri bloklarıyla çalışmayı gerektirmesidir. Amaçlardan birisi de hesaplama yükünü azaltmak olduğu için daha az verinin bulunduğu bir yöntemin kullanılması daha akıllıca olacaktır. Bir imgenin histogramı büyük bir veri kaybına rağmen halen imge hakkında önemli bilgiler içermektedir ve ihtiyaç duyduğumuz daha küçük veri yığınlarını bize sağlayabilmektedir. Ayrıca veri miktarının azalması, problemin kötü koşullu olma olasılığını azaltacak ve hesaplama hatasının getireceği sorunlar azalacaktır.

2.2. Sınıflandırma ve Gruplandırma

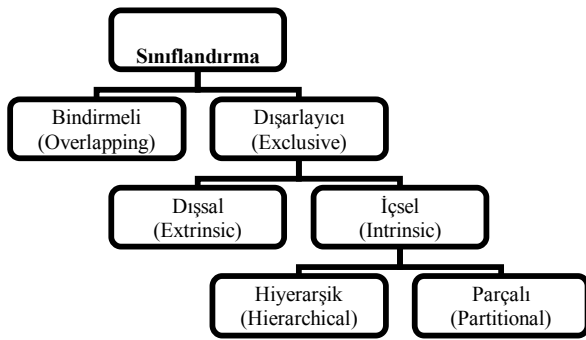
Jain ve Dubes’un [12] tanımına göre grup analizi, nesneleri belirli bir problemin şartları çerçevesinde anlamlı alt kümelere sınıflandırma işlemidir. Bunun sonucu olarak

nesneler, içinde bulundukları popülasyonu karakterize eden etken gösterimi sağlayacak şekilde organize edilir.

Gruplandırma, nesnelerin sonlu kümelerine uygulanan bir sınıflandırma türüdür. Nesneler arasındaki ilişkiler, satır ve sütunların nesnelere karşılık geldiği yakınlık matrisleriyle ifade edilir [12].

Eğer nesneler; desenler veya d-boyutlu metrik uzaydaki noktalar olarak karakterize edilirse, yakınlıklar nokta çiftleri arasındaki uzaklıklar (ör: Öklid uzaklığı) olarak seçilebilir [12]. Nesne çiftleri arasında anlamlı uzaklık ölçümleri ya da yakınlıklar oluşturulamazsa, anlamlı bir grup analizi mümkün olmamaktadır. Yakınlık matrisi, gruplandırma algoritmasında kullanılabilen tek giriş olgusudur [12].

Gruplandırma, sınıflandırmanın özel bir türüdür [12]. Sınıflandırma ve gruplandırma arasındaki ilişki hakkında tartışma Kendall [21] tarafından verilmiştir. Lance ve Williams [27] tarafından önerilen sınıflandırma problemleri ağacı Şekil 2’de gösterilmiştir. Bu ağacın her yaprağı, sınıflandırma probleminin farklı türlerini ifade etmektedir. Dışarılayıcı ve içsel sınıflandırma türleri gruplandırma olarak adlandırılır [12].



Şekil 2: Sınıflandırma türleri ağacı

2.3. Çizge Kuramı

Bir grubun evrensel olarak kabul edilmiş bir tanımı olmamasına rağmen, farklı araştırmacılar [12,13,14] grup kavramını değişik şekillerde açıklamaya çalışmıştır [2,15]. Pelillo’ya [22] göre bir grubun iki koşula uyması gerektiği söylenebilir:

Dahili koşul: bir grup içindeki tüm nesnelerin yüksek benzerlik göstermesi gerekir.

Harici koşul: bir grubun dışındaki tüm nesnelerin içindekilerden yüksek farklılık göstermesi gerekir.

Çizge kuramı girişte n nesneli bir kümenin ve çiftli benzerliklerden oluşan $n \times n$ bir benzerlik matrisinin varlığını kabul ederek, giriş nesnelerini olabilecek en homojen şekilde gruplara ayırmayı amaçlar [22]. Burada $V = \{1, 2, \dots, n\}$ tepe noktaları kümesini, $E \subseteq V \times V$ kenar kümesini, $w: E \rightarrow \mathbb{R}^+$ pozitif ağırlık fonksiyonunu ifade etsin.

