



دانشگاه خوارزمی

دانشکده مهندسی کامپیوتر

**بازیابی تصویر بر اساس محتوا در مقیاس بالا با استفاده از الگوی دودویی محلی و
الگوریتم هشینگ چندی سازی تکرارشونده**

پایان نامه برای دریافت درجه کارشناسی ارشد در رشته مهندسی کامپیوتر
گرایش هوش مصنوعی

مونا شاکردنیوی

استاد راهنما

دکتر جمشید شنبه زاده

تابستان ۱۳۹۴



تأییدی هیأت داوران جلسه‌ی دفاع از پایان‌نامه

نام دانشکده: دانشکده مهندسی کامپیوتر

نام دانشجو: مونا شاکردنیوی

عنوان پایان‌نامه: بازیابی تصویر بر اساس محتوا در مقیاس بالا با استفاده از الگوی دودویی محلی و

الگوریتم هشینگ چندی سازی تکرارشونده

تاریخ دفاع: تابستان ۱۳۹۴

رشته: مهندسی کامپیوتر

گرایش: هوش مصنوعی

ردیف	سمت	نام و نام خانوادگی	مرتبه دانشگاهی	دانشگاه یا مؤسسه	امضا
۱	استاد راهنما	دکتر جمشید شنبه زاده	دانشیار	دانشگاه خوارزمی	
۲	استاد مدعو خارجی	دکتر احسان اله کبیر	استاد	دانشگاه تربیت مدرس	

تأییدی صحت و اصالت نتایج

باسمه تعالی

اینجانب مونا شاکردنیوی به شماره دانشجویی ۹۱۳۸۲۲۶۰۶ دانشجوی رشته مهندسی کامپیوتر مقطع تحصیلی کارشناسی ارشد تأیید می‌نمایم که کلیه‌ی نتایج این پایان‌نامه حاصل کار اینجانب و بدون هرگونه دخل و تصرف است و موارد نسخه‌برداری شده از آثار دیگران را با ذکر کامل مشخصات منبع ذکر کرده‌ام. در صورت اثبات خلاف مندرجات فوق، به تشخیص دانشگاه مطابق با ضوابط و مقررات حاکم (قانون حمایت از حقوق مؤلفان و مصنفان و قانون ترجمه و تکثیر کتب و نشریات و آثار صوتی، ضوابط و مقررات آموزشی، پژوهشی و انضباطی ...) با اینجانب رفتار خواهد شد و حق هرگونه اعتراض درخصوص احقاق حقوق مکتسب و تشخیص و تعیین تخلف و مجازات را از خویش سلب می‌نمایم. در ضمن، مسؤولیت هرگونه پاسخگویی به اشخاص اعم از حقیقی و حقوقی و مراجع ذیصلاح (اعم از اداری و قضایی) به عهده‌ی اینجانب خواهد بود و دانشگاه هیچ‌گونه مسؤولیتی در این خصوص نخواهد داشت.

نام و نام خانوادگی: مونا شاکردنیوی

تاریخ و امضا:

مجوز بهره‌برداری از پایان‌نامه

- بهره‌برداری از این پایان‌نامه در چهارچوب مقررات کتابخانه و با توجه به محدودیتی که توسط استاد راهنما به شرح زیر تعیین می‌شود، بلامانع است:
- ☐ بهره‌برداری از این پایان‌نامه برای همگان بلامانع است.
- ☐ بهره‌برداری از این پایان‌نامه با اخذ مجوز از استاد راهنما، بلامانع است.
- ☐ بهره‌برداری از این پایان‌نامه تا تاریخ ممنوع است.

استاد راهنما: دکتر جمشید شنبه زاده

تاریخ:

امضا:

تقديم به:

خانواده عزيزم

خدا...^۱

به من زیستنی عطا کن که در لحظه مرگ، بر بی‌ثمری لحظه‌ای که برای زیستن گذشته است، حسرت نخورم و مُردنی عطا کن که بر بیهودگیش، سوگوار نباشم. بگذار تا آن را، خود انتخاب کنم، اما آنچنان که تو دوست می‌داری.

تو می‌دانی و همه می‌دانند که شکنجه دیدن بخاطر تو، زندانی کشیدن بخاطر تو و رنج بردن به پای تو تنها لذت بزرگ زندگی من است، از شادی توست که من در دل می‌خندم، از امید رهایی توست که برق امید در چشمان خسته‌ام می‌درخشد و از خوشبختی توست که هوای پاک سعادت را در ریه‌هایم احساس می‌کنم. نمی‌توانم خوب حرف بزنم. نیروی شگفتی را که در زیر کلمات ساده و جمله‌های ضعیف و افتاده، پنهان کرده‌ام دریاب، دریاب.

تو می‌دانی و همه می‌دانند که زندگی از تحمیل لبخندی بر لبان من، از آوردن برق امیدی در نگاه من، از برانگیختن موج شعفی در دل من، عاجز است.

تو، چگونه زیستن را به من بیاموز، چگونه مردن را خود خواهم آموخت.

به من توفیق تلاش در شکست، صبر در نومیدی، رفتن بی‌همراه، جهاد بی‌سلاح، کار بی‌پاداش، فداکاری در سکوت، دین بی‌دنیا، مذهب بی‌عوام، عظمت بی‌نام، خدمت بی‌نان، ایمان بی‌ریا، خوبی بی‌نمود، گستاخی بی‌خامی، قناعت بی‌غرور، عشق بی‌هوس، تنهایی در انبوه جمعیت، و دوست داشتن بی‌آنکه دوست بداند، روزی کن.

اگر تنها ترین تنها شوم، باز خدا هست

او جانشین هم‌داشتن هست...

^۱ مناجاتی از دکتر علی شریعتی.

قدردانی

پاس خداوندگار حکیم را که بالطف بی کران خود، آدمی را زیور عقل آراست.
در آغاز وظیفه خود می دانم از زحمات بی دریغ استاد راهنمای خود، جناب آقای دکتر شنبه زاده، صمیمانه تشکر و قدردانی کنم
که قطعاً بدون راهنمایی های ارزنده ایشان، این مجموعه به انجام نمی رسید.
در پایان، بوسه می زنم بر دستان خداوندگار ان مهربانی، پدر و مادر عزیزم و بعد از خدا، ستایش می کنم وجود مقدس شان
را و تشکر می کنم از خانواده عزیزم به پاس عاطفه سرشار و گرمای امید بخش وجودشان، که بهترین پشتیبان من بودند.

مونا ساگردنوی
تابستان ۱۳۹۴

چکیده

تکنیک هشینگ، الگوریتمی کارا به منظور تخمین نزدیک ترین همسایه در پایگاه داده های بزرگ تصاویر می باشد. یادگیری کد باینری یکی از گام های کلیدی به منظور بهبود عملکرد این الگوریتم ها است و هنوز به عنوان چالش پیش رو در این حوزه مطرح می باشد. ورودی های این الگوریتم ها در عملکرد آن ها تاثیر بسزایی دارند. روش پیشنهادی در این رساله با به کارگیری ورودی های مناسب عملکرد الگوریتم های هشینگ را بهبود بخشیده است. در این روش به منظور استخراج ویژگی های تصویر از روش الگوی دودویی محلی استفاده شده است و سپس این بردار به عنوان ورودی به الگوریتم هشینگ چندی سازی تکرارشونده (ITQ) داده شده است که این امر منجر به تولید کدهای فشرده تر، حافظه کمتر، هزینه محاسباتی پایین تر و عملکرد بهتر گردیده است. دلیل این امر ماهیت دودویی و عملکرد مناسب الگوی دودویی محلی در استخراج ویژگی های تصویر می باشد. در پایان به منظور مرتب سازی نتایج بر اساس بیشترین شباهت به تصویر درخواستی کاربر، تکنیک مرتب سازی مجدد با توجه به اهمیت هر بیت اعمال گردیده و به بیت هایی با اهمیت بیشتری وزن بالاتری تخصیص می بابد و وزن های کمتر به بیت های کم اهمیت نسبت داده می شود و با محاسبه فاصله همینگ وزن دار با توجه به اهمیت هر بیت نتایج بر اساس بیشترین شباهت مرتب می گردد. به منظور بررسی عملکرد روش فوق از دو پایگاه داده ۱۰-CIFAR و MNIST و منحنی precision-recall استفاده شده است. نتایج آزمایشات نشان می دهد روش فوق نسبت به روش های مشابه از نظر طول کد باینری، حافظه و هزینه محاسباتی کارا تر می باشد.

واژگان کلیدی: بازیابی تصاویر بر اساس محتوا، مقیاس بالا، الگوی دودویی محلی، الگوریتم هشینگ چندی سازی تکرار شونده

فهرست مطالب

ذ	فهرست تصاویر
ز	فهرست جداول
ژ	فهرست الگوریتم‌ها
۱	فصل ۱:
۱-۱	۱-۱ مقدمه
۱-۲	۲-۱ بازیابی تصویر
۱-۳	۱-۲-۱ محدودیت های روش بازیابی تصویر بر اساس متن
۱-۴	۳-۱ ساختار سیستم های بازیابی تصویر بر اساس محتوا
۱-۵	۱-۳-۱ تصویر درخواستی
۱-۶	۱-۳-۲ استخراج ویژگی ها
۱-۷	۱-۳-۳ پایگاه داده تصاویر
۱-۸	۱-۳-۴ پایگاه داده ویژگی ها
۱-۹	۱-۳-۵ جست و جو شباهت
۱-۱۰	۱-۳-۶ اندیس گذاری و بازیابی
۱-۱۱	۴-۱ بازیابی تصویر در مقیاس بالا
۱-۱۲	۵-۱ کاربردهای سیستم های بازیابی تصویر بر اساس محتوا
۱-۱۳	۶-۱ اهداف پایان نامه
۱-۱۴	۷-۱ ساختار پایان نامه
۸	فصل ۲:
۸-۱	۱-۲ مقدمه

۲-۲	نمایش تصویر	۸
۲-۲-۱	ویژگی های محلی	۱۰
۲-۲-۲	انبان کلمات بصری	۱۰
۲-۲-۳	بردار فیشر	۱۳
۲-۲-۴	بردار ویژگی های تجمع شده محلی	۱۳
۳-۲	جستجوی نزدیکترین همسایه	۱۳
۲-۳-۱	جستجوی نزدیکترین همسایه ی دقیق	۱۴
۲-۳-۲	جستجوی نزدیکترین همسایه ی تقریبی	۱۵
۲-۳-۳	هشینگ حساس به همسایگی (ال اس اچ)	۱۶
۲-۳-۴	نزدیکترین همسایه ی C- تقریبی تصادفی شده	۱۶
۲-۳-۵	رویکرد تقسیم و تخسیر	۱۷
۲-۳-۶	درخت های کی دی	۱۷

فصل ۳:

۲۰	
۳-۱	مقدمه
۳-۲	الگوی دودویی محلی
۳-۲-۱	تغییرات اخیر الگوی دودویی محلی
۳-۳	ال اس اچ
۳-۳-۱	هش
۳-۳-۲	هش حساس به همسایگی
۳-۳-۳	حل مسأله ی نزدیک ترین همسایه ی تقریبی
۳-۳-۴	نقش نگاشت های تصادفی در ال اس اچ
۳-۳-۵	پیاده سازی هش
۳-۳-۶	بررسی عملکرد
۳-۴	کرنل های ثابت به جابه جایی
۳-۵	چندی سازی تکرار شونده
۳-۵-۱	تکنیک آنالیز اجزای اصلی
۳-۵-۲	تحلیل همبستگی کانونیک
۳-۵-۳	تجزیه مقادیر تکین

۳۸	۳-۵-۴ الگوریتم
۴۱	۳-۶ مرتب سازی نتایج بر اساس اهمیت هر بیت
۴۴	۳-۶-۱ محاسبه وزن
۴۶	فصل ۴:
۴۶	۴-۱ مقدمه
۴۶	۴-۲ روش پیشنهادی
۴۶	۴-۲-۱ مرحله اول، استخراج ویژگی های تصویر
۴۷	۴-۲-۲ مرحله دوم، تولید کد باینری فشرده
۴۸	۴-۲-۳ مرحله سوم، به دست آوردن نزدیک ترین همسایه
۴۸	۴-۲-۴ مرحله ی چهارم، مرتب سازی مجدد بر اساس اهمیت هر بیت
۵۰	فصل ۵:
۵۰	۵-۱ مقدمه
۵۰	۵-۲ مجموعه دادگان
۵۲	۵-۳ معیارهای ارزیابی
۵۲	۵-۳-۱ دقت
۵۲	۵-۳-۲ فراخوانی
۵۲	۵-۳-۳ منحنی دقت در مقابل فراخوانی
۵۲	۵-۴ شرح آزمایشات انجام شده
۵۲	۵-۴-۱ تعیین پارامترهای مساله
۵۴	۵-۴-۲ تأثیر ویژگی الگوی دودویی محلی در بهبود عملکرد الگوریتم
۵۵	۵-۴-۳ مقایسه ی روش پیشنهادی با روش های دیگر
۵۶	۵-۴-۴ نمایش خروجی به کاربر
۵۸	۵-۵ نتیجه گیری
۶۰	۵-۶ کارهای آینده
۶۲	مراجع
۶۷	واژه نامه انگلیسی به فارسی

فهرست تصاویر

- ۱-۱ سیستم بازیابی تصویر بر اساس متن [۲۳]. ۲
- ۲-۱ سیستم بازیابی تصویر بر اساس محتوا [۲۳]. ۲
- ۳-۱ ساختار سیستم بازیابی تصویر بر اساس محتوا. ۴
- ۱-۲ نمونه ای از هیستوگرام تصویر ۹
- ۲-۲ استخراج ویژگی های تصویر. ۱۱
- ۳-۲ یادگیری واژه نامه بصری. ۱۱
- ۴-۲ کوانتیزه کردن بردار ویژگی با استفاده از واژه نامه بصری. ۱۲
- ۵-۲ بازنمایی تصاویر با استفاده از تعداد تکرار کلمات بصری. ۱۲
- ۱-۳ عملگر الگوهای دودویی محلی [۲۹]. ۲۰
- ۲-۳ عملگر LBP با شعاع ها و تعداد همسایگی های مختلف [۱۴] ۲۱
- ۳-۳ الگوی یکنواخت در عملگر LBP با شعاع ۱ و تعداد همسایگی ۸ [۳۰]. ۲۳
- ۴-۳ مثال هایی از مفهوم ریزالگوها در LBP ۲۴
- ۵-۳ نگاشت دودویی دوبعدی LSH. ۲۷
- ۶-۳ انتخاب محورهای جدید برای داده های دو بعدی. ۳۴
- ۷-۳ تصویری از ضرایب افزونگی حاصل از تحلیل همبستگی کانونی ۳۶
- ۸-۳ تصویری از عملکرد الگوریتم کواتیزاسیون تکرارشونده [۱۰]. ۴۰
- ۹-۳ مرتب سازی مجدد نتایج بر اساس اهمیت هر بیت [۸]. ۴۳
- ۱-۴ دیاگرام روش پیشنهادی. ۴۷
- ۱-۵ نمونه ای از تصاویر موجود در پایگاه داده CIFAR-۱۰ ۵۱
- ۲-۵ نمونه ای از تصاویر موجود در پایگاه داده MNIST. ۵۱

- ۳-۵ نمودار دقت بازیابی بر اساس مقادیر مختلف m و ε برای مجموعه داده‌گان MNIST . . ۵۳
- ۴-۵ نمودار دقت بازیابی بر اساس مقادیر مختلف m و ε برای مجموعه داده‌گان CIFAR-۱۰ ۵۴
- ۵-۵ منحنی دقت در مقابل فراخوانی به ازای تعداد تکرار مختلف الگوریتم هشینگ ITQ. . . ۵۵
- ۶-۵ تاثیر الگوی دودویی محلی به عنوان ورودی به الگوریتم هشینگ ITQ ۵۶
- ۷-۵ مقایسه عملکرد الگوریتم پیشنهادی با روش های دیگر در مجموعه داده‌گان CIFAR-۱۰ ۵۸
- ۸-۵ مقایسه عملکرد الگوریتم پیشنهادی با روش های دیگر در مجموعه داده‌گان MNIST . . ۵۹

فهرست جداول

۲۴	۱-۳ لیستی از آخرین تغییرات الگوی دودویی محلی
۴۲	۲-۳ مثالی از k بیت کد باینری تصویر درخواستی با پنج تصویر
۵۷	۱-۵ ۲۵ تصویر نخست بازیابی شده

فهرست الگوریتم‌ها

۳-۱ الگوریتم محاسبه وزن هر بیت در مرتب سازی مجدد بر اساس اهمیت هر بیت. ۴۵

فصل ۱

۱-۱ مقدمه

در این فصل پیش زمینه ای از مفهوم بازیابی تصاویر و کاربردهای آن و هدف از انجام این کار ارائه خواهد شد.

۱-۲ بازیابی تصویر

بازیابی تصویر ، فرآیند جست و جو و بازیابی تصاویر در یک پایگاه داده تصاویری می باشد . بازیابی تصاویر به دو دسته ی کلی تقسیم می شود [۴۲] :

- بازیابی تصویر بر اساس متن^۱

- بازیابی تصویر بر اساس محتوا^۲

در بازیابی تصویر بر اساس متن (شکل ۱-۱)، یک کلمه کلیدی^۳ ، زیرنویس^۴ یا یک نوشته توصیفی^۵ به عنوان داده ماوراء به هر یک از تصاویر موجود در پایگاه داده اضافه می شود بنابراین جست و جو روی این داده ها صورت می گیرد. در این روش کلمه کلیدی یا عبارت به عنوان درخواست ارسال می گردد سپس تصاویری که داده ی ماورای آن مرتبط با درخواست کاربر است نمایش داده می شود . این داده ها اکثرا به صورت دستی

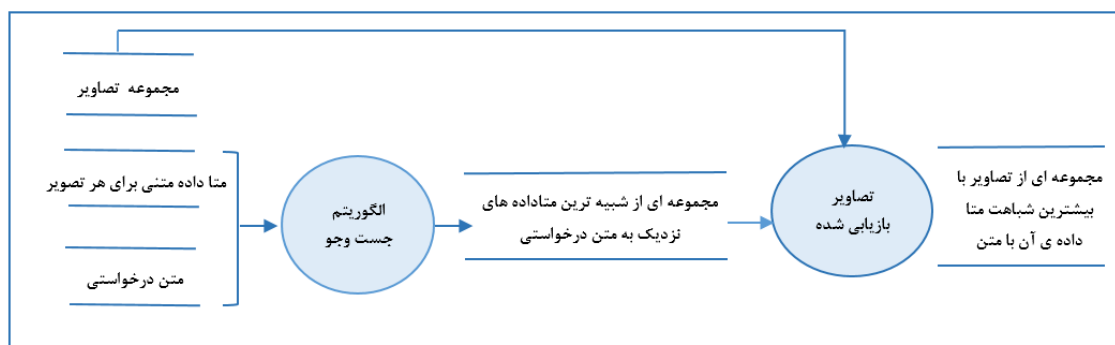
^۱concept based image retrieval

^۲Content based image retrieval

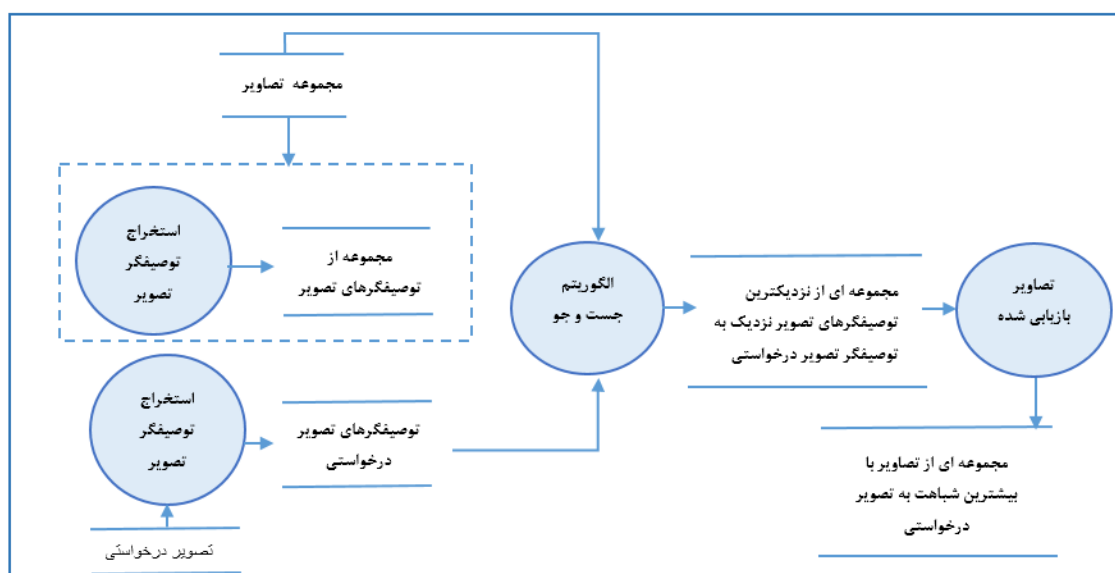
^۳Key word

^۴caption

^۵Description text



شکل ۱-۱: سیستم بازیابی تصویر بر اساس متن [۲۳].



شکل ۱-۲: سیستم بازیابی تصویر بر اساس محتوا [۲۳].

به تصاویر الحاق می‌گردد که این یکی از محدودیت‌های این روش است. در بازیابی تصاویر بر اساس محتوا (شکل ۱-۲)، از داده‌های متنی در فرآیند بازیابی تصاویر استفاده نمی‌شود در این روش، بازیابی براساس میزان شباهت مفاهیم تصاویر به تصویر درخواستی ارسال شده توسط کاربر صورت می‌گیرد. این مفاهیم شامل رنگ، بافت، شکل و یا هر اطلاعات دیگری که از تصویر استخراج می‌گردد می‌باشد بنابراین مفاهیم تصویر باید تجزیه گردد و با یک شیوه مناسب به منظور فرآیند جست و جو نمایش داده شود.

۱-۲-۱ محدودیت های روش بازیابی تصویر بر اساس متن

امروزه اکثر موتورهای جست و جوی تصاویر تحت وب، از الگوریتم های جست و جو بر پایه فراداده استفاده می کنند. مشکل این روش آن است که تعداد زیادی از نتایج، مطلوب ما نمی باشند دلایل این امر:

- زبان: ابهامات؛ مانند چندمعنایی^۶ یا مترادف^۷
- داده های متنی: اگرچه فرآیند جست و جو با تکنولوژی های موجود آسان است اما برای پایگاه داده های بزرگ و یا تصاویری که به صورت خودکار (تصاویر دوربین های مدار بسته) تولید می شوند این روش ها غیر عملی می باشد زیرا نیاز به افرادی است که تصاویر موجود در پایگاه داده را در قالب کلمات توصیف نمایند.
- برجسب گذاری انسان: علاوه بر پرهزینه بودن و ناکارآمد بودن این روش ها برای پایگاه داده های بزرگ می توان به خطا در برجسب گذاری نیز اشاره نمود. از این خطاها می توان به غلط املایی یا برجسب گذاری نادرست اشاره کرد. کاربران متفاوت برجسب های متفاوتی را به تصویر الحاق می کنند و ممکن است تمام مفاهیم موجود در تصویر را برجسب گذاری نکنند.