Gruplanacak olan veriyi iç döngüsü olmayan, yönsüz, kenar ağırlıklı bir çizge, $G = (V, E, w)$ olarak gösterebiliriz. G ’deki tepe noktaları veri değerlerini, kenarlar komşuluk ilişkilerini ve kenar ağırlıkları ise bağlı tepe noktası çiftlerinin benzerliklerini yansıtmaktadır. Girişte verilen $A = \{a_{ij}\}$ benzerlik matrisi ile pozitif ağırlık fonksiyonlarının ilişkisi aşağıdaki gibi olacaktır. [22,23]

$$a_{ij} = \begin{cases} \omega(i, j) & , \text{eğer } (i, j) \in E \\ 0 & \text{diğer durumda} \end{cases} \quad (1)$$

Burada, benzerlik ağırlıklarının üçgen eşitsizliğini bozmadığı ve kendine benzerliklerin $w(i, i)$ tanımlanmadığı varsayılmaktadır [23].

Grup için yapılan tanıma dayanarak (bir grup içindeki tüm nesnelerin yüksek benzerlik göstermesi gerektiği ve bir grubun dışındaki tüm nesnelerin içindekilerden yüksek farklılık göstermesi gerektiği) bir grubun iki temel koşulu sağlaması gerektiği söylenebilir. Bunlar; (a) bir grup yüksek iç homojenliğe sahip olmalıdır, (b) grup içinde kalan elemanlar ile dışında kalan elemanlar arasında yüksek homojensizlik olmalıdır. Elemanlar kenar ağırlıklı bir çizge olarak gösterildiğinde bahsedilen iki koşul “grup içindeki kenarlar arasındaki ağırlıklar büyük, grup elemanları ile dışarıdaki elemanlar arasındaki ağırlıklar küçük olmalıdır” şeklinde yorumlanabilir. Bir grup için yapılan tanıma ulaşmak için, kenar ağırlıklarının atanmasının, tepe noktalarının ağırlıklarının atanmasına giden yolu işaret edeceği ön görüsel düşüncesiyle başlanabilir. Bu bakış açısı, kenar ağırlıklarının atanmasını analiz etmek için daha kolay ve tercih edilebilir bir yol gösterecektir. [22,23]

$i \in S$ olmak üzere $S \subseteq V$ boş olmayan bir tepe noktaları kümesi olsun. Bu durumda, i ’nin S ’e göre ortalama ağırlıklı derecesi

$$\text{awdeg}_S(i) = \frac{1}{|S|} \sum_{j \in S} a_{ij} \quad (2)$$

olarak tanımlanabilir.

Buna ek olarak, eğer $j \notin S$ ise

$$\phi_S(i, j) = a_{ij} - \text{awdeg}_S(i) \quad (3)$$

olarak tanımlanabilir. Burada, tüm $i, j \in V$ ve $i \neq j$ için $\phi_{\{i\}}(i, j) = a_{ij}$ olduğuna dikkat ediniz. Ön görüsel olarak, $\phi_S(i, j)$, i ’nci düğüm ve S ’deki komşuları arasındaki ortalama benzerliğe bağlı olarak i ’nci ve j ’inci düğümler arasındaki benzerliği ölçmektedir [22,23]. Ayrıca, $\phi_S(i, j)$ ’in değeri pozitif veya negatif olabilmektedir.

i ’nin S ’ye göre ağırlığı

$$w_S(i) = \begin{cases} 1 & \text{eğer } |S| = 1 \\ \sum_{j \in S \setminus \{i\}} \phi_{S \setminus \{i\}}(j, i) w_{S \setminus \{i\}}(j, i) & \text{diğer durumda} \end{cases} \quad (4)$$

olarak hesaplanabilir. Burada, tüm $i, j \in V$ ve $i \neq j$ için $w_{\{i,j\}}(i) = w_{\{i,j\}}(j) = a_{ij}$ olduğuna dikkat ediniz. Ayrıca, $w_s(i)$ 'nin basitçe S tarafından tetiklenen alt çizgenin kenarları üzerindeki ağırlıkların bir fonksiyonu olarak hesaplandığını gözlemleyebiliriz. Buna ek olarak S 'nin toplam ağırlığı aşağıdaki gibi olur [23].

$$W(S) = \sum_{i \in S} w_s(i) \quad (5)$$

Ön görüsel olarak, $w_s(i)$, $S \setminus \{i\}$ tepe noktaları arasındaki tüm ortalama benzerliğe bağlı olarak i 'nci tepe noktası ile $S \setminus \{i\}$ tepe noktaları arasındaki tüm ortalama benzerlik hakkında bir ölçü verir.

2.3.1. Baskın kümeler

Boş olmayan tepe noktaları kümesi $S \subseteq V$ için $w(T) > 0$ koşulunu sağlayan bir boş olmayan $T \subseteq S$ kümesi varsa ve aşağıdaki koşulları sağlıyorsa bu küme baskındır [22,23].

- 1) Her $i \in S$ için, $w_s(i) > 0$ (iç homojenlik)
- 2) Her $i \notin S$ için, $w_{S \cup \{i\}}(i) < 0$ (dış homojensizlik)

Burada baskın kümeler aranılan gruplara denk düşmektedir. İkilik sistemdeki matrisler için baskın kümeler *maximal clique*'e denk düşmektedir.

2.4. Performans Değerlendirmesi

Performans değerlendirme için dört ölçüt seçilmiştir. İlk ölçüt [24]'de kullanılan *F-ölçümü* (F)'dür. Bu ölçüt bulunan grupların kalitesini ölçmektedir. 0-1 aralığında yer alan F değeri, $F=1$ iken mükemmel sonucu ifade etmektedir. Bu analiz imge bazında gerçekleştirilmiştir. Bir başka deyişle, imge grupları karşılaştırıldığından dolayı, bu analiz detaylı bir performans analizidir. GT ve DT , sırasıyla olması beklenen ve tespit edilen grup kümelerini ifade etmektedir. [1]

$$F = \frac{1}{Z} \sum_{C_i \in GT} |C_i| \max_{C_j \in DT} \{f(C_i, C_j)\} \quad (6)$$

Burada,

$$f(C_i, C_j) = \frac{2 \times Re(C_i, C_j) \times Pr(C_i, C_j)}{Re(C_i, C_j) + Pr(C_i, C_j)} \quad (7)$$

$$Z = \sum_{C_i \in GT} |C_i| \quad (8)$$

Geri çağırma (Re) ve kesinlik (Pr) fonksiyonları (Aşağıda kullanılacak geri çağırma ve kesinlik değerleri ile karıştırılmaması için geri çağırma ve kesinlik fonksiyonları sırasıyla Re ve Pr olarak ifade edilmiştir) aşağıdaki şekilde tanımlanmıştır.

$$Re(C_i, C_j) = \frac{|C_i \cap C_j|}{|C_i|} \quad (9)$$

$$Pr(C_i, C_j) = \frac{|C_i \cap C_j|}{|C_j|} \quad (10)$$

Benzer şekilde, geri çağırma ve kesinlik değerleri [25]'de performans değerlendirmesi ölçütleri olarak tanımlanmıştır ve bu değerler F_1 ölçütünün hesaplanmasında kullanılmaktadır. Tanımsal olarak "Geri çağırma", doğru tespitlerin toplam olması beklenen grup sayısına oranıdır. "Kesinlik" ise doğru tespitlerin toplam tespit edilen grup sayısına oranıdır. Bir doğru tespit, tespit edilen grupların üyeleri ile olması beklenen grupların üyelerinin %50'den daha fazla oranda üst üste çakışmasıyla belirlenir. Buna ek olarak, bir olması beklenen grup kümesi ancak bir tespit edilen grup kümesi ile eşleştirilebilir ve bir tespit edilen grup kümesi de sadece bir olması beklenen grup kümesi ile eşleştirilebilir. Geri çağırma ve Kesinlik kullanılarak hesaplanan son ölçüt F_1 şu şekilde ifade edilir:

$$F_1 = \frac{2 \times \text{Geri Çağırma} \times \text{Kesinlik}}{\text{Geri Çağırma} + \text{Kesinlik}} \quad (11)$$

3. Deneysel Sonuçlar

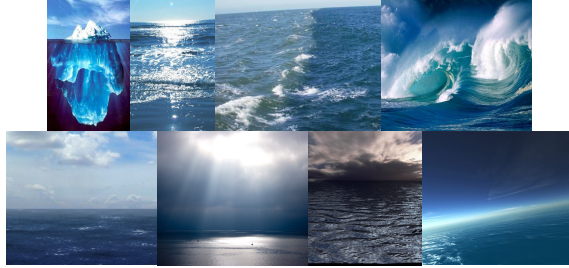
Deneme amaçlı olarak Google™ İmge Arama Motoru üzerinden "ocean" kelimesi aratılmış ve çıkan sonuçlardan ilk 10 sayfada verilenler kaydedilmiştir. Bu 10 sayfa içerisinde Google'ın güncelleme yapmamasından dolayı ulaşılamayan sitelerdeki ve imgenin bağlantısının şifre gerektirdiği imgeler kaydedilmemiştir. Bu yolla, değişik formatlarda (.jpeg, .png, .gif, .bmp, vd.) 146 adet imge elde edilmiştir.