۱-۳-۳ ساختار سیستم های بازیابی تصویر بر اساس محتوا

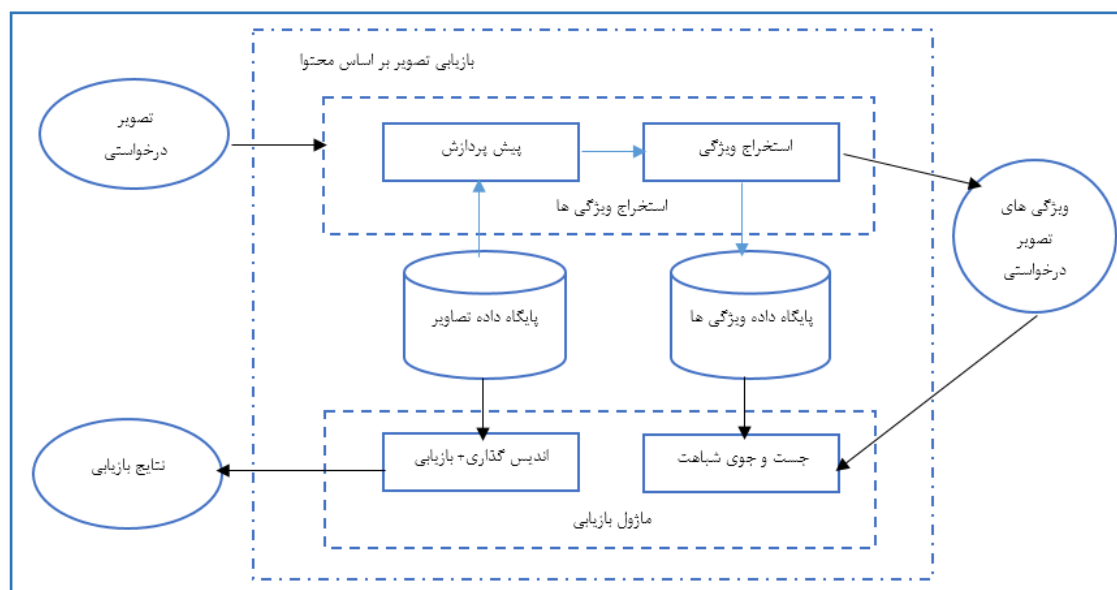
شکل ۱-۳ سیستم بازیابی تصویر بر اساس محتوا را نمایش می دهد فرآیند آموزش این سیستم به صورت آفلاین انجام می شود و با پیکان های آبی نمایش داده شده است و فرآیند درخواست کاربر به صورت آنلاین انجام می گردد و با پیکان های مشکی نمایش داده شده است.

۱-۳-۱ تصویر درخواستی

این بخش، ورودی (تصویری است که به عنوان درخواست برای عملیات جست و جو گرفته شده است) را برای سیستم توصیف می کند. کاربر با انتخاب یک تصویر به دنبال یافتن تصاویر مشابه به آن در پایگاه داده می باشد.

^۶polysemy

^۷synonymy



شکل ۱-۳: ساختار سیستم بازیابی تصویر بر اساس محتوا.

۱-۳-۲ استخراج ویژگی ها

این بخش یکی از اجزای مهم این سیستم می باشد. در این بخش ویژگی های تصویر با هدف نمایش مفاهیم تصویر در قالب عددی استخراج می شود.

۱-۳-۳ پایگاه داده تصاویر

شامل مجموعه ای از تصاویر می باشد که در این سیستم نتایج درخواست کاربر را فراهم می کند.

۱-۳-۴ پایگاه داده ویژگی ها

مکانی است که کلیه اطلاعات استخراج شده از تصاویر آموزشی در آن ذخیره می شود. بردار ویژگی های کلیه تصاویر آموزشی استخراج و ذخیره می گردد. سیستم از این داده ها به منظور مقایسه ویژگی ها در فرآیند جست و جو استفاده می کند. ویژگی ها ممکن است در قالب یک لیست یا بردارهای ویژگی و یا ساختمان داده های پیچیده تری مانند درخت و... ذخیره گردد.

۵-۳-۱ جست و جو شباهت

اساس این جزء، مقایسه می باشد و نتایج را بر مبنای تکنیکی که به منظور مقایسه استفاده شده است برمی گرداند. سیستم میزان شباهت را برای هر یک از تصاویر موجود در پایگاه داده فراهم میکند. این فرآیند بر اساس معیارهای شباهت که به منظور محاسبه اختلاف میان توصیفگرهای تصویر درخواستی و پایگاه داده ویژگی ها می باشد انجام می گیرد. از دید سیستم، شباهت دو تصویر وابسته به اختلاف در فضای ویژگی است. بنابراین زمانی که اختلاف میان توصیفگرهای تصویر کم باشد شباهت آن دو بیشتر می باشد.

۶-۳-۱ اندیس گذاری و بازیابی

در این بخش سیستم تعدادی از تصاویر را به عنوان نتیجه انتخاب و به کاربر نمایش می دهد.

۴-۱ بازیابی تصویر در مقیاس بالا

دیتاست های تصویر در مقیاس بالا شامل هزاران، میلیون ها یا بیشتر تصویر می باشند. نمونه هایی از این دیتاست ها صفحات وب^۸، تصاویر پزشکی^۹، فریم های ویدئو استخراج شده از دوربین های کاوشگر و... است. جست و جوی تصاویر در مقیاس بالا با استفاده از تکنیک های تطابق ناشیانه^{۱۰} (تطابق تصویر درخواستی با تمامی تصاویر موجود در دیتاست) یک روش عملی^{۱۱} نمی باشد. این روش در صورتی که پایگاه داده شامل چند صد تصویر باشد مناسب است دلایل رد روش فوق به شرح ذیل می باشد:

- تعداد بالای تصاویر موجود در دیتاست.
- ابعاد بالای توصیفگرهای تصویر که معمولاً شامل ده ها یا صدها بعد برای هر یک از تصاویر می باشد.
- تکنیکی که به منظور تطابق جفت تصویرها مورد استفاده قرار می گیرد.
- پیچیدگی معیار فاصله ای که به منظور تطابق توصیفگرها مورد استفاده قرار می گیرد که یک عامل اصلی در افزایش سرعت و دقت بازیابی محسوب می گردد.

^۸www^۹Medical image^{۱۰}Brute force matching^{۱۱}practical

۵-۱ کاربردهای سیستم های بازیابی تصویر بر اساس محتوا

بازیابی تصاویر بر اساس محتوا مولد بسیاری از برنامه های کاربردی در دنیای واقعی است. این برنامه ها شامل شناسایی اشیاء^{۱۲}، پیدا کردن کپی های غیر مجاز تصاویر در صفحات وب^{۱۳}، جست و جو کردن فریم های ویدئو^{۱۴}، دسته بندی کردن تصاویر^{۱۵}، تصاویر پزشکی می باشد. کاربردهای دیگر آن شامل گالری های هنری و مدیریت موزه ها، مهندسی معماری، طراحی داخلی، مدیریت منابع زمین، سیستم های اطلاعات جغرافیایی، مدیریت پایگاه داده های علمی، پیش بینی آب و هوا، طراحی مد، مدیریت پایگاه داده های تجاری، اجرای قانون و پیشگیری جرم، آرشیو تصاویر و سیستم های اجتماعی و سرگرمی های خانگی می باشد. یکی از کاربردهای جذاب بازیابی تصاویر بر اساس محتوا، موتور جست و جوی تصاویر آنلاین می باشد به عنوان نمونه ای از این موتورهای جست و جو تصاویر می توان به Image Reverse images، TinyEye google Search اشاره نمود.

۶-۱ اهداف پایان نامه

اهداف این پایان نامه به شرح ذیل می باشد:

- ارائه نمایی کلی از حوزه ی بازیابی تصویر
- مروری بر آخرین الگوریتم ها و متدهای ارائه شده به منظور رفع مشکلات مطرح شده در حوزه ی بازیابی تصاویر در مقیاس بالا
- ارائه روشی جدید و کارا به منظور بازیابی تصاویر بر اساس محتوا در مقیاس بالا

۷-۱ ساختار پایان نامه

در فصل ۲ به مرور روش هایی که به منظور حل چالش های بازیابی تصویر بر اساس محتوا در مقیاس بالا پرداخته است اشاره دارد. این فصل به بیان روش هایی که به منظور نمایش ویژگی های تصویر و روش هایی که به منظور جست و جوی نزدیک ترین همسایه ارائه شده است پرداخته است.

^{۱۲}Object recognition

^{۱۳}Finding partial image duplicate on the web

^{۱۴}Searching individual video frame

^{۱۵}Image classification

فصل ۳ به تشریح الگوریتم های به کار رفته در این کار پرداخته است.

فصل ۴ به بیان روش پیشنهادی در این پایان نامه اشاره دارد.

فصل ۵ به تشریح آزمایشات انجام شده، نتایج به دست آمده و مقایسه و ارزیابی روش پیشنهادی با روش های دیگر پرداخته است. در انتها به بیان نتیجه گیری و کارهای پیشنهادی برای آینده اشاره دارد.

فصل ۲

۱-۲ مقدمه

این فصل به مرور تحقیقات و کارهای اخیر انجام شده به منظور حل چالش های موجود در حوزه ی بازیابی تصویر بر اساس محتوا در مقیاس بالا می پردازد. از چالش های موجود در این حوزه می توان به مساله ابعاد بالای توصیفگر های تصویر و پیچیدگی و سرعت الگوریتم های جست و جو اشاره نمود که در پایگاه داده های تصاویر با ابعاد بالا منجر به افزایش حافظه ی مورد نیاز و کاهش سرعت بازیابی می گردند. به منظور رفع مشکلات فوق راهکارهایی ارائه گردیده است که منجر به کاهش هزینه ی محاسباتی، کاهش حافظه ی مورد استفاده و افزایش سرعت بازیابی گردد. در بخش اول به بیان الگوریتم های ارائه شده به منظور نمایش ویژگی های تصویر پرداخته شده است و در بخش دوم روش هایی که به منظور جست و جوی تصاویر ارائه شده بیان گردیده است.

۲-۲ نمایش تصویر

متدهای نمایش تصویر قدیمی^۱ از رنگ ، بافت یا شکل تصویر برای توصیف مفاهیم آن استفاده می کند. به عنوان نمونه در سیستم QBIC^۲ از هیستوگرام رنگ تصویر، برای نمایش تصاویر استفاده می شود. هیستوگرام رنگ تعداد وقوع هر رنگ در تصویر را نمایش می دهد (شکل ۱-۲)

این موضوع بسیار بدیهی است که هیستوگرام رنگ، شامل اطلاعاتی درباره ی ساختار تصویر نمی باشد و تنها توزیع رنگ در تصویر را توصیف می کند. متدهای نمایش تصویر که اخیرا ارائه شده است از تکنیک

^۱traditional

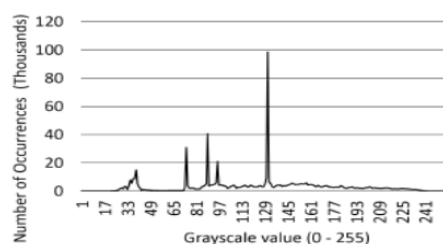
^۲Query By Image Content



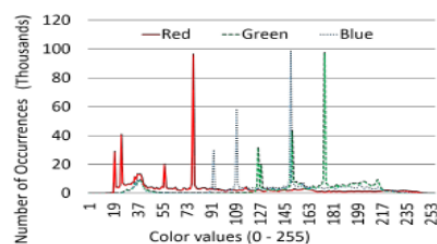
(ب) تصویر در سطح خاکستری



(آ) تصویر رنگی



(د) هیستوگرام سطح خاکستری



(ج) هیستوگرام رنگی

شکل ۲-۱: نمونه ای از هیستوگرام تصویر

های پیشرفته ای مانند توصیفگرهای نقاط کلیدی تصویر^۳ و یا توصیفگرهای نقاط مورد علاقه^۴ استفاده می کنند. نقاط کلیدی تصویر به نقاطی اطلاق می گردد که در جایگاه های خوش تعریف^۵ در تصویر قرار دارد و نسبت به تغییرات سراسری و محلی در حوزه ی مقیاس^۶، چرخش^۷، نور و... ثابت هستند. تعدادی از نقاط مورد علاقه با قابلیت اطمینان و درجه ی بالایی از تکرارپذیری^۸ محاسبه می شوند این نقاط کلیدی به منظور تولید یک نمایش عددی^۹ برپایه ی اطلاعات پیکسل های موجود در ناحیه ای از تصویر پیرامون نقطه کلیدی استفاده می شود. نتیجه ی این نمایش عددی بردارهای توصیفگر نقاط کلیدی^{۱۰} یا بردار ویژگی^{۱۱} نامیده می شود. تجزیه کردن یک تصویر به ویژگی ها^{۱۲} یا نقاط کلیدی محلی مورد علاقه^{۱۳} تکنیکی است که به

^۳ Image key point descriptor

^۴ Interest point descriptors

^۵ Well defined position

^۶ scaling

^۷ rotate

^۸ reproducibility

^۹ Numerical representation

^{۱۰} Keypoint descriptor vecotr

^{۱۱} Feature vector

^{۱۲} features

^{۱۳} Local keypoint of interest

صورت گسترده برای نمایش تصاویر دیجیتال مورد استفاده قرار می گیرد. بسیاری از کاربردهای کامپیوتر شامل آن هایی که مرتبط با بازیابی تصویر براساس محتوا می باشند وابسته به حضور پایدار ویژگی های نماینده یک تصویر می باشند. این کار هدف اصلی شناساگرهای نقاط کلیدی تصویر و توصیفگرهای آن است. هدف شناساگر نقاط کلیدی ایده آل، شناسایی نواحی از تصویر است که در صورت اعمال تغییرات در نور، زاویه دید یا مقیاس، دوباره همان نقاط به عنوان نقاط کلیدی شناسایی گردد. به طور کلی این نواحی تصویر نسبت به تمام جابه جایی های^{۱۴} ممکن یک تصویر مقاوم هستند.

۱-۲-۲ ویژگی های محلی

ویژگی های محلی محتویات بصری نواحی مورد علاقه را توصیف می کند [۲۸]. یک توصیفگر محلی خوب باید جداساز و مقاوم در برابر تغییرات نمایش که منجر به تغییرات شناساگر می گردد باشد ضعف این روش این است که یک تصویر با مجموعه زیادی از توصیفگرها، توصیف می شود و در جستجو برای بازیابی تصاویر مشابه باید حجم زیادی از مقایسات صورت گیرد زیرا توصیفگرهای تصویر درخواستی با توصیفگرهای تصاویر موجود در دیتاست مقایسه می گردد.

۲-۲-۲ انبان کلمات بصری

هدف این روش جایگزینی هر یک از توصیفگرهای محلی یک تصویر با یک کلمه بصری که از لغات^{۱۵} از پیش تعریف شده است. هدف این روش به کارگیری تکنیک های بازیابی متن در بازیابی بر اساس محتوا می باشد.

مراحل اصلی مدل

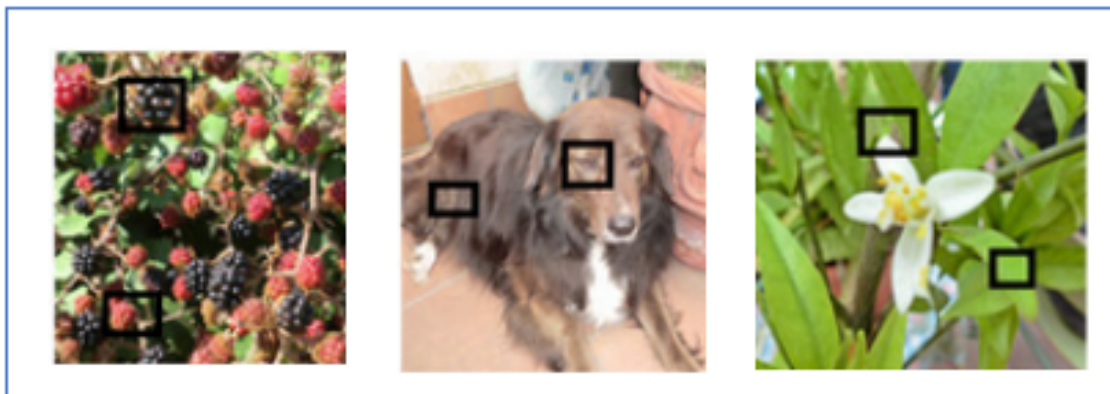
مراحل اصلی در مدل انبان کلمات بصری به صورت زیر است:

استخراج ویژگی

در مرحله استخراج ویژگی باید از قسمت هایی از تصویر به عنوان کلمات بصری نمونه گیری کنیم (شکل ۲-۲).

^{۱۴}transformation

^{۱۵}vocabulary



شکل ۲-۲: استخراج ویژگی های تصویر.



شکل ۲-۳: یادگیری واژه نامه بصری.

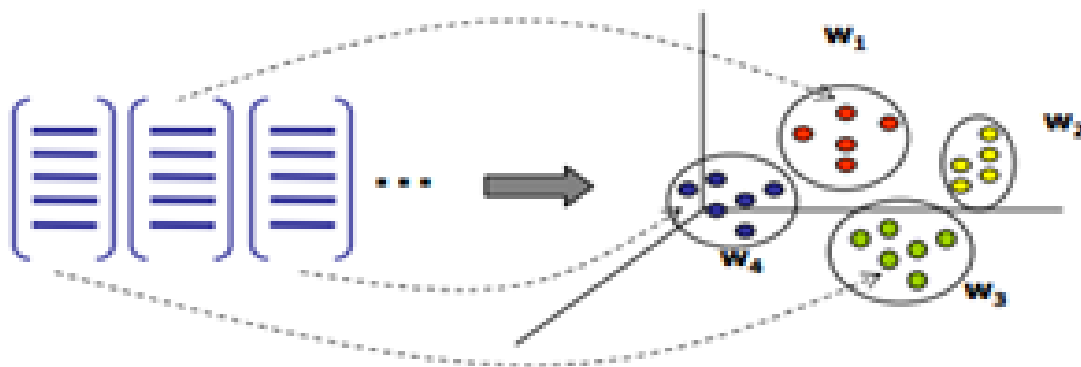
یادگیری واژه نامه بصری

انواع روش های خوشه بندی راهکارهایی هستند که در یادگیری واژه نامه^{۱۶} از آن استفاده می شود و معمولاً نوعی یادگیری بدون ناظر هستند یکی از معمول ترین روش ها برای یادگیری واژه نامه الگوریتم k-means است بعد از اجرای الگوریتم k-means مرکز هر خوشه^{۱۷} به عنوان کلمه کلیدی^{۱۸} در نظر گرفته می شود (شکل ۲-۳).

^{۱۶}codebook

^{۱۷}centroid

^{۱۸}codewords



شکل ۲-۴: کوانتیزه کردن بردار ویژگی با استفاده از واژه نامه بصری.



شکل ۲-۵: بازنمایی تصاویر با استفاده از تعداد تکرار کلمات بصری.

کوانتیزه کردن بردارهای ویژگی با استفاده از واژه نامه بصری

در این مرحله از واژه نامه برای چندی سازی بردارهای ویژگی استفاده می شود چندی ساز برداری یک بردار ویژگی را به عنوان ورودی گرفته و آن را به نزدیک ترین بردار واژه ^{۱۹} درون واژه نامه انتساب می دهد معیار فاصله می تواند اقلیدسی یا هر متریک دیگری باشد (شکل ۲-۴).

بازنمایی تصاویر با استفاده از تعداد تکرار کلمات بصری

در این مرحله ویژگی ها را بر اساس آنچه در مرحله قبل به آن انتساب داده ایم شمارش می کنیم و یک بردار ویژگی جدید یا یک هیستوگرام از کلمات بصری تشکیل می دهیم (شکل ۲-۵).

^{۱۹}codevector

۲-۲-۳ بردار فیشر

کرنل های فیشر چگونگی انحراف مجموعه ای از توصیفگر ها را از توزیع متوسط توصیف می کنند [۱۹]. با استفاده از یک مدل سراسری پارامتریک مدل می شود. کرنل های فیشر به منظور طبقه بندی تصاویر [۲۳] و جستجوی تصاویر در مقیاس بالا مورد استفاده قرار می گیرند. هم چنین اثبات شده است که بردار فیشر برگرفته از روش انبان کلمات بصری است. به طوری که روش انبان کلمات بصری تعداد توصیفگرهایی که به هر ناحیه نسبت داده شده است می شمارد بردار فیشر موقعیت تقریبی توصیفگرها در هر ناحیه را کد می کند. نتایج حاصل از [۲۱] نشان می دهد بهره گیری از بردار فیشر و تکنیک هشینگ دارای عملکرد بهتری نسبت به مدل انبان کلمات بصری است.

۲-۲-۴ بردار ویژگی های تجمع شده محلی

این روش برگرفته از نام *Vector of Locally Aggregated Descriptors* می باشد مانند روش انبان کلمات بصری در ابتدا یک واژه نامه $\{\mu_1, \mu_2, \dots, \mu_k\}$ با استفاده از یک الگوریتم خوشه بندی مثل k-means ایجاد می شود سپس هر یک از توصیفگرهای محلی x_t به نزدیکترین کلمه بصری موجود در واژه نامه نسبت داده می شوند سپس از رابطه زیر:

$$v_i = \sum_{x_t: NN(x_t)=i} x_t - \mu_t \quad (۱-۲)$$

بردار $V = [v_1^T, v_2^T, \dots, v_k^T]$ به دست می آید و برای بهبود کارایی این روش نرمال سازی انجام می شود بعد از انجام این پروسه دو توصیفگر سراسری V_1 و V_2 که مربوط به دو تصویر می باشند می توانند با استفاده از ضرب داخلی با هم مقایسه شوند.

۲-۳ جستجوی نزدیکترین همسایه

مجموعه ای از نقاط به نام P در فضای بعد بالا داریم. می خواهیم ساختمان داده ای بسازیم تا با گرفتن هر نقطه ی ورودی q ، نزدیکترین نقاط به q در P را بدهد. این مسأله جستجوی نزدیکترین همسایه ^{۲۰} نام دارد و اهمیت زیادی در بخش های مختلف علوم کامپیوتر (شامل شناسایی الگو، جستجو در داده های مالی، مדיا، آمار محاسباتی، و داده کاوی) دارد. بسیاری از این کاربردها با مجموعه های بزرگ (مثلاً یک پایگاه داده

^{۲۰} nearest neighbor search

شامل صفحات وب میتواند بیش از یک بلیون متن داشته باشد) سروکار دارند. علاوه بر این، بعد نقطه ها نیز می تواند بسیار بزرگ باشد (مثلاً چند صد یا حتی چند هزار بعد). بنابراین، طراحی الگوریتم هایی که بتوانند با اندازه ی پایگاه داده و تعداد ابعاد داده به خوبی مقیاس پیدا کنند، بسیار پراهمیت است. مسأله ی نزدیکترین همسایه یک نمونه از دسته ی بزرگتری از مسائل به نام مسائل مجاورت^{۲۱} می باشد. این ها مسائلی هستند که تعریفشان وابسته به فاصله ی بین نقاط ورودی می باشد. ابتدا مسأله ی جستجوی نزدیکترین همسایه ی دقیق^{۲۲} را تعریف می کنیم.