Elde edilen imgelerin histogramları hesaplanmış ve bu histogramlar arasındaki Öklid uzaklıkları hesaplanarak elde edilen değerlerin tersi, benzerlik matrisinin elemanları olarak seçilmiştir. Bu noktada histogramların çözünürlükten bağımsız hale gelmesi için histogramların altında kalan alan 1 olacak şekilde normalize edilmiştir. Seçimin Öklid uzaklığı olarak belirlenmesinin nedeni, denenen 17 farklı çeşit uzaklık ölçütü içerisinde en iyi sonucu vermiş olmasıdır. Elde edilen benzerlik matrisi, daha sonra baskın kümeler yöntemi kullanılarak imgelerin gruplandırılmasında kullanılmıştır.

Google™ imge arama motorundan elde edilen 146 imgenin gruplandırma sonuçları üzerinde bir performans değerlendirmesi yapabilmek için, aynı imgelerin elle gruplandırılması işlemi gerçekleştirilmiştir. Bu gruplandırma gerçekleştirilirken algoritmanın nasıl işlediğini bilen bir kişi (T1) ve algoritmanın nasıl işlediğini bilmeyen bir kişi (T2) tarafından iki ayrı gruplandırma yapılmıştır. T2'nin oluşturduğu gruplara ait iki örnek Şekil 3 ve Şekil 4'te gösterilmiştir. Ayrıca bu 146 imge içerisinde sadece .jpeg imgeleri ayırarak (125 imge) algoritmayı bilen iki ayrı kişi (J1 ve J2) tarafından gruplandırma yapılmıştır. Elle gruplandırma yapan J1 ve T1 kişileri aynı kişilerdir. Bu gruplandırmalar, olması beklenen grup kümelerini (GT) oluşturmaktadır.

Sadece .jpeg türünde imgelerin seçilmesinin nedeni, diğer türlerdeki imgelerin çoğunluğunun farklı renk derinliklerinde bulunması ve bu farklı renk derinliklerinin histogram

üzerindeki değerlerin belirgin seviyelerde toplanmasına neden olmasıdır. Bu durum histogramda pek çok renk seviyesinde değerlerin sıfır olmasına ve Öklid uzaklığında büyük farkların çıkmasına neden olmaktadır. Başarım seviyesinin düşmesinin ana nedenlerinden birisi bu farklılıktır.

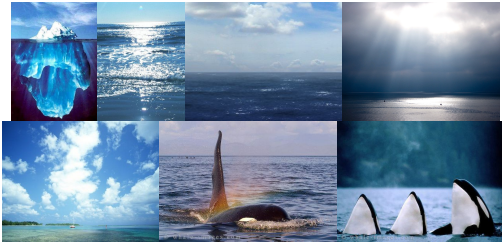


Şekil 3: T2'nin oluşturduğu bir gruptaki örnek imgeler



Şekil 4: T2'nin oluşturduğu bir gruptaki imgeler

Tüm imgelerin gruplandırılması işlemi sonucunda, kullanılan algoritma imgeleri 24 farklı grup olarak ayırırken, T1 ve T2 sırasıyla 35 ve 30 gruba ayırmıştır. Sadece .jpeg türünde imgeler gruplandırıldığında algoritma 19, J1 ve J2 ise sırasıyla 30 ve 32 grup oluşturmuştur. Algoritmanın kullanımı sonucunda oluşturulan gruplara ait iki örnek Şekil 5 ve Şekil 6'da gösterilmiştir.



Şekil 5: Algoritmanın Şekil 3'de oluşturulan gruba en yakın bulduğu sonuç gruptaki örnek imgeler



Şekil 6: Algoritmanın Şekil 4'te oluşturulan gruba en yakın bulduğu sonuç gruptaki imgeler

Algoritma ve kullanıcılar tarafından elde edilen gruplandırma sonuçları doğrultusunda GT ve DT kümeleri belirlenmiştir. Yapılan performans değerlendirmesi sonucunda Bölüm 2.4'de verilen performans ölçütleri hesaplanarak Tablo 1'de verilmiştir.