۲-۳-۱ جستجوی نزدیکترین همسایه ی دقیق

با داشتن مجموعه ی P از نقاط در فضای R^d به دنبال ساختمان داده ای می باشیم که با گرفتن هر نقطه ی ورودی q ، نقطه ای در P که نزدیکترین فاصله را تا q دارد پیدا کند. مسأله کامل تعریف نشده است مگر اینکه فاصله ی بین هر دو نقطه ی p و q تعریف شود. معمولاً فرض می شود که فاصله توسط نرم L_s ایجاد می شود. یعنی فاصله ی p و q به صورت $\|p - q\|_s$ است که در آن $\|X\| = (\sum_{i=1}^d [x_i]^s)^{1/s}$ است. تعریف دیگر و کلی تر از فاصله نیز امکان پذیر است. یک راه حل ساده و پیش پا افتاده برای این مسأله به شرح زیر است: با گرفتن نقطه ی ورودی q ، فاصله ی q تا همه ی نقاط در P را حساب کن و نقطه ای که کمترین فاصله را داشت اعلام کن. این روش که بررسی خطی^{۲۳} نام دارد، پیچیدگی زمانی $\Theta(dn)$ دارد. این برای مجموعه داده های کوچک قابل قبول می باشد ولی به شدت برای مجموعه داده های بزرگتر غیر کاراست. هدف تحقیقات در این حوزه این است که الگوریتمی طراحی شود که این مسأله را در زمان زیرخطی (یا حتی لگاریتمی) حل کند. مسائل مجاورت جزء اولین مسائلی بودند که در هندسه ی محاسباتی مورد بررسی قرار گرفتند. نتیجه ی این تحقیقات این بوده که تعداد زیادی راه حل های کارا برای حالتی که نقاط در فضایی با بعد ثابت قرار داشته باشند ارائه شده است. به عنوان مثال، مسأله ی نزدیکترین همسایه برای هر نقطه ی ورودی با $o(\log n)$ زمان و $o(n)$ حافظه قابل حل است. نتایج مشابه برای مسائل دیگر نیز به دست آمده است. متأسفانه با رشد بعد، این الگوریتم ها کارایی خود را از دست می دهند. به طور دقیق تر، پیچیدگی زمانی و حافظه ای آنها به صورت نمایی نسبت به بعد افزایش می یابد. به خصوص، مسأله ی نزدیکترین همسایه راه حلی به زمان $o(d \log n)$ برای هر نقطه ورودی ولی با $n^{o(d)}$ حافظه دارد. اگر بخواهیم پیچیدگی حافظه را خطی یا نزدیک خطی کنیم، بهترین کران بالایی به دست آمده برای پیچیدگی زمانی برای ورودی تصادفی $\min(2^{o(d)}, dn)$ است، که حتی برای d هایی که خیلی بزرگ نباشند تقریباً نسبت به n خطی است. علاوه بر

^{۲۱}proximity problems

^{۲۲}exact nearest neighbor search

^{۲۳}linear scan

این، وابستگی نمایی حافظه و/یا زمان به بعد (که مصیبت بعد^{۲۴} نام دارد) از دیدگاه کاربردی نیز بررسی شده است. به خصوص، دانسته شده است که بسیاری از ساختمان داده های پرکاربرد (که از حافظه های خطی یا نزدیک خطی استفاده می کنند) زمانی که بعد از یک مقدار آستانه ای بیشتر می شود (معمولاً بین ۱۰ الی ۱۲ که به تعداد نقاط داده بستگی دارد)، از خود پیچیدگی زمانی خطی نسبت به نشان می دهند. عدم موفقیت در از بین بردن وابستگی نمایی به بعد باعث شده که بسیاری از محققان ادعا کنند که هیچ پاسخ کارایی برای این دسته از مسائل زمانی که بعد به اندازه ی کافی بزرگ باشد وجود ندارد. در همین راستا، سؤال مطرح می شود: در صورتی که اجازه بدهیم پاسخ ها تقریبی باشند، آیا ممکن است که وابستگی نمایی از بین برود؟ مفهوم تقریب زدن برای جستجوی نزدیکترین همسایه به خوبی تعریف میشود: به جای اینکه خروجی را نقطه p که از همه به q نزدیک تر است اعلام کند، الگوریتم اجازه دارد که هر نقطه ای که در فاصله ی $(1 + \epsilon)$ برابر فاصله ی p تا q قرار دارد را اعلام کند. تعریف های مشابهی قابل اعمال به مسائل مجاورت دیگر نیز می باشد. توجه کنید که این رویکرد مشابه طراحی الگوریتم های تقریبی کارا برای مسائل ان پی هارد می باشد. در سالهای اخیر، تعدادی از محققان نشان دادند که در واقع در بسیاری از حالات، تقریب باعث کاهش وابستگی به بعد از نمایی به چندجمله ای می شود.

۲-۳-۲ جستجوی نزدیکترین همسایه ی تقریبی

تقریباً همه ی الگوریتم های مربوط به مسائل مجاورت در فضاهایی با بعد بالا ابتدا مسئله را به مسئله ی یافتن نزدیک ترین همسایه های تقریبی کاهش می دهند. بنابراین ما نیز با تعریف این مسئله شروع می کنیم.

جستجوی نزدیک ترین همسایه ی تقریبی یا (r, c) -NN

با داشتن مجموعه ی P از n نقطه در فضای متری، $M=(X, D)$ به دنبال طراحی ساختمان داده ای می باشیم که بتواند عمل زیر را انجام دهد: برای هر نقطه ی ورودی $q \in X$ ، اگر $p \in P$ وجود داشته باشد به قسمی که $D(p, q) \leq r$ باشد، نقطه ی $p \in P$ را پیدا کنید طوری که $D(p, q) \leq cr$ باشد. معیار متری همینگ: یک معیار متری (Σ, D) که در آن Σ مجموعه ای از نمادهاست و برای هر $(p, q) \in \Sigma$ ، $D(p, q)$ برابر با تعداد $i \in 1, \dots, d$ است به طوری که $p_i \neq q_i$ باشد.

^{۲۴}curse of dimensionality

رویکرد تصویر کردن تصادفی

اولین الگوریتم ها برای (r, c) -NN در ابعاد بالا به کمک تکنیک تصویر کردن تصادفی به دست آمده اند. این تکنیک زمانی کاربرد دارد که معیار متری D توسط نرم L_p ایجاد شده باشد، زمانی که $p \in [0, 2]$ باشد. ابتدا به حالتی می پردازیم که تمام نقاط و نقطه ی ورودی بردارهای دودویی از $\{0, 1\}^d$ باشند و D فاصله ی همینگ باشد (یعنی معیار متری ای که توسط نرم L_1 ایجاد شود).

کاهش بعد

یک روش تصادفی شده می باشد که بعد فضای همینگ را از d به $k = O(\log n / \epsilon^2)$ کاهش دهد به طوری که با $(1 + \epsilon)$ ضریب فاصله های بین نقطه های داده و نقطه ی ورودی را حفظ می کند. این روش به ما اجازه می دهد تا مساله (r, c) -NN را در فضای d بعدی به مساله ی (r, c) -NN در فضای k بعدی کاهش دهیم. رویکرد مشابهی را می توان برای حل مساله نزدیک ترین همسایه در فضای اقلیدسی به کار برد. به خصوص حل مسئله ی $(r, c + \epsilon) - NN$ در L_2^d به کمک $n(\frac{1}{d})^{O(d)}$ حافظه قابل انجام است.

۲-۳-۳. هشینگ حساس به همسایگی (ال اس اچ)

ال اس اچ یک الگوریتم تصادفی شده است. از خاصیت تصادفی در ساخت ساختمان داده استفاده می شود. علاوه بر این، الگوریتم برای حل مسئله ی همسایه ی نزدیک^{۲۵} به کار می رود نه مسئله ی نزدیکترین همسایه. مسئله ی همسایه ی نزدیک نسخه ی تصمیم مسئله ی نزدیکترین همسایه است. تعریف مسئله به صورت زیر می باشد. در این تعاریف نقطه ی p ، همسایه ی R نزدیک q می باشد. اگر فاصله از p به q حداکثر R باشد.

۲-۳-۴. نزدیکترین همسایه ی c -تقریبی تصادفی شده

با داشتن مجموعه ی P از نقاط در فضای R^d و پارامترهای $0 < R < \delta$ ، ساختمان داده ای بساز که با گرفتن هر نقطه ی ورودی q ، هر همسایه ی R نزدیک q در p را با احتمال $1 - \delta$ اعلام کند. توجه کنید که در تعریف دوم، احتمال های اعلام همسایه های R نزدیک مختلف ممکن است از هم مستقل نباشد. علاوه بر این، توجه کنید که در هر دو مسئله مقدار R در زمان پیش پردازش دانسته شده است. بنابراین، با مقیاس کردن مختصات همه ی نقاط می توانیم فرض کنیم که $R=1$ است. ایده ی اصلی الگوریتم ال اس اچ این است که

^{۲۵}decision version

نقاط داده را با استفاده از چند تابع هش طوری هش کند که برای هر تابع، احتمال برخورد^{۲۶} برای نقاطی که به هم نزدیک ترند بسیار بالاتر از نقاطی باشد که از هم دورترند. در این صورت می توان همسایه های نزدیک را با هش کردن نقطه ی ورودی q و بازیابی نقاط موجود در جعبه ی شامل q به دست آورد.

۲-۳-۵ رویکرد تقسیم و تخسیر

تکنیک های کاهش بعد و ال اس اچ محدودیت هایی دارند. به خصوص، نمی توان برای حل مسئله ی نزدیک ترین همسایه ی تحت نرم از آن ها استفاده کرد. خوشبختانه، این نرم خاصیت های جالب دیگری دارد که طراحی ساختمان داده برای نزدیک ترین همسایه ی تقریبی را میسر می سازد. تنها الگوریتم شناخته شده برای حل (r, c) -NN تحت نرم L_∞^d دارای پارامترهای زیر برای هر $\rho > 0$ می باشد:

• ضریب تقریب: $c = O(4 \lceil \log(1 + \rho) \log 4d \rceil)$ اگر $\rho = \log c$ آن گاه $c=3$

• حافظه: $dn^{1+\rho}$

• زمان: $O(d \log n)$

ایده ی اصلی الگوریتم استفاده از رویکرد تقسیم و تخسیر است. به خصوص، ابرصفحه های H که شامل همه ی نقاطی که یک مختصشان (مثلاً i امین) با هم برابر باشد را در نظر بگیرید. الگوریتم سعی می کند ابرصفحه ی ای پیدا کند که دارای این خاصیت باشد که مجموعه ی نقاط $P_L \subset P$ که در سمت چپ آن و در فاصله ی $r \geq$ از آن قرار دارد خیلی کوچک تر از مجموعه ی P_M از نقاطی باشد که در فاصله ی r از آن قرار دارد. شرایط مشابهی باید برای مجموعه ی P_R از نقطه ها در طرف راست H نیز برقرار باشد. اگر چنین H ای وجود داشته باشد، P را به دو مجموعه ی $P_1 = P_L \cup P_M$ و $P_2 = P/P_L$ تقسیم می کنیم و ساختمان داده را به صورت بازگشتی روی p_1 و P_2 می سازیم. به راحتی می توان دید که هنگام پردازش نقطه ی ورودی q ، وابسته به این که q در کدام سمت H قرار داشته باشد، الگوریتم را به صورت بازگشتی روی P_1 و P_2 ادامه می دهیم. علاوه بر این، می توان اثبات کرد که افزایش در استفاده از حافظه که ناشی از تکرار P_M می باشد خیلی زیاد نیست.

۲-۳-۶ درخت های کی دی

درخت کی دی ساختمان داده ای است که توسط جان بنتلی در سال ۱۹۷۵ اختراع شد. با آن که این ساختمان داده بسیار قدیمی است ولی هم چنان جزء ساختمان داده های پرکاربرد برای جستجو در فضاهای چندبعدی

^{۲۶}collision

^{۲۷} به خصوص جستجو در حافظه ی اصلی به شمار می رود.

با داشتن مجموعه ای از n نقطه در فضای d بعدی، درخت کی دی به صورت بازگشتی به شرح زیر ساخته می شود. ابتدا، میانه ی مقادیر مختص i ام نقاط را پیدا میکنیم (از $i=1$ شروع می کنیم). یعنی مقدار M ای محاسبه میشود که ۵۰ درصد نقاط دارای مختص i ام بزرگتر یا مساوی M بوده و ۵۰ درصد دیگر نقطه ها دارای مختص i ام کوچک تر از M باشند. مقدار x ذخیره می شود و مجموعه ی P به P_L و P_R افراز می شود به قسمی که P_L فقط شامل نقاطی می شود که مختص i امشان کوچک تر یا مساوی M باشد و $|P_R| = |P_L| \pm 1$. این پروسه به صورت بازگشتی روی P_L و P_R تکرار می شود با این تفاوت که i با $i=1$ جایگزین می شود (یا ۱ اگر $d=i$ باشد). زمانی که مجموعه ی نقاط یک نود اندازه ی یک داشته باشد، الگوریتم بازگشتی خاتمه می یابد. ساختمان داده حاصل یک درخت دودویی است با n برگ و عمق $\lceil \log n \rceil$. به خصوص برای $d=1$ ، یک درخت جستجو دودویی متوازن به دست می آوریم. با توجه به این که میانه ی n نقطه در زمان $O(n)$ به دست می آید، کل ساختمان داده در زمان $O(n \log n)$ قابل ساخت است. چند راه حل برای یافتن نزدیک ترین همسایه های نقطه ی ورودی q در مجموعه نقاط P وجود دارد. اولین راه حل در همان مقاله ی سال ۱۹۷۵ ارائه شده است. یک راه حل را ما در اینجا توضیح می دهیم. جستجو از ریشه ی درخت آغاز می شود و بازگشتی است. در هر زمان، الگوریتم فاصله ی R به نزدیک ترین نقطه به q را نگه می دارد، در ابتدا $R=\infty$ است. در نود برگ (که مثلاً شامل نقطه ی $\hat{P}$ است) الگوریتم چک می کند که $\|q - \hat{p}\| < R$ باشد. اگر این طور بود مقدار R ، $\|q - \hat{p}\|$ قرار داده می شود و $\hat{p}$ به عنوان کاندید نزدیک ترین همسایه ذخیره می شود. برای یک نود داخلی (غیر برگ) الگوریتم به ترتیب زیر عمل می کند فرض کنید M مقدار میانه ذخیره شده در نود باشد که برای مختص های i ام حساب شده است. الگوریتم چک می کند که مختص i ام کوچک تر یا مساوی M است اگر کوچک تر بود، الگوریتم روی نود چپ ادامه می دهد، در غیر این صورت روی نود راست ادامه می دهد. بعد از اتمام الگوریتم بازگشتی، الگوریتم تستی را انجام می دهد یعنی چک می کند که اگر کره ای با شعاع R شامل نقطه ای در R^d می باشد که مختص i ام آن در طرف مخالف M نسبت به q باشد. اگر این طور بود، الگوریتم بازگشتی روی فرزند (سرزده ی) نود فعلی ادامه پیدا می کند. در غیر این صورت الگوریتم پایان می یابد. در پایان، الگوریتم آخرین کاندید نزدیک ترین همسایه را به عنوان جواب اعلام می کند.

نشان داده شده است اگر داده و نقطه ی ورودی به صورت مستقل و تصادفی از یک توزیع تصادفی انتخاب شوند، آن گاه الگوریتم در زمان $G(d) \log n$ برای یک تابع G انجام می شود. ولی $G(d)$ نسبت به d نمایی است (برای d هایی که نسبت به $\log n$ کوچک هستند) ولی در بدترین حالت پیچیدگی زمانی $O(dn)$ است. زمانی که بعد فضای جستجو از چند بعد بیشتر شود مجبور می شویم که تقریباً همه ی نودهای مجموعه داده را

^{۲۷} multidimensional spaces

بررسی کنیم.

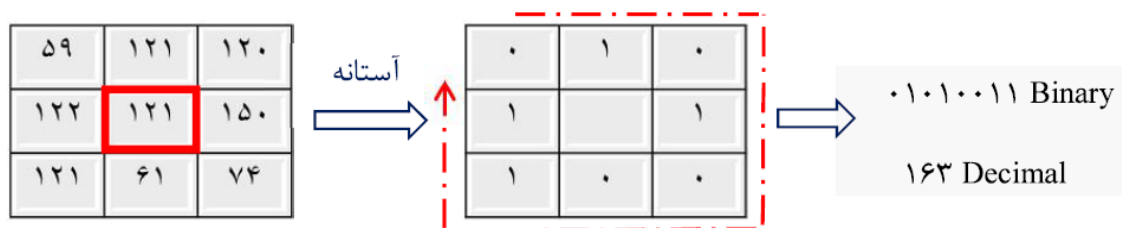
فصل ۳

۱-۳ مقدمه

در این فصل به بیان پیشینه ی الگوریتم های به کار رفته در این پایان نامه اشاره می گردد و عملکرد آن ها به تفصیل بیان می گردد.

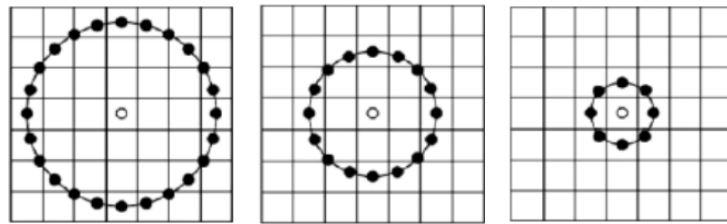
۲-۳ الگوی دودویی محلی

عملگر LBP اصلی به عنوان توصیفگری قدرتمند برای بافت تصویر معرفی شده است [۲۹]. این عملگر برای هر پیکسل با توجه به برچسب پیکسل های همسایگی 3×3 یک عدد دودویی تولید می کند. برچسب ها با آستانه سازی مقدار پیکسل های همسایه با مقدار پیکسل مرکزی به دست می آیند. به این صورت که برای پیکسل های با مقدار بزرگتر یا مساوی مقدار پیکسل مرکزی برچسب ۱ و برای پیکسل های با مقادیر کوچکتر از مقدار پیکسل مرکزی برچسب ۰ قرار می گیرد. سپس این برچسب ها به صورت چرخشی در کنار هم قرار گرفته و یک عدد ۸ بیتی تشکیل می دهند. نحوه ی کار این عملگر در شکل ۱-۳ آمده است. محدودیت عملگر LBP



شکل ۱-۳: عملگر الگوهای دودویی محلی [۲۹].

پایه همسایگی کوچک (3×3) آن است که باعث می شود بر تصاویر با مقیاس بزرگ تسلط نداشته باشد. بدین منظور بعداً این عملگر با توسعه ای بر اندازه همسایگی مطرح شد که به صورت دایره ای با شعاع R پیکسل بر روی P پیکسل نشان داده می شود. این عملگر به صورت $LBP_{P,R}$ نشان داده می شود و حداکثر می تواند 2^P مقدار مختلف، مطابق با 2^P الگوی دودویی تولید شده توسط P پیکسل موجود بر روی شعاع همسایگی R تولید کند. برای مثال می توانید به شکل ۳-۲ مراجعه کنید که نحوه ی انتخاب پیکسل های همسایه در این نوع الگوی دودویی محلی را به ازای سه شعاع مختلف را نشان می دهد. بعد از برچسب گذاری یک تصویر



(آ) ($P = 8, R = 1$) (ب) ($P = 16, R = 2$) (ج) ($P = 24, R = 3$)

شکل ۳-۲: عملگر LBP با شعاع ها و تعداد همسایگی های مختلف [۱۴]

توسط عملگر LBP، هیستوگرامی از برچسب ها به صورت زیر تعریف می شود:

$$H_i = \sum_{xy} I(f_i(x, y) = i), i = 0, \dots, n-1 \quad (1-3)$$

که n تعداد برچسب های تولید شده توسط عملگر LBP و تابع I به صورت رابطه زیر تعریف شده است:

$$I(A) = \begin{cases} 1 & A \text{ is True} \\ 0 & A \text{ is False} \end{cases} \quad (2-3)$$

اپراتور الگوی دودویی محلی پایه برای تغییر شکل های مقیاس خاکستری یکنواخت، تغییر نمی کند که می تواند رده تمرکز پیکسل را در مجاورت محلی حفظ نماید و تصویری از برچسب های الگوی دودویی محلی محاسبه شده در یک منطقه به عنوان توصیفگر بافت نشان دهد.

این اپراتور می تواند مقادیر خروجی متفاوتی مربوط به الگوهای محلی متفاوت را توسط P پیکسل مجاور نشان دهد. در صورتی که تصویر بچرخد این پیکسل های اطراف، در هر کدام از قسمت های مجاور به همین ترتیب در مجاورت دایره حرکت می کنند و به این ترتیب نتیجه آن یک مقدار الگوی دودویی متفاوت می باشد

و تنها استثنای الگوها یک و صفر خواهد بود. به منظور حذف چرخش تصویر، یک الگوی دودویی محلی ثابت چرخشی در [۳۰] پیشنهاد شده است. که آن را در فرمول زیر مشاهده می کنید:

$$LBP_{P,R}^{ri} \min\{ROR(LBP_{P,R}, i) | i = 0, 1, \dots, p-1\} \quad (3-3)$$

که در آن $ROR(x, i)$ به عنوان یک متغیر سمت راست دایره ای کوچک در تعداد P بیت به صورت (x, i) اعلان می گردد. اپراتور $LBP_{P,R}^{ri}$ مقادیر آماری الگوهای انتشاری متغیر غیر چرخشی گوناگون را نشان می دهد که به ریزویژگی های مشخص شده درون تصویر بستگی دارد و در حقیقت می تواند الگوهایی را نشان دهد که به عنوان آشکار ساز ویژگی^۱ لحاظ شده است. با این وجود نشان داده شده است که چنین اپراتور الگوی محلی چرخشی ثابت ضرورتاً نمی تواند اطلاعات متمایزی را ارائه دهد زیرا تناوبات الگوهای انتشاری، جداگانه تلفیق شده و همان طور که در $LBP_{P,R}^{ri}$ نشان داده شده تفاوت های عمده ای با یکدیگر داشته و فقط قابلیت تعیین مقادیر خام فضای زاویه دار در فواصل ۴۵ درجه امکان پذیر می گردد.