Tablo 1: Performans ölçütleri

	Kesinlik	Geri Çağırma	F	F1
T1	0.6842	0.4483	0.4626	0.5417
T2	0.4583	0.3667	0.3882	0.4074
J1	0.7623	0.6340	0.4606	0.6923
J2	0.7500	0.5143	0.4699	0.6102

4. Tartışma

Yapılan bu çalışmada elde edilen performans değerleri görece düşük gözükmesine rağmen aslında beklenilenin üstünde başarılar göstermektedir. Olması beklenen grup kümelerinin kişiye özel olması ve algoritmanın bu kişiye özelliklerinden uzak olması başarımın bu seviyede olmasının asıl nedenidir. Bu tür arama sonuçlarında evrensel bir çözümün kolaylıkla elde edilemeyeceği bir gerçektir. Fakat bu çalışma bu tür evrensel bir çözüm arayışına atılan ilk adımdır. Benzer yapıda bir çalışmanın daha önce yapılmamış olması başarımlar konusunda bir karşılaştırma yapılmasını olanaksız kılmaktadır.

Ek olarak bu çalışmada histogramlardan benzerlik matrisine geçiş aşamasında kullanılan uzaklık ölçütü her ne kadar mantıklı gözükse bir çözüm olsa da olması gereken optimum çözüm değildir. Daha uygun bir ölçütün bulunması durumunda başarımın artacağı söylenebilir.

Yapılan karşılaştırmalarda görülebileceği gibi (.jpeg imgelerin aynı bit derinliğinde bulunduğu göz önüne alınarak) aynı bit derinliğindeki imgelerin ayrı ayrı gruplandırılmasının yapılması, farklı bit derinliğindeki imgelerin uyarlanarak aynı bit derinliğine getirilmesi yoluyla gruplandırılmasından daha başarılı sonuçlar verecektir.

Daha başarılı sonuçların elde edilebilmesi için kullanılabilecek yöntemlerden biri de bütün imgelerin benzer büyüklükte bir çözünürlüğe getirilmesi ve gruplandırmanın bu noktadan sonra yapılmasıdır. Bununla birlikte bu tür bir işlemin yapılmasının işlemsel maliyeti oldukça yüksek olacağından, bu yöntem çok tercih edilebilir değildir. Temelde histogramların elde edilmesi bu çözünürlük farklılığı sorununu bir noktaya kadar çözmektedir. Aynı imgenin farklı çözünürlükteki kopyalarının histogramları çıkartıldığında oldukça benzer bir görünüm gösterdikleri gözlemlenir. Bu nedenle bu tür bir yaklaşım bu çalışmada göz önüne alınmamıştır.

5. Teşekkür

Yazarlar hata ayıklama, arama motorlarından elde edilen imge veri tabanlarının oluşturulması ve elle imge gruplama

çalışmaları sırasındaki katkılarından ötürü Arş. Gör. Levent Özparlak'a teşekkür eder.