در [۳۰] تعدادی از الگوهای بنیادی به نام الگوهای یکنواخت^۲ نام گذاری شده اند که عبارتند از الگوهایی که در آن ها با چرخش بیت ها حداکثر ۲ تغییر ۰ به ۱ یا برعکس داشته باشیم. به عنوان مثال ارقام ۰۰۰۰۰۰۰۰ و ۰۰۱۱۱۰۰۰۱ یکنواخت هستند. ذخیره ی الگوهای یکنواخت عملگر LBP جدیدی به نام $LBP_{P,R}^{u2}$ تولید می کند که برای همسایگی ۸ پیکسل ۵۹ الگو خواهد داشت، به طوری که برای ۵۸ بین^۳ اول هیستوگرام تعریف شده مربوط به الگوهای یکنواخت و بین ۵۹ ام برای ذخیره ی مجموع دیگر الگوهای غیر یکنواخت است. این عدد برای LBP استاندارد ۲۵۶ متناظر با تعداد حالت های مختلف عدد دودویی ۸ بیتی است. ۵۸ الگوی یکنواخت تولید شده توسط عملگر LBP پایه در شکل ۳-۳ نشان داده شده اند.

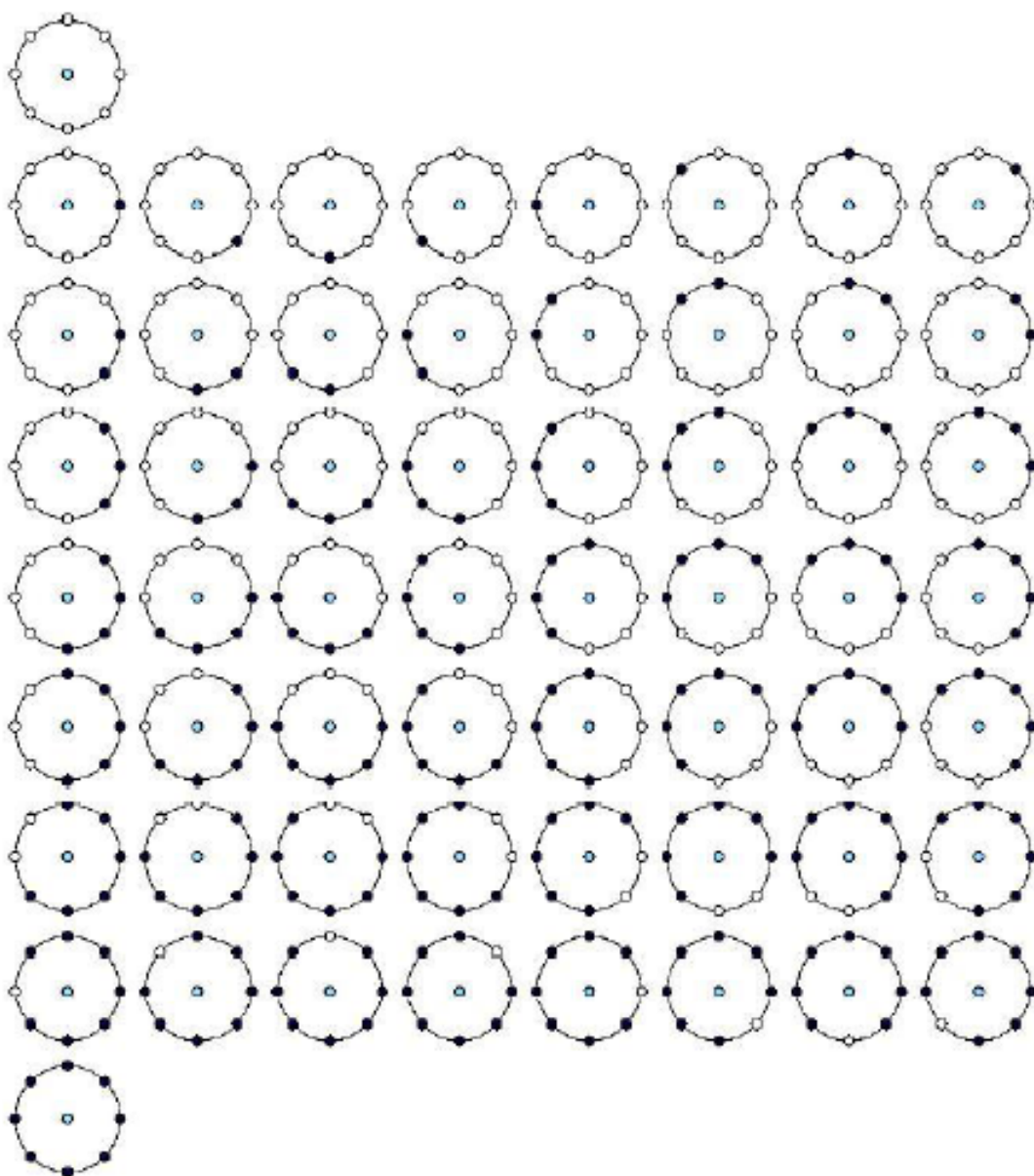
آزمایشات بر روی تصاویر نشان می دهد که ۹۰٪ الگوهای (۸ و ۱) و ۷۰٪ الگوهای (۱۶ و ۲) یکنواخت هستند [۳۰]. آزمایشات مشابه بر روی تصاویر چهره نشان می دهد که ۹۰,۶٪ الگوهای (۸ و ۱) و ۸۵,۲٪ الگوهای (۸ و ۲) یکنواخت هستند.

هیستوگرام LBP^{u2} به دست آمده از تصویر شامل اطلاعاتی درباره ی ریزالگوهای محلی از قبیل لبه، نقطه، نواحی صاف و ... روی تصویر است. بدین ترتیب بافت تصویر که همان چین و چروک های آن می باشد توصیف خواهد شد. شکل ۳-۴ نمونه هایی از ریزالگوهای توصیف شده در هیستوگرام LBP^{u2} را به خوبی نشان می دهد، در این شکل دایره های مشکی نشان دهنده عدد ۰ و دایره های سفید نشان دهنده ی عدد ۱ هستند.

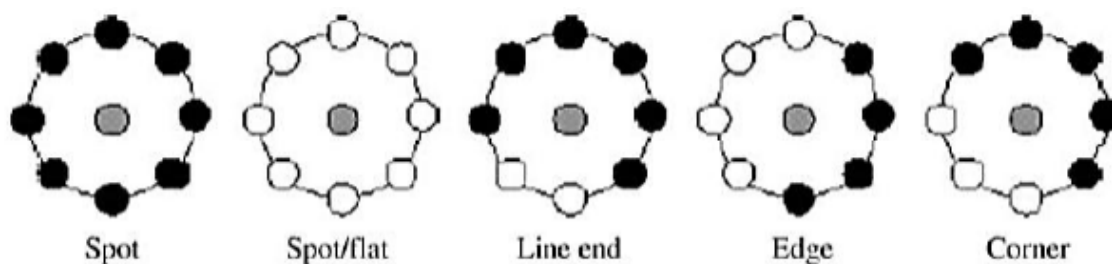
^۱ Feature Detector

^۲ uniform

^۳ bin



شکل ۳-۳: الگوی یکنواخت در عملگر LBP با شعاع ۱ و تعداد همسایگی ۸ [۳۰].



شکل ۳-۴: مثال هایی از مفهوم ریزالگوها در LBP

جدول ۳-۱: لیستی از آخرین تغییرات الگوی دودویی محلی

منبع	ویژگی ها	تغییرات	گزینش
[۵]، [۳۵]، [۲۲]	در نظر گرفتن تأثیرات پیکسل های مرکزی، ارائه الگوهای ساختاری کامل	LBP Improved	بهبود قدرت جداسازی
[۴۷]	ترکیب الگوهای غیریکنواخت در الگوهای یکنواخت	LBP Hamming	
[۱۷]، [۱۵]	جداسازی الگوهای محلی یکسان، افزایش ابعاد	LBP Extended	
[۱۱]	در نظر گرفتن علامت و مقدار نواحی محلی	LBP Completed	افزایش تحمل در برابر تغییرات
[۳۹]	مقاومت نسبت به تغییرات سطح خاکستری	Patterns Ternary Local	
[۳]	مقاومت نسبت به تغییرات یکنواخت سطح خاکستری، افزایش پیچیدگی محاسباتی	LBP Soft	
[۲۶]	استخراج اطلاعات ناهمگن، مقاومت نسبت به چرخش	LBP Elongated	انتخاب همسایه
[۲۵]، [۴۸]	گرفتن اطلاعات ساختار کوچک و کلان	LBP Multi-Block	
[۴۵]	کدگذاری نوع پیچ اطلاعات متنی	LBP Patch Three/Four	
[۳۲]، [۶]	گسترش الگوی دودویی محلی به داده های سه بعدی	LBP ۳D	گسترش به سه بعدی
[۵۰]، [۵۱]	توصیف بافت پویا، افزایش ابعاد توصیفگر	(LBP-TOP) LBP Volume	
[۲۴]، [۳۷]، [۴۰]، [۴۶]، [۳۶]، [۱۲]، [۴۹]	ترکیب فواید گابور و الگوی دودویی محلی، افزایش زمان محاسباتی و ابعاد توصیفگر	Wavelet Gabor and LBP	
[۱۶]، [۱۳]	به کارگیری فواید توصیفگر SIFT، کاهش طول ابعاد توصیفگر	SIFT and LBP	ترکیب با دیگر روش ها
[۲]	مقاومت نسبت به چرخش	Fourier Histogram LBP	

۳-۲-۱ تغییرات اخیر الگوی دودویی محلی

متد LBP در سال های اخیر به منظور بهبود عملکرد در کاربردهای مختلف دستخوش تغییراتی گردیده است. این تغییرات بر جنبه های مختلف الگوی دودویی محلی اصلی تمرکز کرده است.

- بهبود قدرت جداسازی آن
- افزایش تحمل در برابر تغییرات^۴
- نحوه ی انتخاب همسایه
- گسترش به داده های سه بعدی
- ترکیب آن با دیگر روش ها

در جدول ۳-۱ لیستی از آخرین تغییرات الگوی دودویی محلی آورده شده است.

^۴robust

۳-۳ ال اس اچ

قبل از این که مفهوم ال اس اچ (هش حساس به همسایگی) را توضیح بدهیم ابتدا به مفهوم هش می پردازیم:

۳-۳-۱ هش

با ساختن یک جدول هش، (یعنی ساختمان داده ای که به ما اجازه می دهد به سرعت از یک نماد به یک مقدار برسیم)، زمانی که یک ورودی می گیریم، می توانیم یک تابع شبه تصادفی از آن نماد را استفاده کنیم که به ما یک عدد صحیح بدهد که به عنوان اندیس جدول هش استفاده کنیم. بدین صورت نمادی با تعداد زیادی حروف و بیت به یک اندیس نسبتاً کوچک از جدول هش نگاشت می شود. برخورد زمانی رخ می دهد که دو نقطه به یک مقدار هش شوند. یک جدول هش ای که خوب طراحی شده باشد دارای زمان جستجوی $O(1)$ و حافظه $O(N)$ خواهد بود که N تعداد اقلام جدول است. علاوه بر این، یک تابع هش مناسب دو نماد که به هم نزدیک هستند را به دو جعبه ی مختلف هش می کند. این کار باعث می شود که جدول هش برای یافتن «تطبیقهای دقیق» مناسب باشد. برای یافتن تطبیق های نزدیک (یا تقریبی) به صورت کارا، می توان از ال اس اچ استفاده کرد که در ادامه توضیح می دهیم.

۳-۳-۲ هش حساس به همسایگی

^۵خانواده ای از تابع های هش $h : \{0, 1\}^d \rightarrow U$ را (r_1, r_2, P_1, P_2) حساس گوئیم

که برای $(r_1 < r_2, P_1 < P_2)$ اگر برای هر $q, p \in \{0, 1\}^d$

$$\bullet \text{ اگر } D(p, q) \leq r_1 \text{ آنگاه } pr[h(q) = h(p)] \geq P_1$$

$$\bullet \text{ اگر } D(p, q) > r_2 \text{ آنگاه } pr[h(q) = h(p)] \leq P_2$$

مفهوم ال اس اچ می تواند برای هر فضای متری D تعریف شود. ولی، برای فضای همینگ، خانواده ی ال اس اچ بسیار ساده است، کافی است که همه ی تابع های $h_i, i = 1, \dots, d$ را که $h_i(p) = p_i, p \in \{0, 1\}^d$ باشد برداریم. از آن جا که $Pr[h(q) = h(p)] = 1 - D(p, q)/d$ می توان نتیجه گرفت که این خانواده از توابع هش حساس است.

یک هشینگ حساس به همسایگی استاندارد، نگاشت ها را به صورت تصادفی انجام می دهد. هر یک از توابع

^۵Hash functions

هش به صورت مستقل با انتخاب یک بردار تصادفی w_k از یک توزیع گاوسی با میانگین صفر و کوواریانس مشخص تولید می شود. سپس تابع هش به صورت زیر در نظر گرفته می شود:

$$h_k(x) = \text{sign}(xw_k) \quad (۴-۳)$$

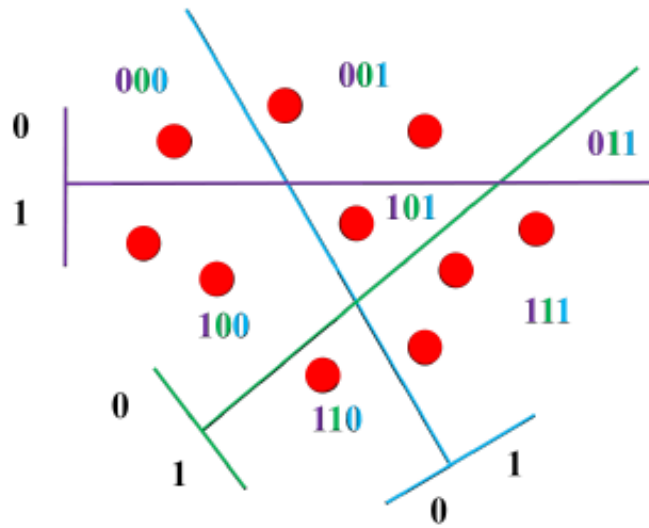
$$Y = \text{sign}(XW) \quad (۵-۳)$$

اگر هر بیت در کد با استفاده از نگاشت خطی تصادفی و آستانه تصادفی محاسبه گردد، آن گاه به دنبال افزایش تعداد بیت ها، فاصله ی همینگ میان کدها به سمت کسینوس زاویه ی میان آنها همگرا می شود. این بدان معناست که فاصله ی همینگ میان کدها به فاصله ی اقلیدسی میان آیت ها نزدیک می شود. بنابراین این متد توپولوژی داده های ورودی را به دنبال افزایش تعداد بیت ها حفظ می کند. نگاشت تصادفی در LSH و دیگر متدها استفاده می شود. این متدها مستقل از دیتاست می باشند و هیچ فرض پیشین درباره ی توزیع داده ها در نظر نمی گیرند ولی برای دستیابی به دقت بازایی خوب نیاز به تعداد کدهای بیشتری دارند به طور کلی یادگیری توابع دودویی هش مزایای زیر را دارد:

- کدها بیشتر فشرده می شوند.
- کلاس های عمومی تر از معیارهای شباهت می تواند حفظ شود.

در شکل ۳-۵ نمونه ای از نگاشت دودویی دو بعدی برای LSH نمایش داده شده است: اگر خانواده ی ال اس اچ با فاصله ی بزرگی بین P_1 و P_2 داشته باشیم، مسئله ی $NN - (r_1, r_2/r_1)$ را می توان به صورت زیر حل کرد. هنگام پیش پردازش، تمام نقاط داده مانند P به جعبه ی $h(p)$ هش می شوند. برای پاسخ گویی به نقطه ی ورودی q ، الگوریتم تمام نقاطی که در جعبه ی $h(q)$ قرار دارند را بازایی می کند و چک می کند که کدام به q نزدیک تر است. اگر فاصله ی بین P_1 و P_2 به اندازه ی کافی بزرگ باشد، می توان نشان داد که این روش زمان زیر خطی برای پاسخ گویی به نقطه ی ورودی نیاز دارد. متأسفانه فاصله ای که توسط خانواده ی ال اس اچ به وجود می آید به اندازه ی کافی بزرگ نیست، ولی با به هم چسباندن چند تابع هش می توان این فاصله را بزرگ کرد.

برای k و L ای که بعداً مشخص می شود، L تابع را $g_j(q) = (h_{1,j}(q), \dots, h_{k,j}(q))$ ، که در آن $1 \leq j \leq L$ ، $h_{t,j}$ (برای $1 \leq j \leq L$ ، $t \leq j$ به صورت مستقل و یکنواخت از H انتخاب شده اند را در نظر می گیریم. در زمان پیش پردازش، هر $p \in P$ (نقاط مجموعه داده) را در جعبه ی $g_j(p)$ برای $j = 1, \dots, L$ قرار می دهیم.



شکل ۳-۵: نگاشت دودویی دوبعدی LSH.

برای پاسخ گویی به نقطه ی ورودی در q ، جعبه های $g_1(q), \dots, g_L(q)$ جستجو می کنیم. دوراه حل زیر وجود دارد:

- جست وجو را پس از یافتن اولین L' نقطه متوقف کنیم.
- جست وجو را تا زمان بازیابی همه ی نقاط از جعبه ها ادامه دهیم.

نشان داده شده است که در اولین راه حل، اگر $L' = 3L$ باشد می تواند مساله همسایه ی نزدیک c تقریبی تصادفی را با پارامترهای R و $1 < \delta$ حل کند. برای این کار باید مقدار L را $\Phi(n^\rho)$ قرار دهیم که در آن $\rho = \frac{\ln \lambda / p_1}{\ln \lambda / p_r}$ ، توجه کنید که این باعث می شود الگوریتم در زمان متناسب با n^ρ اجرا شود که زمانی که $P_1 > P_r$ باشد نسبت به n زیر خطی است.

از طرف دیگر، در عمل، راه حل دوم پرکاربردتر به نظر می رسد. مخصوصاً به خاطر این که در آن نیازی به مشخص کردن پارامترهای اضافی نیست. علاوه بر این، هیچ دلیلی برای مشخص کردن نیست، زیرا که الگوریتم قادر است مساله ی همسایه ی نزدیک دقیق را حل کند، که در ادامه توضیح می دهیم. می توان نشان داد که راه حل دوم می تواند مساله ی همسایه ی نزدیک دقیق تصادفی را برای پارامتر R حل کند. مقدار δ وابسته به انتخاب پارامترهای k و L می باشد، برای هر δ می توان پارامترهای k و L را طوری تعیین کرد که احتمال خطا کوچکتر از δ باشد. پیچیدگی زمانی نیز وابسته به k و L است: می توان در بدترین حالت $\phi(n)$ باشد، ولی برای بسیاری از مجموعه های داده و انتخاب های پارامترها میتواند در زمان زیرخطی اجرا شود. کیفیت تابع ال اس اچ توسط پارامتر ρ تعیین می شود. نشان داده شده است که اگر فاصله بر اساس L_1

اندازه گیری شود، آن گاه خانواده ی توابع (R, cR, P_1, P_2) - حساس وجود دارد به طوری که $\rho = 1/c$.

۳-۳-۳ حل مسأله ی نزدیک ترین همسایه ی تقریبی

الگوریتم بالا را می توان با تغییرات کمی برای حل مسئله ی نزدیک ترین همسایه ی تقریبی به کار برد، بدون این که پیچیدگی الگوریتم تغییر پیدا کند. مسأله ی نزدیک ترین همسایه ی تقریبی تحت نرم L_p ، برای $p \in [1, 2]$ ، می تواند به همین مسئله در فضای همینگ کاهش پیدا کند (تبدیل شود). این کار به راحتی برای L_d انجام می شود. اگر فرض کنیم که مختصات نقاط p دارای مقادیر $\{1, \dots, M\}$ باشند، اگر $U(p) = (U(p_1), \dots, U(p_d))$ را تعریف کنیم که در آن $U(x)$ یک رشته از $1 \times M$ و x دنبال می شود باشد، آن گاه خواهیم داشت $\|p - q\|_1 = D(U(p), U(q))$. در حالت کلی، M می تواند بسیار بزرگ باشد، ولی می تواند مسئله ی نزدیک ترین همسایه ی تقریبی به $d^{(1)}$ کاهش یابد. بنابراین می توانیم $NN(r, c)$ را تحت L_1 به $NN(r, c)$ در فضای همینگ کاهش دهیم. برای به دست آوردن الگوریتم برای نرم L_p زمانی که $p \in [1, 2]$ ، باشد از این موضوع استفاده می کنیم که $L_1^{(d)}$ با تغییرات محدودی^۶ جاسازی^۷ شود. از طرف دیگر برای $p=2$ ، می توان مسأله را مستقیم در فضای اقلیدسی حل کرد.

۴-۳-۳ نقش نگاشت های تصادفی در ال اس اچ

ال اس اچ مبتنی بر این ایده ی ساده است که اگر دو نقطه به هم نزدیک هستند، آن گاه پس از نگاشته شدن^۸، این نقاط هم چنان به هم نزدیک می مانند. برای شرح بیشتر مطلب از توضیح نگاشت تصادفی آغاز می کنیم که یک نقطه را از یک فضای با بعد بالا به یک زیرفضای با بعد پایین نگاشت می کند.

نگاشت های تصادفی: ضرب داخلی

در قلب ال اس اچ، نگاشت اسکالر (یا همان ضرب داخلی) $h(V) = V^T \cdot X$ می باشد که V نقطه ای ورودی در فضای با بعد بالا بوده و X برداری است که مولفه های آن به صورت تصادفی از یک توزیع گاوسی مثلاً $N(0, 1)$ انتخاب شده است. این نگاشت اسکالر سپس به چند جعبه ی هش گسسته سازی می شود، با این هدف که اقلام نزدیک در فضای اصلی به یک جعبه هش شوند. تابع هش ای که نتیجه می شود عبارت است

^۶bounded distortion

^۷embedding

^۸projection

از:

$$h^{X+b}(V) = \lceil \frac{X^T \cdot V + b}{w} \rceil \quad (۳-۶)$$

که در آن w عرض هر جعبه ی نتیجه ی گسسته سازی است و یک متغیر تصادفی است که بین 0 و w توزیع یکنواخت دارد و کمک می کند تا خطای گسسته سازی بدون اضافه کردن به هزینه ی عملکرد، راحت تر تحلیل شود. برای آن که عملکرد نگاشت برای کار ما مناسب باشد، می بایست نقاط نزدیک را به مکان های نزدیک به هم تصویر کند. به عبارت دیگر:

- برای هر دو نقطه p و q در R^d که نزدیک به هم هستند، احتمال بالای P_1 ای وجود داشته باشد که در جعبه ی یکسان هس شوند:

$$P_H[h(p) = h(q)] \geq P_1 \text{ for } \|p - q\| \leq R_1 \quad (۳-۷)$$

- برای هر دو نقطه p و q در R^d که از هم دور هستند، احتمال کم $P_2 < P_1$ ای وجود داشته باشد که در جعبه ی یکسان هس شوند:

$$P_H[h(p) = h(q)] \leq P_2 \text{ for } \|p - q\| \geq cR_1 = R_2 \quad (۳-۸)$$

در این تعاریف $\|\cdot\|$ نرم L_2 است و $R_2 > R_1$ است. توجه شود که به واسطه ی خطی بودن ضرب داخلی، اختلاف بین دو نقطه ی هس شده $\|h(p) - h(q)\|$ اندازه ای دارد که توزیع آن با $\|p - q\|$ متناسب است و بنابراین $P_1 > P_2$.