6. Kaynakça

- [1] Sakarya, U., ve Telatar, Z., "Graph-based multilevel temporal video segmentation", *Multimedia Systems*, Springer Berlin / Heidelberg, 14(5): 277-290, 2008.
- [2] Pavan, M., ve Pelillo, M., "Dominant sets and hierarchical clustering", *Proc. Ninth IEEE International Conference on Computer Vision (ICCV'03)*, 1: 362-369, 2003.
- [3] Perona, P., ve Freeman, W., "A factorization approach to grouping", *In Computer Vision – ECCV'98 (H. Burkhardt and B. Neumann, ed.)*, Springer-Verlag, 655-670, Berlin, 1998.
- [4] Sarkar, S., ve Boyer, K.L., "Quantitative measures of change based on feature organization: Eigenvalues and eigenvectors", *Computer Vision and Image Understanding*, 71(1): 110-136, 1998.
- [5] Aksoy, S., ve Haralick, R.M., "Graph-theoretic clustering for image grouping and retrieval", *Proc. IEEE Computer Society Conference on Computer Vision and Pattern Recognition (CVPR'99)*, 1: 63-68, 1999.
- [6] Shi, J., ve Malik, J., "Normalized cuts and image segmentation", *IEEE Trans. Pattern Analysis and Machine Intelligence*, 22(8): 888-905, 2000.
- [7] Gdalyahu, Y., Weinshall, D., ve Werman, M., "Self-organization in region: stochastic clustering for image segmentation, perceptual grouping, and image database organization", *IEEE Trans. Pattern Analysis and Machine Intelligence*, 23(10): 1053-1074, 2001.
- [8] Zahn, C.T., "Graph-theoretical methods for detecting and describing gestalt clusters", *IEEE Trans. of Computers*, C-20(1): 68-86, 1971.
- [9] Wu, Z., ve Leahy, R., "An optimal graph theoretic approach to data clustering: theory and its application to image segmentation", *IEEE Trans. Pattern Analysis and Machine Intelligence*, 15(11): 1101-1113, 1993.
- [10] Hubert, L.J., "Some applications of graph theory to clustering", *Psychometrika*, 38: 435-475, 1974.
- [11] Matula, D.W., "Graph theoretic techniques for cluster analysis algorithms", *In Classification and Clustering (J. Van Ryzin, ed.)*, Academic Press, Inc., 95-129, New York, 1977.
- [12] Jain, A.K., ve Dubes, R.C., "Algorithms for clustering data", Prentice Hall, 320 p., Englewood Cliffs, New Jersey, 1988.
- [13] Auguston, J.G., ve Minker, J., "An analysis of some graph theoretical clustering techniques", *J. ACM*, 17(4): 571-588, 1970.
- [14] Raghavan, V.V., ve Yu, C.T., "A comparison of the stability characteristics of some graph theoretic clustering methods", *IEEE Trans. Pattern Analysis Machine Intelligence*, 3: 393-402, 1981.
- [15] Pavan, M., ve Pelillo, M., "A new graph-theoretic approach to clustering and segmentation", *Proc. IEEE Computer Society Conference on Computer Vision and Pattern Recognition (CVPR'03)*, 1: 145-152, 2003.
- [16] Sevil, S., Zitouni, H., İkizler, N., Özkan, D., ve Duygulu, P., "Resim arama sonuçlarının çizge tabanlı bir yöntemle yeniden sıralanması", *IEEE 16. Sinyal İşleme, İletişim ve Uygulamaları Kurultayı (SIU 2008)*, Didim, Türkiye, 2008.
- [17] "Görüntü aramaları için Google sıkça sorulan sorular", http://www.google.com.tr/intl/tr/help/faq_images.html. Erişim Tarihi: 28.12.2008
- [18] Ben Haim, N., Babenko, B., ve Belongie, S.J., "Improving web-based image search via content based clustering", *SLAM'06*, 106-111, 2006.
- [19] Schroff, F., Criminisi, A., ve Zisserman, A., "Harvesting image databases from the web", *Proc. 11th International Conference on Computer Vision (ICCV'07)*, 1-8, 2007.
- [20] Deselaers, T., Keysers, D., ve Ney, H., "Clustering visually similar images to improve image search engines", *Informatiktage der Gesellschaft für Informatik*, Springer, 302-305, 2003.
- [21] Kendall, M.G., "Discrimination and classification", *In Multivariate Analysis (P.R. Krishnaiah, ed.)*, Academic Press, Inc., 165-185, New York, 1966.
- [22] Pelillo, M., "Clustering and image segmentation", *Primo Workshop Annuale del Dipartimento di Informatica Università Ca' Foscari di Venezia*, Italy, 2006.
- [23] Pavan, M., "A new graph-theoretic approach to clustering, with applications to computer vision", *Ph.D. thesis, Università di Bologna*, Padova, Venezia, 112 s., Italy, 2004.
- [24] Peng, Y., Ngo, C.W., "Clip-based similarity measure for querydependent clip retrieval and video summarization", *IEEE Trans. Circuits Syst. Video Technol.*, 16: 612-627, 2006.
- [25] Zhai, Y., Shah, M., "Video scene segmentation using Markov chain Monte Carlo", *IEEE Trans. Multimedia* 8, 686-697, 2006.
- [26] Hanjalic, A., "Towards theoretical performance limits of video parsing", *IEEE Trans. Circuits Syst. Video Technol.*, 17: 261-272, 2007.
- [27] Lance, G.N., ve Williams, W.T., "A general theory of classificatory sorting strategies: II, Clustering systems", *Computer Journal*, 10: 271-277, 1967.