های تصادفی: k ضرب داخلی

برای بزرگ کردن اختلاف بین P_1 و P_2 ، ضرب داخلی را به صورت موازی انجام می دهیم. این کار نسبت احتمالی را که نقاط با فاصله های متفاوت به یک جعبه هس شوند را افزایش می دهد زیرا $(P_1/P_2)^k > 1$. نگاشت نهایی با انجام k ضرب داخلی مستقل برای تبدیل نقطه ی ورودی v به k عدد حقیقی به دست می آید. همان طور که برای نگاشت اسکالر انجام شده، k ضرب داخلی را با هدف این که نقاط مشابه به جعبه ی یکسان هس شوند، گسسته سازی می کنیم.

افزایش عرض گسسته سازي w باعث افزایش تعداد نقاطی می شود که به هر جعبه هش می شوند. برای رسیدن به نتیجه ی نهایی همسایه ی نزدیک می بایست یک جستجوی خطی روی همه ی نقاطی که در جعبه ی یکسان با نقطه ی ورودی قرار می گیرند انجام دهیم. بنابراین تغییر w یک موازنه بین جدول بزرگ تر با جستجوی خطی کوچک تر، یا یک جدول کوچک تر با تعداد بیشتری نقاط در جستجوی نهایی.

برای هر مجموعه از k ضرب داخلی، اگر نقطه ورودی و نزدیک ترین همسایه در یک جعبه قرار بگیرد، ما موفق شده ایم. این با احتمال P_1^k اتفاق می افتد که با استفاده از ضرب داخلی های بیشتر کاهش پیدا می کند. برای کاهش تأثیر گسسته سازي نادرست در هر نگاشت، L نگاشت مستقل تشکیل می دهیم و همسایه ها را با کمک همه ی آنها به دست می آوریم. این بدین دلیل است که احتمال اینکه نزدیک ترین همسایه واقعی برای همه ی نگاشت ها در جعبه ی نادرست قرار گیرد بسیار کم است. با افزایش L می توانیم نزدیک ترین همسایه ی واقعی را با هر احتمالی به اندازه ی دلخواه بالا، به دست آوریم.

۳-۳-۵ پیاده سازی هش

عمل نگاشت و گسسته سازي هر نقطه را در جعبه ای قرار می دهد که با k اندیس صحیح مشخص می شود. با توجه به این که این فضای k بعدی تنگ^۹ است، می توانیم از روش های معمول (و دقیق) هشینگ استفاده کنیم تا به صورت کارا بفهمیم که چه زمانی نقاط در یک جعبه افتاده اند. ما در اینجا روش $E^2 LSH$ را شرح می دهیم ولی روش های پیچیده تر نیز وجود دارد.

جستجوی ساده برای یافتن نقاطی از مجموعه داده که در جعبه ی یکسان با نقطه ورودی قرار دارند میتواند $O(\log N)$ زمان صرف کند ولی ما این را به $O(1)$ کاهش می دهیم چون از توابع هش معمولی استفاده می کنیم. ابتدا از یک تابع هش معمولی استفاده می کنیم تا نگاشت k بعدی گسسته سازي شده را به یک اندیس خطی تبدیل کنیم

$$T_1 = (\sum_i H_i k_i) \bmod P_1 \quad (9-3)$$

که در آن H_i وزن های صحیح هستند و P_1 سایز جدول هش است (یک عدد اول بزرگ). هدف این جدول هش این است که هر نقطه ی جداگانه را در فضای k بعدی در یک خانه ی جداگانه از P_1 خانه جدول قرار دهد، تا این که برخوردی بین نقاط متفاوت رخ ندهد (مقدار یکسان هش به آنها نسبت داده نشود)، با آن که جدول هش خوب اقلام را به خوبی به طور یکنواخت توزیع می کند، ولی با کوچک شدن جدول (برای آن که

^۹sparse

بتوان آن را در حافظه ی اصلی جاي داد) ريسک برخورد نقطه هاي نامرتبط افزايش می يابد. برای حل اين مشکل هش دومی به نام T_2 نقاط k -بعدی را به کار می گیريم تا بررسی کنیم که نقاطی که از جدول هش بازيايی شده اند واقعا با نقاط ورودی تطبيق دارند يا خير. T_2 همان شکل T_1 را دارد ولی وزن ها وسایزهای مختلفی استفاده می کند. مقادير دومين هش را در جعبه هایی که توسط اولین هش انتخاب شده اند ذخيره می کنیم. در زمان بازيايی، مقدار ذخيره شده را با اندیس جعبه مقایسه می کنیم تا تطبيق هاي واقعی را در صورت وجود بیاییم. با توجه به اینکه مقادير هش دوم کوتاه هستند (مثلاً مقادير ۱۶ - بيتی) مقایسه شان خیلی سریع تر از مقایسه ی k -بعدی نقاط اولیه می باشد. علاوه بر این، احتمال برخورد پی در پی تحت هر دوی T_1 و T_2 به شدت کاهش می يابد، حتی زمانی که جدول هش کوچک است.

۳-۳-۶ بررسی عملکرد

دقت

دقت ال اس اچ توسط احتمال این که نزدیک ترین همسایه ی واقعی را پیدا می کند تعیین می شود. برای تحلیل این موضوع، مفهوم توزیع s - مقاوم را معرفی می کنیم. توزیع s - مقاوم است اگر برای هر مجموعه از متغیرهای تصادفی، $X_1, \dots, X_n$ مستقل با توزیع یکسان D و هر مجموعه از اعداد حقیقی $v_1, \dots, v_n$ متغیر تصادفی $\sum_i V_i X_i$ توزیع یکسان با توزیع متغیر تصادفی $(\sum_i |V_i|^s)^{1/s}$ داشته باشد که در آن X از D انتخاب شده است. برای $s=2$ (یعنی همان نرم L_2)، توزیع گاوسی s - مقاوم است. برای تحلیل دقت نگاشت هایمان، توجه کنید که نگاشت های دو نقطه ی نزدیک p و q که با فاصله ی $u = \|p - q\|$ از هم جدا شده اند، همیشه به هم نزدیک خواهد بود. ولی به دلیل گسسته سازی ممکن است در دو جعبه ی متفاوت قرار گیرند. احتمال این که این دو نقطه با گسسته سازی به یک جعبه نسبت داده شوند عبارت است از:

$$p(u) = Pr_{a,b}[h_{a,b}(p) = h_{a,b}(q)] = \int_0^w \frac{1}{u} f_s\left(\frac{t}{u}\right) \left(1 - \frac{t}{w}\right) dt \quad (10-3)$$

که در آن f_s تابع چگالی احتمال H است. برای هر جعبه با عرض w ، این احتمال با افزایش فاصله ی u کاهش می يابد. احتمال P_H را در فضای L_2 با $R_1 = w$ به صورت زیر حساب می کنیم:

$$P = 1 - 2F(-w/c) - \frac{2}{\sqrt{2\pi}w/c} \times (1 - e^{-(w^2/2c^2)}) \quad (11-3)$$

در این جا $F(\cdot)$ تابع چگالی احتمال تجمعی یک متغیر تصادفی گاسی است و c نسبت فاصله هاست اگر $c=1$ قرار دهیم، P_1 را به دست می آوریم.

احتمال این که نقطه ای خاص با نقطه ی ورودی در یک جعبه قرار گیرد عبارت است از $(P_1)^k$ ، از طرف دیگر احتمال این که همه ی نگاشت ها نتوانند برخوردی بین نقطه ی ورودی و نزدیک ترین همسایه ی واقعی تولید کنند عبارت است از $(1 - P_1^k)^L$. اگر بخواهیم که احتمال این که ال اس اچ نتواند نزدیک ترین همسایه ی واقعی را پیدا کند بیش تر از δ نباشد، با داشتن یک مقدار برای k ، بایستی L حداقل مقدار زیر باشد:

$$L = \frac{[\log \delta]}{\log(1 - P_1^k)} \quad (12-3)$$

الگوریتم E^2LSH بهترین مقدار را برای k با آزمودن هزینه محاسبات برای نمونه های در مجموعه داده به دست می آورد.

سرعت

زمان لازم برای یافتن نزدیک ترین همسایه، زمان T_g برای محاسبه ی نگاشت ها و هش کردن، به اضافه ی زمان T_c مورد نیاز برای جست و جو در جعبه ها برای برخوردها می باشد. با توجه به این که KL نگاشت وجود دارد، $O(nkL)T_g$ است، که در آن n بعد فضای داده ی اصلی است. از طرف دیگر T_c به صورت خطی با امید ریاضی تعداد برخوردها (یعنی $T_c = O(dLN_c)$ که در آن d میانگین تعداد نقاط در هر جعبه می باشد) افزایش می یابد. N_c ، که امید تعداد برخوردهای یک نگاشت می باشد، توسط عبارت زیر داده شده است:

$$N_c = \sum_{q' \in D} p^k(\|q - q'\|) \quad (13-3)$$

که در آن $p(\cdot)$ احتمال کمک هر نقطه به ایجاد برخورد می باشد و D کل نقاط در پایگاه را نشان می دهد. به راحتی می توان دید که T_g متناسب با k افزایش می یابد، در حالی که T_c کاهش می یابد چون $P_k < P$ برای $P < 1$ و $k > 1$.

۳-۴ کرنل های ثابت به جابه جایی

کدگذاری دودویی حساس به همسایگی با استفاده از کرنل های ثابت به جابه جایی ویژگی هایی را استخراج می کند که تضمین شده است نگاشت دودویی داده اصلی نسبت به جابه جایی مقاوم هستند. ویژگی های تصادفی فوریه^{۱۰} به صورت زیر تعریف می شوند:

$$\phi_{w,b}(x) = \sqrt{2} \cos(w^T x + b) \quad (14-3)$$

که در آن $w \sim P_k$ و $b \sim Unif[0, 2\pi]$ ، به عنوان نمونه برای کرنل $K(S) = e^{-\gamma \|s\|^2/2}$ و $w \sim$ $Normal(0, \gamma I)$ می توان نشان داد $E_{w,b}[\phi_{w,b}(x)\phi_{w,b}(y)] = K(x, y)$. کد باینری به صورت زیر محاسبه می شود:

$$sign(\phi_{w,b}(X) + t) \quad (15-3)$$

که در آن t آستانه تصادفی و $t \sim Unif[-1, 1]$ می باشد.

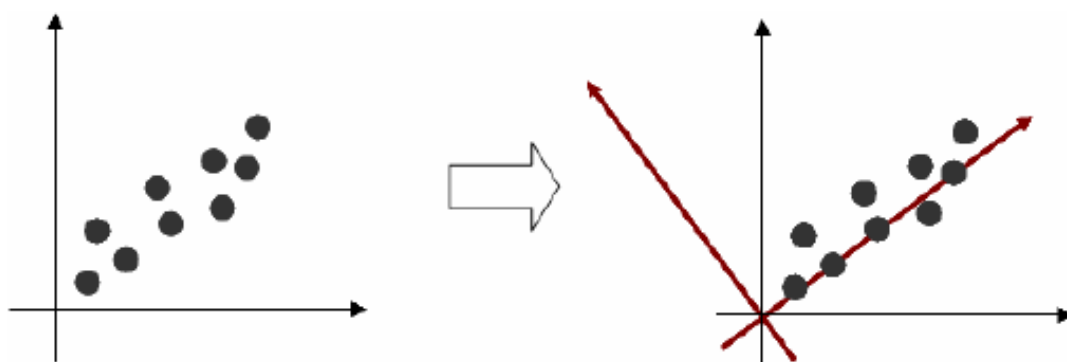
۳-۵ چندی سازی تکرار شونده

برای تشریح الگوریتم هشینگ چندی سازی تکرار شونده [۹]، [۱۰] نیاز به درک مفاهیم تکنیک آنالیز اجزای اصلی، تحلیل همبستگی کانونیک و تجزیه مقادیر تکین می باشد بدین منظور در ابتدا به بیان این مفاهیم می پردازیم.

۳-۵-۱ تکنیک آنالیز اجزای اصلی

در این روش محورهای مختصات جدیدی برای داده ها تعریف شده و داده ها بر اساس این محورهای مختصات جدید بیان می شوند. اولین محور باید در جهتی قرار گیرد که واریانس داده ها ماکسیمم شود (یعنی در جهتی که پراکندگی داده ها بیشتر است). دومین محور باید عمود بر محور اول به گونه ای قرار گیرد که واریانس داده ها ماکسیمم شود. به همین ترتیب محورهای بعدی عمود بر تمامی محورهای قبلی به گونه ای قرار می گیرند که داده ها در آن جهت دارای بیشترین پراکندگی باشند در شکل ۳-۶ این مطلب برای داده های دو بعدی

^{۱۰} Random Fourier features



شکل ۳-۶: انتخاب محوره‌های جدید برای داده‌های دو بعدی.

نشان داده شده است. روش PCA به نام‌های دیگری نیز معروف است مانند:

• Transform Loeve Karhunen (KLT)

• Transform Hotelling

• (EOF) Function Orthogonal Empirical

الگوریتم PCA

مرحله ۱ - انتخاب داده

مرحله ۲ - کم کردن میانگین از داده‌ها در این مرحله میانگین هر بعد از مقادیر آن بعد کم می‌شود تا میانگین داده‌ها در هر بعد صفر گردد.

مرحله ۳ - محاسبه ماتریس کواریانس

مرحله ۴ - محاسبه بردارهای ویژه و مقادیر ویژه در این مرحله بردارهای ویژه و مقادیر ویژه ماتریس کواریانس محاسبه می‌شود. طبق قضایای جبر خطی، یک ماتریس متقارن $n \times n$ دارای n بردار ویژه مستقل و n مقدار ویژه می‌باشد. مقادیر ویژه میزان پراکندگی داده‌ها در راستای بردار ویژه مربوطه را نشان می‌دهد. می‌توان گفت بردار ویژه‌ای که دارای بزرگترین مقدار ویژه است مولفه اصلی داده‌های موجود می‌باشد.

مرحله ۵ - انتخاب مولفه‌ها و ساختن Vector Feature در این مرحله مفهوم کاهش ابعاد داده وارد می‌شود. بردار ویژه‌ای که در مرحله‌ی قبل به دست آمده است بر اساس مقادیر ویژه‌ی آنها از بزرگ به کوچک مرتب می‌شود (مقادیر ویژه‌ی ماتریس کواریانس همگی بزرگتر یا مساوی صفر هستند). بدین ترتیب مولفه‌های داده‌ها از پراهمیت به کم اهمیت مرتب می‌شوند. در این مرحله به منظور کاهش داده‌ها، مولفه‌های کم

اهمیت حذف می شود. البته این کار با از دست دادن مقدار کمی اطلاعات همراه است. کاری که در این مرحله باید انجام شود ایجاد یک Vector Feature است که در واقع ماتریسی از بردارها می باشد. این ماتریس شامل آن بردارهای ویژگی ای است که ما می خواهیم آنها را نگه داریم.

$$FeatureVector = (eig_1 eig_2 eig_3 \dots eig_n) \quad (3-16)$$

مرحله ۶- به دست آوردن داده های جدید در آخرین مرحله از PCA فقط باید ترانهاده ی ماتریس Feature Vector که در مرحله ی قبل به دست آمده را در ترانهاده ی داده های نرمال سازی شده ضرب شود.

$$FinalData = RowFeatureVector \times RowDataAdjust \quad (3-17)$$

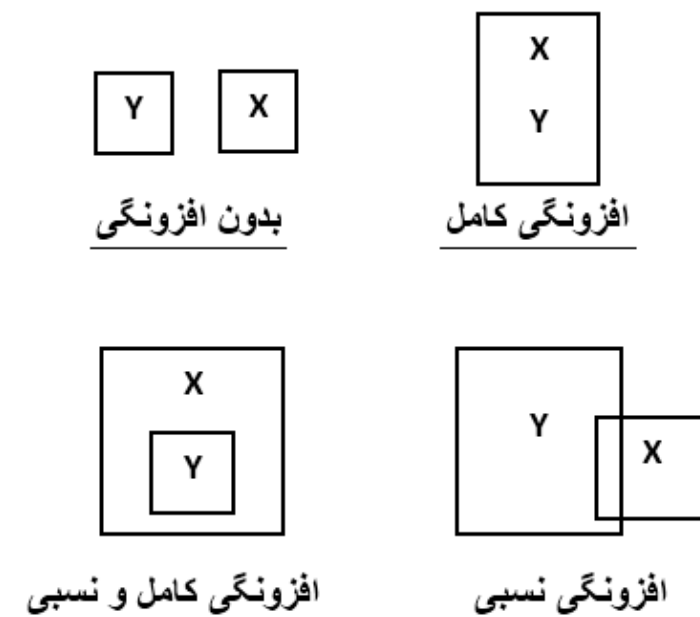
که RowFeatureVector ماتریسی است که بردارهای ویژه در سطرها ی آن به ترتیب مقادیر ویژه از بالا به پایین قرار گرفته اند و RowDataAdjust ماتریسی است که شامل داده هایی است که میانگین هر بعد از آن بعد کم شده است. در این ماتریس، داده ها در ستون های آن ذخیره شده و هر سطر آن مربوط به یک بعد است.

۳-۵-۲ تحلیل همبستگی کانونیک

تحلیل همبستگی کانونیک یکی از متدوال ترین روشهای تحلیل چند متغیره بوده و هدف آن تعیین ارتباط خطی بین متغیرهای چند بعدی است. در این روش در نظر گرفتن یک مجموعه از متغیرها به عنوان متغیرهای مستقل و مجموعه ی دیگر به عنوان متغیرهای وابسته می تواند بسیار مفید باشد. از جمله امتیازات روش تحلیل همبستگی کانونیک در مقایسه با تحلیل همبستگی معمولی^{۱۱} این است که روش OCA به سیستم مختصاتی که در آن تعریف شده است، وابسته می باشد. این بدان معنی است که حتی اگر رابطه ای قوی بین دو متغیر چند بعدی وجود داشته باشد، این ارتباط ممکن است با انتخاب سیستم مختصات استفاده شده در روش OCA دیده نشود. در حالی که در یک سیستم مختصات^{۱۲} دیگر، این رابطه خطی میزان همبستگی بالایی را ارائه دهد. روش CCA سیستم مختصاتی را می یابد که در آن میزان همبستگی دارای مقدار بهینه است. در بعضی از دسته داده های چند متغیره، متغیرها به طور طبیعی به دو گروه مجزا مثل $X_1, X_2, X_3, \dots, X_p$ و $Y_1, Y_2, Y_3, \dots, Y_p$ تقسیم می شوند. هدف اصلی CCA، ساختن دو گروه جدید از مولفه ها به صورت $U=aX$

^{۱۱} Ordinary Correlation Analysis (OCA)

^{۱۲} Coordinating System



شکل ۳-۷: تصویری از ضرایب افزونگی حاصل از تحلیل همبستگی کانونی .

و $V=bY$ است که ترکیب خطی از متغیرهای اولیه هستند. همبستگی کانونی با دو مجموعه از داده ها آغاز می شود که شامل بردارهایی از مشاهدات انجام شده بر کلیه متغیرها می باشد. هدف همبستگی کانونی، با ایجاد X بعنوان یک بردار m بعدی از متغیرهای پیش بین و Y بعنوان یک بردار P بعدی از متغیرهای ملاک، دستیابی به یک ترکیب خطی از متغیرهای پیش بین است که دارای حداکثر همبستگی با یک ترکیب خطی از متغیرهای ملاک می باشد. تحلیل همبستگی کانونی اندازه رابطه بین دو مجموعه از متغیرها را با ضرایب افزونگی^{۱۳} تعیین می کند. ضرایب افزونگی، درجه همپوشانی بین دو مجموعه متغیر را نشان می دهد؛ به بیان دقیق تر ضرایب افزونگی، نمایه میانگین درجه انحراف در یک مجموعه متغیر می باشد که با متغیرهای کانونی در سایر مجموعه ها سهمیم یا قابل پیش بینی از آنها می باشد.

استفاده از یک مجموعه متغیر برای پیش بینی مجموعه دوم از متغیرها بیانگر این است که مجموعه دوم با وجود مجموعه اول زائد^{۱۴} است. یک شاخص نامتقارن به این معنا است که ضریب افزونگی یک مجموعه از متغیرها عموماً با ضریب افزونگی مجموعه دیگری از متغیرها برابر نمی باشد. ضرایب افزونگی یا بطور مجزا یا بصورت ادغام شده در توابع کانونی آزمون می شوند. شکل ۳-۷ برای درک بهتر، تصویری از ضرایب افزونگی حاصل از تحلیل همبستگی کانونی را نشان می دهد. بنابراین همانطور که پیشتر ذکر شد تحلیل

^{۱۳}Redundancy Coefficients

^{۱۴}Redundant

همبستگی کانونی اندازه رابطه بین دو مجموعه از متغیرها را با ضرایب افزونگی (Rd) تعیین می کند. ضرایب افزونگی، درجه همپوشانی بین دو مجموعه متغیر را نشان می دهد؛ به بیان دقیق تر ضرایب افزونگی، نمایه میانگین درجه انحراف در یک مجموعه متغیر می باشد که با متغیرهای کانونی در سایر مجموعه ها سهمیم یا قابل پیش بینی بوسیله آنها می باشد.

۳-۵-۳ تجزیه مقادیر تکین

یکی از مهم ترین ابزارهای تفکیک و تجزیه در تسهیل حل دستگاه های بزرگ خطی، روش تجزیه مقادیر منفرد است. این روش، خود تعمیمی از روش تجزیه مقادیر ویژه ی اپراتورهای مربعی است. با این تفاوت که در این روش هر ماتریس با هر بعد را می توان به حاصل ضرب سه ماتریس که یکی از آن ها قطری و دو ماتریس دیگر نیز ماتریس هایی اورتوگونال و یا معکوس پذیر می باشد، تبدیل نمود. به عبارتی دیگر در این روش برای هر ماتریس با هر بعد دو ماتریس یافت می گردد، به طوری که با اعمال دو ماتریس یکی از چپ و دیگری از راست، یک ماتریس قطری حاصل می شود. از این روش در زمره ی روشهای قطری سازی^{۱۵} نیز می باشد. اگر معادله گسسته شده مساله به صورت زیر باشد:

$$Ax = b \quad A \in R^{m \times n}, x \in R^n, b \in R^m \quad (3-18)$$

آن گاه یکی از ابزارهای کار با این مساله روش تجزیه مقادیر تکین ماتریس A است. برای آسانی کار فرض می نمائیم که $m \geq n$ باشد. در این حالت تجزیه مقادیر منفرد ماتریس A به صورت ذیل خواهد بود:

$$A_{mn} = U_{mm} \Sigma_{mm} V_{nn}^T = \sum_{i=1}^n (u_i \delta_i v_i^T) \quad (3-19)$$

در معادله ۳-۱۹ $U_{mm} = u_1, u_2, \dots, u_m$ و $V_{mm} = v_1, v_2, \dots, v_m$ ماتریس هایی با ستون های اورتونمال است،

$$V^T V = I_{nn}, U^T U = I_{mm} \quad (3-20)$$

^{۱۵}Bi-diagonalization

ستون های ماتریس های U و V را به ترتیب بردارهای منفرد چپ و راست ماتریس A می نامند. عناصر قطر اصلی ماتریس $\Sigma_{mn} = [\sigma_i \delta_{ij}]_{1 \leq i \leq m, 1 \leq j \leq n}$ غیر منفی بوده و غالباً به ترتیب ذیل قرار می گیرند:

$$\sigma_1 \geq \sigma_2 \geq \dots \geq \sigma_n \geq 0 \quad (21-3)$$

مقادیر قطر اصلی Σ ، مقادیر منفرد ماتریس A و نسبت σ_1/σ_n عدد شرط ماتریس نامیده می شود، از روابط $AA^T = U\Sigma^T\Sigma U^T$ و $A^T A = V^T\Sigma^T\Sigma V$ دیده می شود که ماتریس A به شدت به تجزیه مقادیر ویژه ماتریس های AA^T و $A^T A$ وابسته است. این وابستگی یکتایی تجزیه مقادیر منفرد را نشان می دهد که از یکتایی تجزیه مقادیر ویژه ماتریس های متقارن نتیجه می گردد. در ارتباط با مسائل بدو وضع گسسته دومشخصه اصلی برای تجزیه مقادیر منفرد اغلب یافت می گردد: مقادیر منفرد به تدریج بدون گسستگی به سمت صفر میل میکند و افزایش ابعاد A ، تعداد مقادیر منفرد کوچک را افزایش خواهد داد. در حالی که i با افزایش اندیس کاهش می یابد بردارهای منفرد u_i, v_i نیز با افزایش i مرتباً تغییر علامت می دهند. هم چنین با توجه به رابطه ی اصلی تجزیه مقادیر منفرد در ابتدای بخش روابط زیر نیز حاصل می گردد:

$$\begin{cases} Av_i = \sigma_i u_i, & \|Av_i\|_2 = \sigma_i. \\ A^T u_i = \sigma_i v_i, & \|A u_i\|_2 = \sigma_i. \end{cases} \quad i = 1, 2, \dots, n \quad (22-3)$$

۳-۵-۴ الگوریتم

این روش یکی دیگر از طرح های ساده و کارآمد برای یافتن مناسبترین چرخش برای داده هایی به مرکز صفر به منظور به حداقل رساندن خطای کواتیزاسیون به هنگام نگاشت داده ها به رئوس ابرمکعب دودویی می باشد. به منظور یافتن کدهایی با بیشترین واریانس و ناهمبستگی میان آن ها، بردارهای ورودی با یک روش بدون ناظر مانند تکنیک آنالیز اجزای اصلی^{۱۶} نگاشت می شوند اگرچه می توان از تکنیک های با ناظر مانند CCA^{۱۷} نیز استفاده نمود.

اگر $W \in R^{d \times q}$ ماتریس ضرایب باشد که با روش PCA به دست آمده است آن گاه هر بیت $k = 1, \dots, q$ می تواند با رابطه ی زیر محاسبه نمود:

$$h_k(x) = \text{sgn}(xw_k) = \text{sgn}(v) \quad (23-3)$$

^{۱۶}PCA

^{۱۷}Canonical correlation analysis

کلیه ی فرآیند رمزنگاری^{۱۸} به صورت:

$$Y = \text{sgn}(XW) = \text{sgn}(V) \quad (24-3)$$

اگر W راه حل بهینه باشد آن گاه WR ، که در آن R یک ماتریس متعامد $q \times q$ است نیز بهینه می باشد. بنابراین داده های نگاشت شده $V=XR$ نیز به صورت متعامد تبدیل می شوند. ITQ به صورت متعامد داده ها را تبدیل می کند تا خطای چندی سازی را به حداقل برساند.

$$Y = \text{sgn}(XWR) = \text{sgn}(VR) \quad (25-3)$$

فرض کنید $v \in R^d$ یک بردار در فضای نگاشت باشد. بسیار ساده است که نشان داده شود $\text{sgn}(v)$ یکی از رئوس ابر مکعب $\{-1, 1\}^q$ می باشد که در فضای اقلیدسی به v نزدیک است. فقدان چندی سازی اختلاف میان بردار v و نگاشت آن به کد دودویی در ابرمکعب $\{-1, 1\}^q$ می باشد:

$$\|\text{sgn}(v) - v\|_2 \quad (26-3)$$

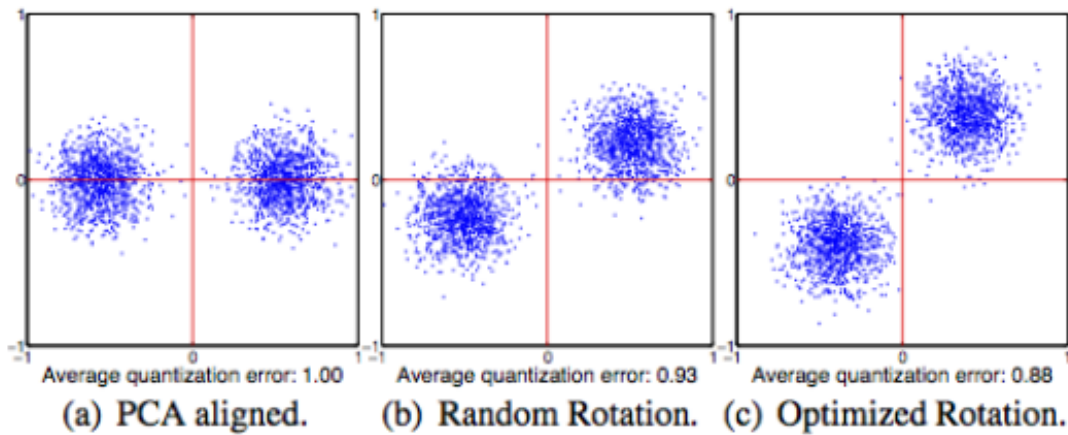
در صورتی که مقدار این فقدان کمتر باشد کدهای باینری بهتری تولید می گردد بنابراین در این روش به دنبال یافتن چرخش متعامدی می باشیم به گونه ای که نقاط نگاشت شده به کدهای باینری نظیرشان نزدیک باشند.

$$\min(Q(Y, R)) \quad (27-3)$$

$$Q(Y, R) = \|Y - VR\|_F \quad (28-3)$$

شکل ۸-۳ نمایش از عملکرد الگوریتم ITQ را نشان می دهد ایده ی اصلی این روش کوانتیزه کردن هر یک از نقاط داده به نزدیک ترین رئوس ابرمکعب $(\pm 1, \pm 1)$ می باشد. در شکل اول بردارهای x و y مطابق با جهت PCA داده ها می باشند. کوانتیزاسیون نقاط داده در یک خوشه را به رئوس متفاوت نسبت می دهد. شکل دوم چرخش تصادفی داده ها در این روش واریانس داده ها متعادل شده است و خطای چندی سازی کمتر شده

^{۱۸}encoding



شکل ۳-۸: تصویری از عملکرد الگوریتم کواتیزاسیون تکرار شونده [۱۰].

است در شکل سوم چرخش بهینه با استفاده از ITQ به دست آمده و خطای چندی سازی کمتر شده است. در این جا کار را با اعمال الگوریتم PCA روی نقاط داده آغاز می کنیم و مساله یادگیری یک کد باینری خوب و مناسب را به صورت کمینه کردن خطای چندی سازی داده های تصویر شده به هر یک از رؤوس ابرمکعب فرمول می کنیم. در این بخش به مساله یادگیری کدهای باینری بدون هیچ اطلاعات با نظارت به برجسب کلاس ها اشاره می کنیم در ابتدا یک کاهش ابعاد خطی به داده ها اعمال می گردد سپس چندی سازی دودویی را در فضای نتیجه اعمال می کنیم.

بیشینه کردن واریانس

ما به دنبال کدهای کارایی هستیم که در آن واریانس هر بیت بیشینه است و بیت ها دو به دو ناهمبسته هستند ما می توانیم به این هدف با بیشینه کردن تابع هدف زیر دست یابیم:

$$\tau(w) = \sum_k \text{var}(h_k(x)) = \sum_k \text{var}(\text{sgn}(xw_k)) \quad (29-3)$$

$$\frac{1}{n} B^T B \quad (30-3)$$

زمانی واریانس بیشینه می گردد که توابع کدگذاری دقیقاً بیت های متوازن تولید کنند یعنی زمانی که دقیقاً برای $h_k(x) = 1$ و برای نیمه دیگر از داده ها $h_k(x) = -1$

شرط بیشینه کردن واریانس

یک تابع هش با بیشینه آنتروپی $H(h_k(x))$ باید واریانس مقادیر هش را بیشینه کند و بر عکس:

$$H(h_k(x)) \leftrightarrow \max var(h_k(x)) \quad (3-31)$$

فرض می کنیم h_k احتمال انتساب مقدار $h_k(x) = 1$ به نقطه داده برابر p باشد و احتمال انتساب $h_k(x) = 0$ برابر $1-p$ مقدار آنتروپی $H(h_k(x))$ از رابطه ی

$$H(h_k(x)) = -p \log(p) - (1-p) \log(1-p) \quad (3-32)$$

خیلی ساده است که نشان دهیم زمانی مقدار آنتروپی بیشینه است که بخش بندی به صورت متعادل انجام شود $p = \frac{1}{2}$ در ادامه نشان می دهیم بخش بندی متعادل منجر به بیشینه کردن واریانس بیت ها می شود.

$$E(x) = M = 2P - 1 \quad (3-33)$$

$$var[h_k(x)] = (h_k(x) - M)^2 = 4P(1-P) \quad (3-34)$$

بیشینه مقدار واریانس در $p = \frac{1}{2}$ اتفاق می افتد که این مقدار به مفهوم بخش بندی متعادل است.

۳-۶ مرتب سازی نتایج بر اساس اهمیت هر بیت

برای بسیاری از متدهای هشینگ مانند LSH، SH، ... ItQH، توابع هش مستقل می باشند. هر یک از توابع هش یک بیت از کد باینری را تولید می کند بنابراین k کد باینری نیز از هم مستقل می باشند. برای این k بیت ما نمی توانیم تصمیم گیری انجام دهیم که کدام یک از این کدها نسبت به دیگر کدها دارای اهمیت بیشتری است بنابراین هنگام محاسبه فاصله ی همینگ به هر یک از بیت ها وزن یکسان تعلق می گیرد [۸]. در واقع می توان هر یک از توابع هشینگ را مانند یک کلاسند دو کلاسه در نظر گرفت در فرآیند کلاس بندی، بخشی از تصاویر که به دسته ی ۱ تعلق می گیرد با کد باینری ۱ و دسته ی دیگر به ۰ تعلق می گیرد با کد باینری ۰ نمایش داده می شود.

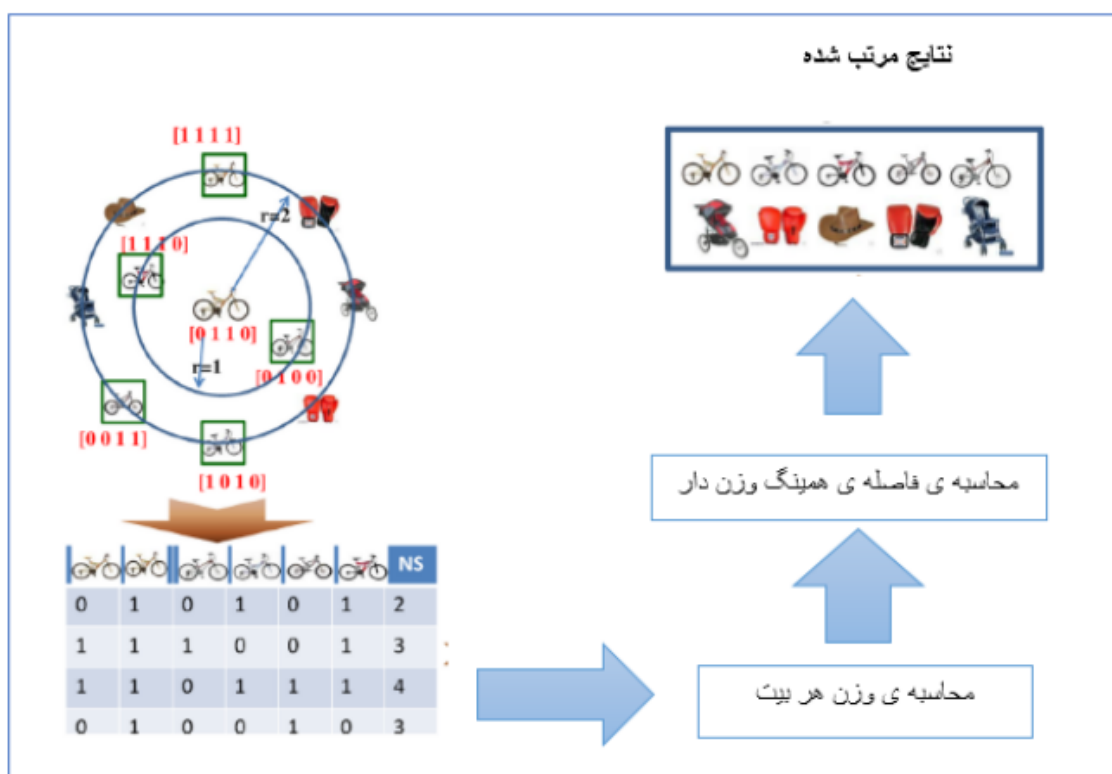
جدول ۳-۲: مثالی از k بیت کد باینری تصویر درخواستی با پنج تصویر

bit	q	I_1	I_2	I_3	I_4	I_5	NS	w
b_1	۰	۱	۱	۱	۰	۱	۱	w_1
b_2	۱	۰	۱	۰	۰	۱	۲	w_2
b_3	۱	۱	۰	۱	۱	۰	۳	w_3
b_4	۰	۰	۰	۰	۱	۰	۴	w_4

اگر برای یک تابع هش تعداد زیادی از تصاویر مشابه دارای کد یکسان با آن در تصویر درخواستی باشند بنابراین این تابع هش از اهمیت بیشتری برخوردار است بنابراین باید وزن بیشتری به کد باینری متعلق به این تابع هش تخصیص داده شود در ادامه نحوه ی محاسبه اهمیت هر بیت توضیح داده می شود. در ابتدا کد باینری تصاویر مشابه به دست آمده با کد باینری تصویر درخواستی به منظور مشخص نمودن اهمیت هر بیت مقایسه می گردد. جدول ۳-۲ مثالی از k بیت کد باینری تصویر درخواستی با پنج تصویر $\{I_i\}_{i=1}^5$ که فاصله ی همینگ آن ها با تصویر درخواستی برابر ۲ است را نمایش می دهد. NS تعداد تصاویری است که دارای کد باینری یکسان با تصویر درخواستی کاربر در آن بیت می باشد و $w = \{w_1, \dots, w_k\}$ وزن آن بیت می باشد.

با مشاهده جدول فوق می توان دریافت که در برخی از بیت ها مانند بیت چهارم تصاویر بیشتری کد هش یکسانی را به اشتراک گذاشته اند و در برخی از بیت ها مانند بیت اول تصاویر کمتری کد هش یکسان را به اشتراک قرار داده اند. برای b_4 مقدار $NS = 4$ ، برای b_3 مقدار $NS = 3$ ، برای b_2 مقدار $NS = 2$ و برای b_1 مقدار $NS = 1$ بنابراین ما به این نتیجه می رسیم که اهمیت بیت چهارم بیشتر از بیت سوم، دوم و اول است بنابراین $w_4 > w_3 > w_2 > w_1$ مرتب سازی جدد نتایج براساس اهمیت هر بیت در شکل زیر نمایش داده شده است. برای تصویر درخواستی تکنیک بازخورد^{۱۹} به منظور یافتن تصاویر مشابه از تصاویر بازگردانده شده اعمال می شود. در ابتدا پایگاه داده تصاویر بر پایه فاصله ی همینگ جست و جو می شود و تصاویری با کمترین فاصله ی همینگ نسبت به تصویر درخواستی استخراج می شود سپس تصاویری که بیشترین شباهت را با تصویر درخواستی کاربر دارند انتخاب می شوند نحوه ی تخصیص وزن به هر یک از هش بیت ها در بخش بعد توضیح داده می شود.

^{۱۹}Feedback technique



شکل ۳-۹: مرتب سازی مجدد نتایج بر اساس اهمیت هر بیت [۸].

۳-۶-۱ محاسبه وزن

فرض کنیم تصویر درخواستی کاربر q و m تعداد تصاویری که بیشترین شباهت را به تصویر درخواستی کاربر دارد و با استفاده از تکنیک هشینگ محاسبه شده است. این m تصویر می توانند بر پایه ی فاصله ی همینگ به صورت یک شکل متحدالمركز نمایش داده شوند و مرکز آن تصویر درخواستی و تصاویر دیگری است که دارای فاصله ی همینگ ۰ می باشند. کد باینری m تصویر به صورت $H = \{H_1, \dots, H_m\}$ که در آن $H_i = \{H_i^{(1)}, H_i^{(2)}, \dots, H_i^{(k)}\}$ می باشد کد باینری تصویر درخواستی به صورت $H_q = \{H_q^{(1)}, H_q^{(2)}, \dots, H_q^{(k)}\}$ می باشد وزن K بیت هش کد تصویر درخواستی به صورت $w = \{w_1, \dots, w_k\}$ می باشد که مقدار اولیه $w_k = 1$ است.

در این بخش مقادیر H_i و H_q بیت به بیت با هم مقایسه می شوند. برای کد باینری K ام اگر $H_i^{(k)} = H_q^{(k)}$ آن گاه وزن w_k افزایش می یابد در غیر این صورت این وزن کاهش می باشد. وزن w به صورت بازگشتی که تعداد تکرار آن برابر با m می باشد به روزرسانی می گردد. برای زامین تکرار w_k با رابطه ی ۳-۳۵ محاسبه می شود:

$$w_k = \begin{cases} \varepsilon w_k & \text{if } H_q^k \neq H_i^k \\ (1 + \varepsilon) w_k & \text{otherwise} \end{cases} \quad 0 < \varepsilon < 1 \quad (3-35)$$

در این رابطه ε یک پارامتر است که درباره ی مقدار آن در بخش مشاهدات بحث می شود. بعد از m تکرار می توانیم مقدار w را به دست آوریم و از آن برای محاسبه ی فاصله ی همینگ وزن دار استفاده کنیم و با مرتب سازی نتایج فاصله ی همینگ وزن دار تصاویری که بیشترین شباهت را به تصویر درخواستی کاربر دارند نمایش داده شود. نحوه ی محاسبه ی این وزن ها در الگوریتم زیر ارائه شده است:

الگوریتم ۳-۱ الگوریتم محاسبه وزن هر بیت در مرتب سازی مجدد بر اساس اهمیت هر بیت.

ورودی: کد باینری تصویر درخواستی کاربر $H_q = \{h_q^{(1)}, h_q^{(2)}, \dots, h_q^{(k)}\}$ ؛ m تصویر که بیشترین شباهت را به تصویر درخواستی کاربر بر اساس فاصله ی همینگ دارند. کد باینری i امین تصویر بازگردانده شده

$H_i = \{H_i^{(1)}, H_i^{(2)}, \dots, H_i^{(k)}\}$ ؛ پارامتر ε .

خروجی: وزن $w = \{w_1, \dots, w_k\}$ برای H_q

۱: مقدار اولیه وزن $w_k = 1, \dots, 1$

۲: برای $m : 1$

۳: برای $k : 1$

۴: آیا k امین بیت H_q و H_i با هم برابر است؟

اگر بله، $w_k = (1 + \varepsilon)w_k$

در غیر این صورت $w_k = \varepsilon w_k$

۵: پایان for

۶: پایان for

فصل ۴

۴-۱ مقدمه

در فصل گذشته به تشریح الگوریتم های به کار رفته در این پایان نامه پرداخته شد. در این فصل برآنیم تا به بیان الگوریتم پیشنهادی و بررسی دقیق تر تکنیک های به کار رفته در این کار بپردازیم. دیاگرام روش پیشنهادی در شکل ۴-۱ نمایش داده شده است. به طور کلی الگوریتم پیشنهادی به چهار مرحله ی کلی استخراج ویژگی های تصاویر، ایجاد کد باینری فشرده، به دست آوردن نزدیک ترین همسایه و مرتب سازی مجدد نتایج تقسیم بندی می شود در این فصل به تشریح کامل هر یک از این مراحل پرداخته می شود.

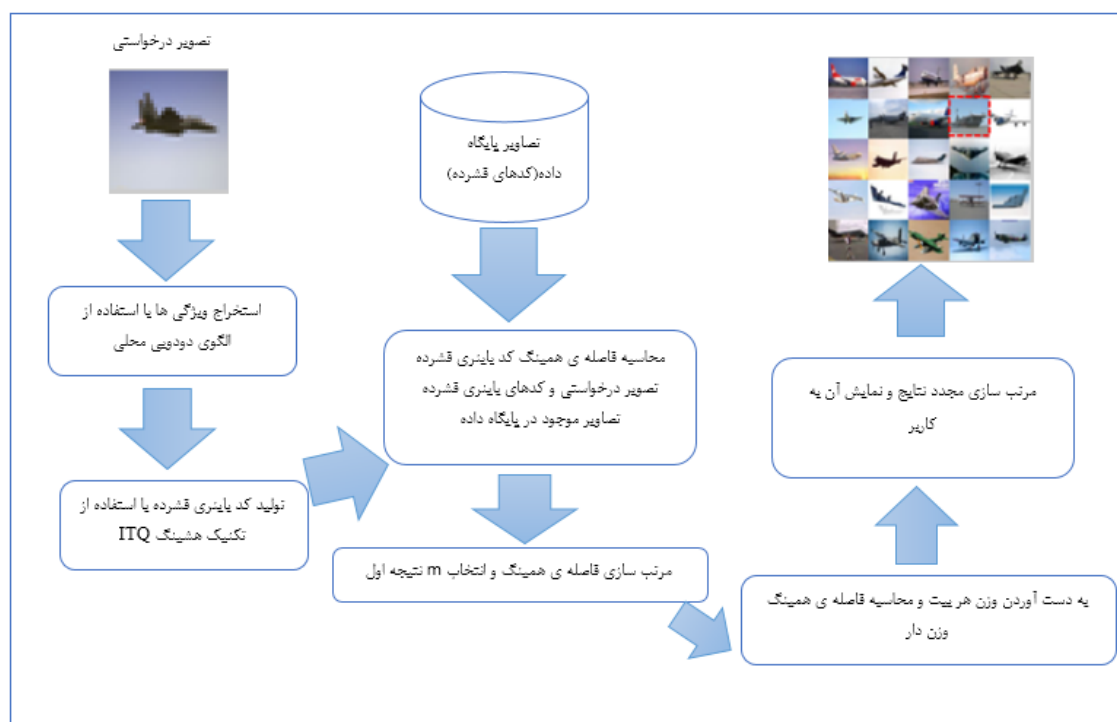
۴-۲ روش پیشنهادی

۴-۲-۱ مرحله اول، استخراج ویژگی های تصویر

به منظور توصیف یک تصویر باید ویژگی های آن استخراج گردد. ویژگی های محلی یک تصویر انتخاب مناسب تری نسبت به ویژگی های سراسری به منظور توصیف یک تصویر می باشد، زیرا ویژگی های محلی قدرت جداسازی بیشتری نسبت به ویژگی های سراسری دارند اما هزینه ی محاسباتی بیشتری نسبت به ویژگی های سراسری دارند. به دلیل حجم بالای تصاویر موجود در پایگاه داده، روشی که به منظور استخراج ویژگی های محلی استفاده می گردد باید دارای ویژگی های ذیل باشد:

- قدرت جداسازی بالا

- هزینه محاسباتی پایین



شکل ۴-۱: دیاگرام روش پیشنهادی.

در میان روش های استخراج ویژگی های محلی، الگوی دودویی محلی به دلیل سرعت تولید، ماهیت دودویی، هزینه محاسباتی پایین و قدرت جداسازی بالا انتخاب مناسبی می باشد. الگوی دودویی محلی به عنوان یک توصیفگر غیر پارامتری مطرح است. از آن جایی که الگوی دودویی محلی هم از مشخصه های آماری و هم ساختاری بافت استفاده می کند یک ابزار قدرتمند به منظور توصیف یک تصویر می باشد. هم چنین این توصیفگر دارای مزایایی چون تحمل پذیری نسبت به تغییرات یکنواخت و سادگی محاسباتی می باشد. با توجه به دلایل ذکر شده، در این مرحله به منظور استخراج ویژگی های تصاویر موجود در پایگاه داده از روش الگوی دودویی محلی استفاده و بردار ویژگی تصاویر تولید گردید.

۴-۲-۲ مرحله دوم، تولید کد باینری فشرده

به دلیل حجم بالای تصاویر موجود در پایگاه داده استفاده از روش تطابق ناشیانه، که در آن یک تصویر با تمام تصاویر موجود در پایگاه داده مقایسه می شود یک روش عملی نمی باشد زیرا این روش نیاز به حافظه ی زیادی دارد و هزینه ی محاسباتی آن بالا است. تکنیک هشینگ روشی مناسب به منظور تخمین زدن نزدیک ترین همسایه می باشد که در آن به ازای هر تصویر یک کد باینری فشرده تولید می شود که در نهایت عملیات مقایسه

تصاویر با استفاده از معیار فاصله ی همینگ انجام می گردد که این امر موجب افزایش سرعت محاسباتی و کاهش حافظه ی مورد نیاز می شود. در این مرحله ما از الگوریتم هشینگ ITQ به منظور تولید کد باینری فشرده استفاده نمودیم این الگوریتم یکی از کارآمدترین الگوریتم ها به منظور کاهش خطای چندی سازی می باشد. بردار ویژگی های تولید شده در مرحله ی قبل به عنوان ورودی این الگوریتم در نظر گرفته می شود و پس از کاهش ابعاد این بردار با استفاده از تکنیک آنالیز اجزای اصلی در داده های بدون برچسب و روش همبستگی کانونی در داده های برچسب گذاری شده، مناسب ترین چرخش برای داده ها با استفاده از روش تجزیه مقادیر منفرد به دست می آید سپس هر یک از داده ها به نزدیک ترین رئوس ابر مکعب باینری تصویر می شوند و کد باینری فشرده برای آن ها به دست می آید.

۴-۲-۳ مرحله سوم، به دست آوردن نزدیک ترین همسایه

به دلیل ماهیت دودویی کدهای تولید شده در مرحله ی قبل، می توان از فاصله ی همینگ به عنوان معیار فاصله به منظور اندازه گیری فاصله ی میان تصویر درخواستی و تصاویر موجود در پایگاه داده استفاده نمود. از آن جایی که این معیار یک عملگر بیتی محسوب می شود از سرعت بالای برخورداری می باشد. این امر در پایگاه داده هایی با مقیاس بالا از اهمیت ویژه ای برخوردار است. در این مرحله فاصله ی کد باینری فشرده تصویر درخواستی با کد باینری فشرده تصاویر موجود در پایگاه داده با استفاده از فاصله ی همینگ به دست می آید و نتایج مرتب سازی می گردد.

۴-۲-۴ مرحله ی چهارم، مرتب سازی مجدد بر اساس اهمیت هر بیت

همان طور که در بخش قبل اشاره گردید، فاصله ی همینگ به منظور اندازه گیری شباهت میان دو کد دودویی مورد استفاده قرار می گیرد. برای k بیت کد دودویی، فاصله ی همینگ یک مقدار صحیح می باشد و حداکثر مقدار آن برابر k می باشد. بنابراین تعداد زیادی از تصاویر فاصله ی همینگ یکسانی با تصویر درخواستی دارند. در معیار اندازه گیری همینگ تمامی بیت ها ارزش یکسانی داشته و با آن ها برخورد یکسانی انجام می شود ولی در عمل برخی از بیت ها از اهمیت بیشتری برخوردار می باشند بدین منظور اهمیت هر بیت در این مرحله با استفاده از الگوریتم مرتب سازی مجدد بر اساس اهمیت هر بیت محاسبه می گردد. بدین منظور ابتدا m تصویر اول به دست آمده از مرحله ی قبل انتخاب می گردد و بیت های کد باینری فشرده آن با بیت های کد باینری تصویر درخواستی مقایسه می گردد و در صورت یکسان بودن به وزن آن بیت افزوده شده و در غیر این صورت از وزن آن کاسته می شود با اجرای این الگوریتم به ازای تصاویر انتخابی، وزن بیت های آن به دست می آید حال با استفاده از معیار فاصله همینگ وزن دار، محاسبات انجام شده و نتایج مرتب می گردد

و n تصویر اول به کاربر نمایش داده می شود.

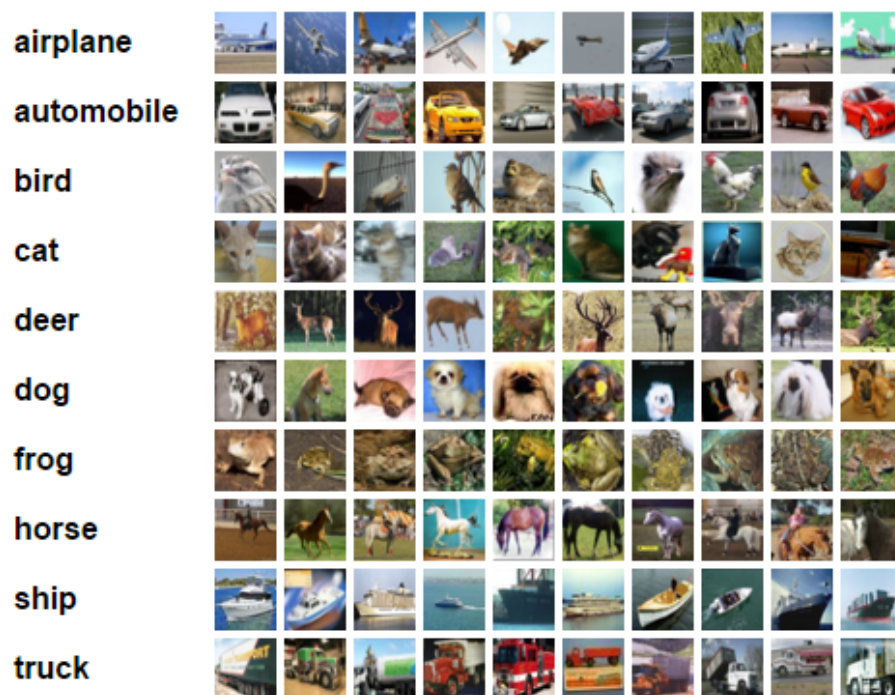
فصل ۵

۵-۱ مقدمه

در این فصل به بررسی مسائلی از قبیل معرفی دادگان مورد استفاده، معیارهای ارزیابی، شرح انجام آزمایشات و ارزیابی و مقایسه روش پیشنهادی با سایر روش های ارائه شده می پردازیم.

۵-۲ مجموعه دادگان

در این پایان نامه به منظور مشاهده و ارزیابی نتایج روش پیشنهادی از دو مجموعه دادگان CIFAR-۱۰ و MNIST استفاده شده است. مجموعه ی CIFAR-۱۰ شامل ۶۰۰۰۰ تصویر رنگی در ابعاد 32×32 می باشد. این تصاویر در ۱۰ کلاس مختلف تقسیم بندی گردیده است. در هر کلاس از این مجموعه ۶۰۰۰ تصویر وجود دارد و این مجموعه شامل ۵۰۰۰۰ داده آموزش و ۱۰۰۰۰ داده تست می باشد. در شکل ۵-۱ نمونه ای از تصاویر موجود در این پایگاه داده نمایش داده شده است. مجموعه دادگان MNIST شامل ارقام ۰ تا ۹ دست نویس با ابعاد 28×28 می باشد این مجموعه شامل ۶۰۰۰۰ داده آموزش و ۱۰۰۰۰ داده تست می باشد در شکل ۵-۲ نمونه ای از تصاویر موجود در پایگاه داده MNIST نمایش داده شده است.



شکل ۵-۱: نمونه ای از تصاویر موجود در پایگاه داده CIFAR-۱۰.



شکل ۵-۲: نمونه ای از تصاویر موجود در پایگاه داده MNIST.

۳-۵ معیارهای ارزیابی

۱-۳-۵ دقت

نسبت تعداد تصاویر مرتبط بازیابی شده به تعداد کل تصاویر بازیابی شده می باشد. به عنوان نمونه اگر ۸ تصویر بازیابی شده باشد و تنها ۴ تصویر متعلق به دسته مورد جست وجو ما باشد دقت آن برابر ۵۰٪ می باشد.

۲-۳-۵ فراخوانی

نسبت تعداد تصاویر مرتبط بازیابی شده به تعداد کل تصاویر موجود در پایگاه داده. به عنوان نمونه اگر ۸ تصویر مرتبط از پایگاه داده ای از تصاویر شامل ۸۰ تصویر بازیابی شود فراخوانی آن برابر ۱۰٪ می باشد.

۳-۳-۵ منحنی دقت در مقابل فراخوانی

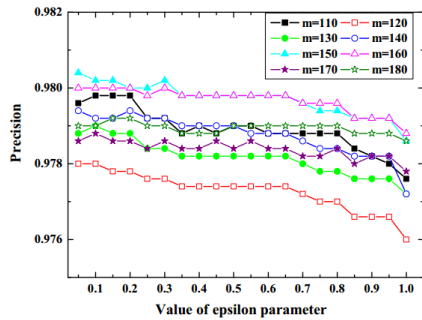
دقت یک امر مهم در فرآیند بازیابی و فراخوانی امری مهم در شناسایی می باشد. این منحنی به منظور نمایش و مقایسه عملکرد الگوریتم ها مورد استفاده قرار می گیرد. این منحنی چگونگی کاهش دقت ، زمانی که کسر بیشتری از تصاویر موجود در پایگاه داده بازیابی می شود را نشان می دهد. منحنی دقت در مقابل فراخوانی ایده آل، دارای مقدار دقت ۱۰۰٪ به ازای تمامی مقادیر فراخوانی می باشد که به معنای بازیابی تمام تصویر مرتبط می باشد.

۴-۵ شرح آزمایشات انجام شده

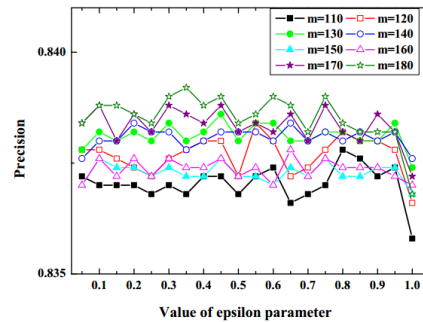
۱-۴-۵ تعیین پارامترهای مساله

تعیین پارامترهای الگوریتم مرتب سازی مجدد بر اساس اهمیت هر بیت

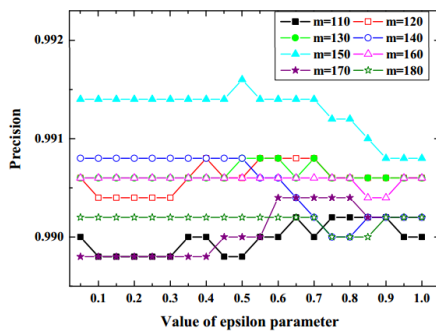
این الگوریتم دارای دو پارامتر m و ε می باشد مقدار مناسب برای این دو پارامتر در هر یک از مجموعه دادگان با اجرای آزمایشات به دست آمد . شکل ۳-۵ تغییرات دقت بازیابی را بر اساس مقادیر مختلف m و ε برای مجموعه دادگان MNIST نشان می دهد. محدوده ε مقادیر بین ۰٫۱ تا ۱ می باشد و مقادیر m بین ۱۱۰ تا ۱۸۰ در نظر گرفته شده است. در شکل (الف) این نمودار برای تعداد ۱۶ بیت، در (ب) برای ۳۲ بیت ، در



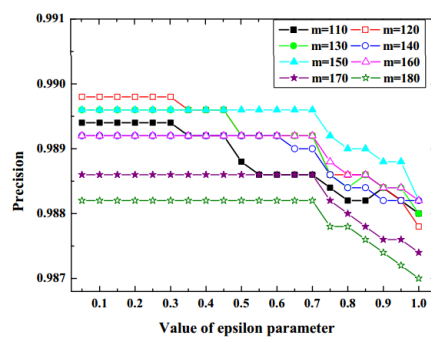
(ب) بیت ۳۲



(آ) بیت ۱۶



(د) بیت ۱۲۸



(ج) بیت ۶۴

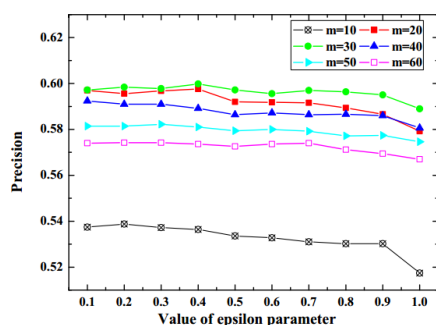
شکل ۳-۵: نمودار دقت بازیابی بر اساس مقادیر مختلف m و ε برای مجموعه دادگان MNIST

(ج) برای بیت ۶۴ و در (د) برای ۱۲۸ بیت کد باینری به دست آمده است. با توجه به نتایج به دست آمده برای مجموعه دادگان MNIST مقدار $m=150$ و $\varepsilon=0.1$ در نظر گرفته شد. شکل ۴-۵ تغییرات دقت بازیابی را بر اساس مقادیر مختلف m و ε برای مجموعه دادگان CIFAR-10 نشان می دهد. محدوده ی مقادیر ε بین ۰٫۱ تا ۱ می باشد و مقادیر m بین ۱۰ تا ۶۰ در نظر گرفته شده است. در شکل (الف) این نمودار برای تعداد ۱۶ بیت، در (ب) برای ۳۲ بیت، در (ج) برای ۶۴ بیت و در (د) برای ۱۲۸ بیت کد باینری به دست آمده است. با توجه به نتایج به دست آمده برای مجموعه دادگان CIFAR-10 مقدار $m=30$ و $\varepsilon=0.3$ در نظر گرفته شد.

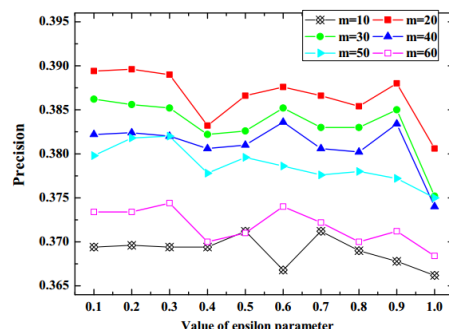
کلیه آزمایشات با پارامترهای به دست آمده در این قسمت انجام شده است.

تعیین پارامتر الگوریتم هشینگ ITQ

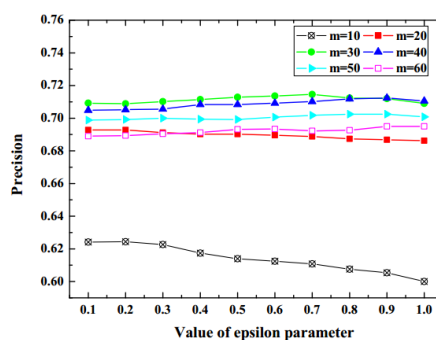
به منظور به دست آوردن تعداد تکرار مناسب برای الگوریتم هشینگ ITQ، به ازای تعداد تکرار ۵۰، ۱ و ۱۵۰ این الگوریتم را تکرار نمودیم و منحنی دقت در مقابل فراخوانی را به دست آوردیم. طبق نتایج به دست آمده



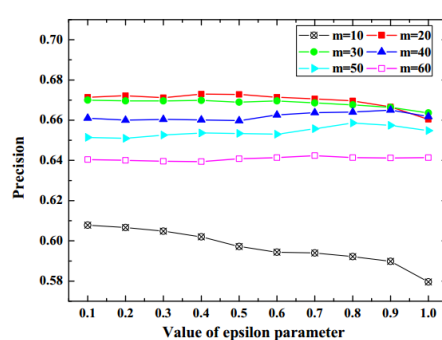
(ب) ۳۲ بیت



(آ) ۱۶ بیت



(د) ۱۲۸ بیت



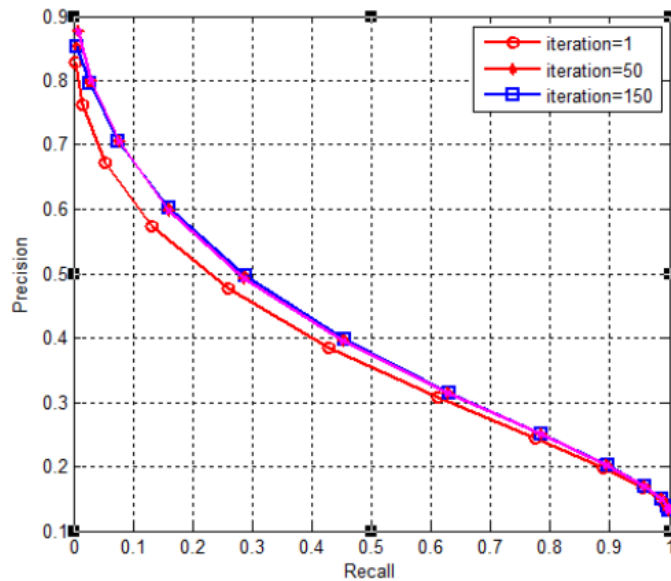
(ج) ۶۴ بیت

شکل ۴-۵: نمودار دقت بازیابی بر اساس مقادیر مختلف m و ϵ برای مجموعه داده‌های CIFAR-۱۰

در شکل ۵-۵ بهترین تکرار برای این الگوریتم ۵۰ می باشد و آزمایشات در این بخش با تعداد تکرار ۵۰ به انجام رسیده است.

۴-۵-۲ تأثیر ویژگی الگوی دودویی محلی در بهبود عملکرد الگوریتم

بردار ویژگی های تصویر به عنوان ورودی الگوریتم هشینگ ITQ، نقشی کلیدی در عملکرد این الگوریتم دارد. در شکل ۵-۶ نموداری ارائه گردیده است که مقایسه ای را بین عملکرد الگوی دودویی محلی و توصیفگر سراسری GIST به عنوان ورودی الگوریتم هشینگ ITQ انجام داده است. همان طور که در نمودار ۵-۶ مشاهده می شود عملکرد این الگوریتم با به کارگیری الگوی دودویی محلی به عنوان ورودی بهبود یافته است. این الگو به دلیل ماهیت دودویی خود، از یک سو هزینه محاسبات را کاهش داده است و از سوی دیگر به دلیل کارایی این روش در استخراج ویژگی های تصویر، موجب افزایش دقت بازیابی گردیده است. با بهره گیری از این روش با تعداد هش بیت کمتری سیستم به دقت بالاتری دست یافته است که این امر در سیستم های بازیابی تصویر بر اساس محتوا در مقیاس بالا بسیار پراهمیت است.



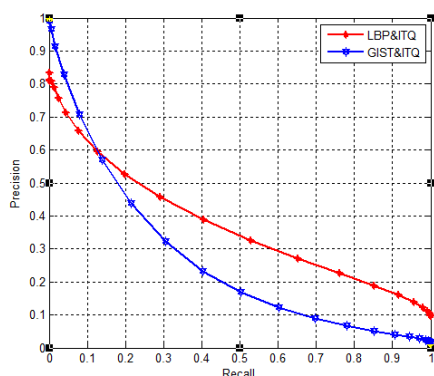
شکل ۵-۵: منحنی دقت در مقابل فراخوانی به ازای تعداد تکرار مختلف الگوریتم هشینگ ITQ.

۳-۴-۵ مقایسه ی روش پیشنهادی با روش های دیگر

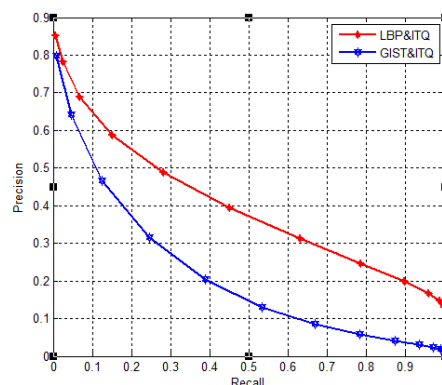
در این بخش الگوریتم پیشنهادی با سه الگوریتم هشینگ پایه ای دیگر که روند هشینگ $B = \text{sgn}(X\tilde{W})$ را دنبال می کند مقایسه شده است که در آن ماتریس تصویر $\tilde{W}$ با روش های مختلف به دست می آید:

- هشینگ حساس به همسایگی: $\tilde{W}$ یک ماتریس تصادفی گاوسی است .
- چرخش تصادفی: در این روش $\tilde{W} = WR$ می باشد که در آن W ماتریس جهت های PCA است و R یک ماتریس متعامد تصادفی است که ایده ی روش ITQ برگرفته از این روش می باشد.
- کرنل های ثابت به جابه جایی: این روش بر پایه ی ویژگی های فوریه ی تصادفی برای تخمین زدن کرنل های گاوسی می باشد.

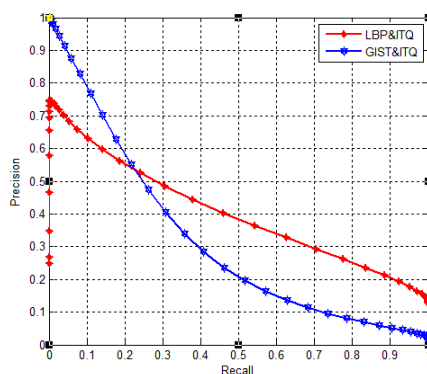
شکل ۵-۷ به مقایسه عملکرد الگوریتم پیشنهادی در مجموعه دادگان CIFAR-۱۰ و شکل ۵-۸ به مقایسه عملکرد الگوریتم پیشنهادی در مجموعه دادگان MNIST پرداخته است و نمودار دقت در مقابل فراخوانی را برای روش های ذکر شده و روش پیشنهادی ترسیم نموده است. همان طور که در این نمودار مشاهده می شود دقت الگوریتم پیشنهادی از روش های ارائه شده بالاتر است . در این روش با تعداد هش بیت کمتر، سیستم به دقت بالاتری دست یافته است .



(ب) منحنی دقت براساس فراخوانی @۳۲ بیت



(آ) منحنی دقت براساس فراخوانی @۱۶ بیت



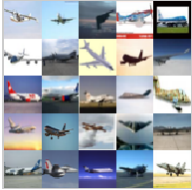
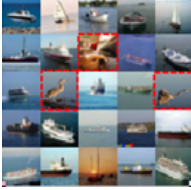
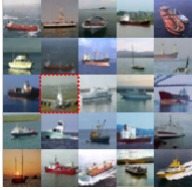
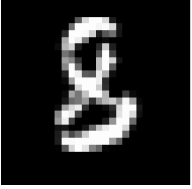
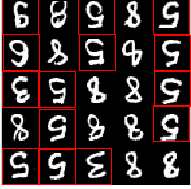
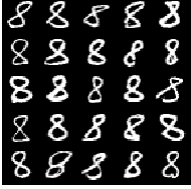
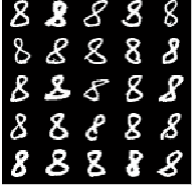
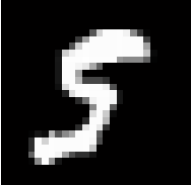
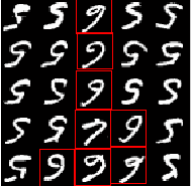
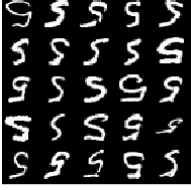
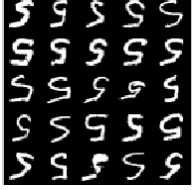
(ج) منحنی دقت براساس فراخوانی @۶۴ بیت

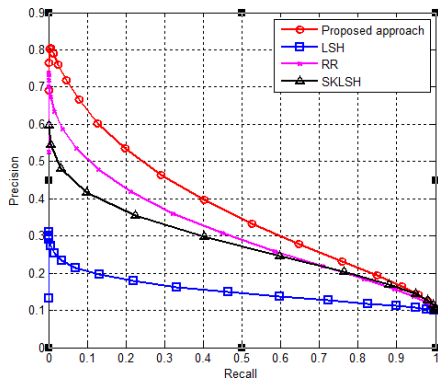
شکل ۵-۶: تاثیر الگوی دودویی محلی به عنوان ورودی به الگوریتم هشینگ ITQ

۴-۴-۵ نمایش خروجی به کاربر

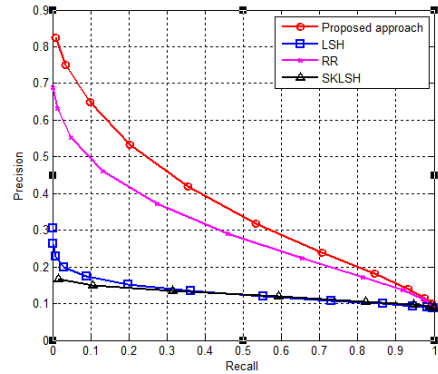
در این بخش با دریافت یک تصویر به عنوان تصویر درخواستی کاربر، با اجرای الگوریتم پیشنهادی، ۲۵ تصویر نخست بازیابی شده را به کاربر نمایش داده ایم. در جدول ۵-۱ نتایج به دست آمده با استفاده از تکنیک آنالیز اجزای اصلی و تحلیل همبستگی کانونی و دقت محاسباتی هر یک نمایش داده شده است. همان طور که مشاهده می شود دقت الگوریتم پس از اعمال مرتب سازی مجدد بر اساس اهمیت هر بیت افزایش یافته است. هم چنین نتایج بر اساس بیشترین شباهت به تصویر درخواستی نمایش داده شده است.

جدول ۵-۱: ۲۵ تصویر نخست بازیابی شده

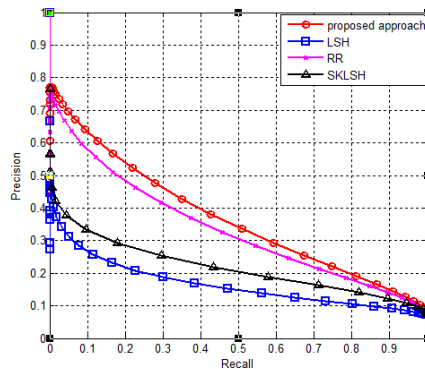
تصویر درخواستی	نتایج با استفاده از PCA	نتایج با استفاده از CCA	نتایج پس از مرتب سازی مجدد
			
دقت	۶۴٪	۹۶٪	۱۰۰٪
			
دقت	۶۰٪	۸۸٪	۹۶٪
			
دقت	۴۴٪	۱۰۰٪	۱۰۰٪
			
دقت	۷۲٪	۱۰۰٪	۱۰۰٪



(ب) منحنی دقت براساس فراخوانی @۳۲ بیت



(آ) منحنی دقت براساس فراخوانی @۱۶ بیت

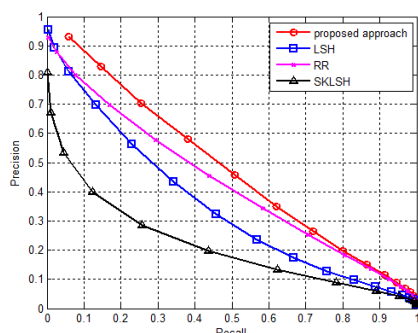


(ج) منحنی دقت براساس فراخوانی @۶۴ بیت

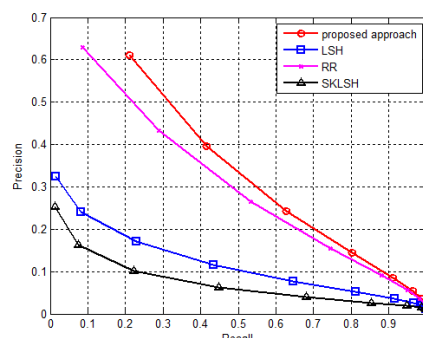
شکل ۵-۷: مقایسه عملکرد الگوریتم پیشنهادی با روش های دیگر در مجموعه دادگان CIFAR-۱۰

۵-۵ نتیجه گیری

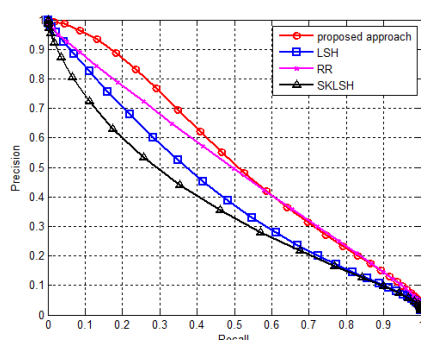
این پایان نامه به ارائه ی روشی کارا به منظور بازیابی تصاویر بر اساس محتوا در مقیاس بالا پرداخته است. این روش بر پایه ی الگوی دودویی محلی و الگوریتم هشینگ ITQ می باشد. در این کار به بررسی نقش کلیدی بردار ویژگی تصویر به عنوان ورودی الگوریتم هشینگ پرداخته است. همان طور که در نتایج آزمایشات در بخش ۵-۴-۲ نشان داده شد به کارگیری بردار ویژگی مناسب موجب بهبود عملکرد الگوریتم هشینگ گردیده است. همان طور که پیش تر مطرح گردید به منظور استخراج ویژگی های تصویر در روش پیشنهادی از الگوی دودویی محلی استفاده شده است. این الگو یکی از قدرتمندترین توصیفگرها برای ارائه ساختارهای محلی به شمار می آید. به دلیل مزایایی چون تحمل پذیری نسبت به تغییرات یکنواخت و سادگی محاسباتی، به صورت موفقیت آمیزی در بسیاری از حوزه های آنالیز تصویر به کار برده می شود. دستیابی به مجموعه ی کوچک ویژگی ها، از متمایزترین خصوصیات الگوی دودویی محلی است که باعث عملکرد بهتر و کاهش ابعاد آن می



(ب) منحنی دقت براساس فراخوانی @۳۲ بیت



(آ) منحنی دقت براساس فراخوانی @۱۶ بیت



(ج) منحنی دقت براساس فراخوانی @۶۴ بیت

شکل ۵-۸: مقایسه عملکرد الگوریتم پیشنهادی با روش های دیگر در مجموعه داده گان MNIST

گردد.

به دلیل بالا بودن حجم تصاویر موجود در پایگاه داده انتخاب روشی کارا برای استخراج ویژگی های یک تصویر بسیار حائز اهمیت می باشد. زیرا روش انتخابی در عین کارایی در استخراج ویژگی ها باید هزینه ی محاسباتی پایینی نیز داشته باشد. الگوی دودویی محلی به دلیل آن که یک روش غیر پارامتری برای استخراج ویژگی ها می باشد از سادگی محاسباتی نیز برخوردار است. بدین ترتیب با بهره گیری از این روش می توانیم به ویژگی های مناسب با هزینه ی پایین دست یابیم.

در ادامه به منظور تخمین نزدیک ترین همسایه از تکنیک ITQ استفاده گردیده است. این روش نیز یکی از روش های ساده و کارا به منظور تخمین نزدیک ترین همسایه می باشد. نتایج آزمایشات در بخش ۳-۴-۵ نشان داده است که روش پیشنهادی نسبت به دیگر روش های ارائه شده از کارایی بالاتری برخوردار است. این بدان مفهوم می باشد که ما به استفاده از تعداد کد هس کمتر توانسته ایم به دقت بالاتری دست یابیم. هس کردن یک تصویر با تعداد کد کمتر منجر به استفاده از حافظه ی کمتر برای نمایش یک تصویر می شود و از سوی دیگر هزینه ی محاسباتی به منظور مقایسه ی دو تصویر را کاهش می دهد. تحقق این امر در پایگاه داده

هایی با ابعاد بالا از اهمیت زیادی برخوردار است زیرا نیاز به تکنیک هایی که حافظه ی مصرفی و هزینه ی محاسباتی را کاهش دهد احساس می گردد.

از آن جا که اهمیت هر بیت در کد باینری تولید شده یکسان نمی باشد با بهره گیری از تکنیک مرتب سازی مجدد نتایج بر اساس اهمیت هر بیت، وزن هر یک از بیت ها به دست آمده و با بهره گیری از معیار فاصله ی همینگ وزن دار نتایج به دست آمده مرتب سازی می گردد. این گام در الگوریتم پیشنهادی به موجب نمایش نتایج بر اساس بیشترین شباهت به تصویر درخواستی در نظر گرفته شده است.

با توجه به نتایج به دست آمده در بخش آزمایشات دقت روش پیشنهادی نسبت به روش های موجود بالاتر است. این روش با به کارگیری تعداد هش کد کمتر توانسته به این دقت دست یابد که این امر نتیجه ای مطلوب به منظور بازیابی تصاویر بر اساس محتوا در مقیاس بالا می باشد.

۵-۶ کارهای آینده

به دنبال مطالعه در یک حوزه، بررسی مشکلات موجود در آن و ارائه راه حل برای آن سپس تست و آزمایش راه حل پیشنهادی، با مشکلات موجود در این مسیر بیشتر آشنا شده و راه حل هایی به ذهن محقق خطوط می کند که به دلیل محدودیت زمانی، امکان بررسی آن برای فرد وجود ندارد. هدف از این بخش ارائه مشکلات این مسیر و راه حل های محقق می باشد به گونه ای که انگیزه ی ادامه این راه را برای خود نویسنده یا خوانندگان فراهم آورد. در این بخش به ارائه ی چندین نمونه که می تواند در تکمیل این کار مؤثر واقع شود می پردازم.

اغلب روش های بافت محور به منظور استخراج ویژگی های یک تصویر حساسیت زیادی نسبت به تغییرات نورپردازی شدید داشته و هم چنین نسبت به تصاویر مبهم و مختل شده حساسیت نشان می دهند به منظور رفع مشکلات ارائه شده می توانیم قبل از مرحله ی استخراج ویژگی ها یک مرحله پیش پردازش روی تصویر انجام داده و سپس خروجی این مرحله را به عنوان ورودی الگوی دودویی محلی در نظر بگیریم. البته در پایگاه دادهایی که در این کار مورد استفاده قرار گرفته است مشکلات فوق وجود ندارد ولی به منظور عمومی سازی این الگوریتم می توان این مرحله را به مراحل موجود اضافه نمود.

به منظور بهبود عملکرد الگوی دودویی محلی در سالهای اخیر تغییرات زیادی به این الگوریتم اعمال شده است. این تغییرات بر جنبه های مختلف الگوی دودویی محلی اصلی تمرکز کرده است که به منظور بهبود قدرت جداسازی، افزایش تحمل در برابر تغییرات، نحوه ی انتخاب همسایه، گسترش به داده های سه بعدی و ترکیب با دیگر روش ها است که لیستی از این تغییرات در جدول ۳-۱ آورده شده است. به عنوان یکی از کارهای پیشنهادی برای آینده می توان به کارگیری این روش ها در الگوریتم پیشنهادی در مرحله ی نخست اشاره کرد و به بررسی و مقایسه نتایج به دست آمده با این روش ها پرداخت.

به کارگیری پایگاه داده هایی با حجم تصاویر بیشتر که شامل میلیون ها تصویر باشد پیشنهاد دیگر ارائه شده در این بخش است. می توان عملکرد روش پیشنهادی را در این پایگاه داده ها بررسی و مقایسه نمود.

مراجع

- [1] Ahonen, T., Hadid, A., and Pietikainen, M. Face description with local binary patterns: Application to face recognition. *Pattern Analysis and Machine Intelligence, IEEE Transactions on* 28, 12 (2006), 2037–2041.
- [2] Ahonen, T., Matas, J., He, C., and Pietikäinen, M. Rotation invariant image description with local binary pattern histogram fourier features. in *Image Analysis*. Springer, 2009, pp. 61–70.
- [3] Ahonen, T., and Pietikäinen, M. Soft histograms for local binary patterns. in *Proceedings of the Finnish signal processing symposium, FINSIG* (2007), volume 5, p. 1.
- [4] Andoni, A., and Indyk, P. Near-optimal hashing algorithms for approximate nearest neighbor in high dimensions. in *Foundations of Computer Science, 2006. FOCS'06. 47th Annual IEEE Symposium on* (2006), IEEE, pp. 459–468.
- [5] Bai, G., Zhu, Y., and Ding, Z. A hierarchical face recognition method based on local binary pattern. in *2008 congress on image and signal processing* (2008), IEEE, pp. 610–614.
- [6] Fehr, J. Rotational invariant uniform local binary patterns for full 3d volume texture analysis. in *Finnish signal processing symposium (FINSIG)* (2007).
- [7] Fergus, R., Weiss, Y., and Torralba, A. Semi-supervised learning in gigantic image collections. in *Advances in neural information processing systems* (2009), pp. 522–530.
- [8] Fu, H., Kong, X., and Wang, Z. Binary code reranking method with weighted hamming distance. *Multimedia Tools and Applications* (2014), 1–18.
- [9] Gong, Y., and Lazebnik, S. Iterative quantization: A procrustean approach to learning binary codes. in *Computer Vision and Pattern Recognition (CVPR), 2011 IEEE Conference on* (2011), IEEE, pp. 817–824.

- [10] Gong, Y., Lazebnik, S., Gordo, A., and Perronnin, F. Iterative quantization: A procrustean approach to learning binary codes for large-scale image retrieval. *Pattern Analysis and Machine Intelligence, IEEE Transactions on* 35, 12 (2013), 2916–2929.
- [11] Guo, Z., Zhang, L., and Zhang, D. A completed modeling of local binary pattern operator for texture classification. *Image Processing, IEEE Transactions on* 19, 6 (2010), 1657–1663.
- [12] He, L., Zou, C., Zhao, L., and Hu, D. An enhanced lbp feature based on facial expression recognition. in *Engineering in Medicine and Biology Society, 2005. IEEE-EMBS 2005. 27th Annual International Conference of the* (2006), IEEE, pp. 3300–3303.
- [13] Heikkilä, M., Pietikäinen, M., and Schmid, C. Description of interest regions with local binary patterns. *Pattern recognition* 42, 3 (2009), 425–436.
- [14] Huang, D., Shan, C., Ardabilian, M., Wang, Y., and Chen, L. Local binary patterns and its application to facial image analysis: a survey. *Systems, Man, and Cybernetics, Part C: Applications and Reviews, IEEE Transactions on* 41, 6 (2011), 765–781.
- [15] Huang, D., Wang, Y., and Wang, Y. A robust method for near infrared face recognition based on extended local binary pattern. in *Advances in Visual Computing*. Springer, 2007, pp. 437–446.
- [16] Huang, D., Zhang, G., Ardabilian, M., Wang, Y., and Chen, L. 3d face recognition using distinctiveness enhanced facial representations and local feature hybrid matching. in *Biometrics: Theory Applications and Systems (BTAS), 2010 Fourth IEEE International Conference on* (2010), IEEE, pp. 1–7.
- [17] Huang, Y., Wang, Y., and Tan, T. Combining statistics of geometrical and correlative features for 3d face recognition. in *BMVC* (2006), Citeseer, pp. 879–888.
- [18] Indyk, P., and Motwani, R. Approximate nearest neighbors: towards removing the curse of dimensionality. in *Proceedings of the thirtieth annual ACM symposium on Theory of computing* (1998), ACM, pp. 604–613.
- [19] Jaakkola, T., Haussler, D., et al. Exploiting generative models in discriminative classifiers. *Advances in neural information processing systems* (1999), 487–493.
- [20] Jégou, H., Douze, M., Schmid, C., and Pérez, P. Aggregating local descriptors into a compact image representation. in *Computer Vision and Pattern Recognition (CVPR), 2010 IEEE Conference on* (2010), IEEE, pp. 3304–3311.

- [21] Jégou, H., Perronnin, F., Douze, M., Sanchez, J., Perez, P., and Schmid, C. Aggregating local image descriptors into compact codes. *Pattern Analysis and Machine Intelligence, IEEE Transactions on* 34, 9 (2012), 1704–1716.
- [22] Jin, H., Liu, Q., Lu, H., and Tong, X. Face detection using improved lbp under bayesian framework. in *Multi-Agent Security and Survivability, 2004 IEEE First Symposium on* (2004), IEEE, pp. 306–309.
- [23] Kamel, A., Mahdi, Y. B., and Hussain, K. F. Multi-bin search: improved large-scale content-based image retrieval. in *Image Processing (ICIP), 2013 20th IEEE International Conference on* (2013), IEEE, pp. 2597–2601.
- [24] Lei, Z., Liao, S., He, R., Pietikainen, M., and Li, S. Z. Gabor volume based local binary pattern for face representation and recognition. in *Automatic Face & Gesture Recognition, 2008. FG'08. 8th IEEE International Conference on* (2008), IEEE, pp. 1–6.
- [25] Liao, S., Zhu, X., Lei, Z., Zhang, L., and Li, S. Z. Learning multi-scale block local binary patterns for face recognition. in *Advances in Biometrics*. Springer, 2007, pp. 828–837.
- [26] Liao, S., and Chung, A. C. Face recognition by using elongated local binary patterns with average maximum distance gradient magnitude. in *Computer Vision–ACCV 2007*. Springer, 2007, pp. 672–679.
- [27] Lowe, D. G. Distinctive image features from scale-invariant keypoints. *International journal of computer vision* 60, 2 (2004), 91–110.
- [28] Mikolajczyk, K., and Schmid, C. A performance evaluation of local descriptors. *Pattern Analysis and Machine Intelligence, IEEE Transactions on* 27, 10 (2005), 1615–1630.
- [29] Ojala, T., Pietikäinen, M., and Harwood, D. A comparative study of texture measures with classification based on featured distributions. *Pattern recognition* 29, 1 (1996), 51–59.
- [30] Ojala, T., Pietikäinen, M., and Mäenpää, T. Multiresolution gray-scale and rotation invariant texture classification with local binary patterns. *Pattern Analysis and Machine Intelligence, IEEE Transactions on* 24, 7 (2002), 971–987.
- [31] Oliva, A., and Torralba, A. Modeling the shape of the scene: A holistic representation of the spatial envelope. *International journal of computer vision* 42, 3 (2001), 145–175.

- [32] Paulhac, L., Makris, P., and Ramel, J.-Y. Comparison between 2d and 3d local binary pattern methods for characterisation of three-dimensional textures. in *Image Analysis and Recognition*. Springer, 2008, pp. 670–679.
- [33] Perronnin, F., and Dance, C. Fisher kernels on visual vocabularies for image categorization. in *Computer Vision and Pattern Recognition, 2007. CVPR'07. IEEE Conference on* (2007), IEEE, pp. 1–8.
- [34] Raginsky, M., and Lazebnik, S. Locality-sensitive binary codes from shift-invariant kernels. in *Advances in neural information processing systems* (2009), pp. 1509–1517.
- [35] Ruiz-del Solar, J., and Quinteros, J. Illumination compensation and normalization in eigenspace-based face recognition: A comparative study of different pre-processing approaches. *Pattern Recognition Letters* 29, 14 (2008), 1966–1979.
- [36] Shan, S., Zhang, W., Su, Y., Chen, X., and Gao, W. Ensemble of piecewise fda based on spatial histograms of local (gabor) binary patterns for face recognition. in *Pattern Recognition, 2006. ICPR 2006. 18th International Conference on* (2006), volume 4, IEEE, pp. 606–609.
- [37] Singh, R., Vatsa, M., and Noore, A. Integrated multilevel image fusion and match score fusion of visible and infrared face images for robust face recognition. *Pattern Recognition* 41, 3 (2008), 880–893.
- [38] Sivic, J., and Zisserman, A. Video google: A text retrieval approach to object matching in videos. in *Computer Vision, 2003. Proceedings. Ninth IEEE International Conference on* (2003), IEEE, pp. 1470–1477.
- [39] Tan, X., and Triggs, B. Enhanced local texture feature sets for face recognition under difficult lighting conditions. in *Analysis and Modeling of Faces and Gestures*. Springer, 2007, pp. 168–182.
- [40] Tan, X., and Triggs, B. Fusing gabor and lbp feature sets for kernel-based face recognition. in *Analysis and Modeling of Faces and Gestures*. Springer, 2007, pp. 235–249.
- [41] Tsaneva, M., Krezhova, D., and Yanev, T. Development and testing of a statistical texture model for land cover classification of the black sea region with modis imagery. *Advances in Space Research* 46, 7 (2010), 872–878.

- [42] Vasconcelos, N. From pixels to semantic spaces: Advances in content-based image retrieval. *Computer*, 7 (2007), 20–26.
- [43] Wang, J., Shen, H. T., Song, J., and Ji, J. Hashing for similarity search: A survey. *arXiv preprint arXiv:1408.2927* (2014).
- [44] Weiss, Y., Torralba, A., and Fergus, R. Spectral hashing. in *Advances in neural information processing systems* (2009), pp. 1753–1760.
- [45] Wolf, L., Hassner, T., and Taigman, Y. Descriptor based methods in the wild. in *Workshop on Faces in 'Real-Life' Images: Detection, Alignment, and Recognition* (2008).
- [46] Yan, S., Wang, H., Tang, X., and Huang, T. Exploring feature descriptors for face recognition. in *Acoustics, Speech and Signal Processing, 2007. ICASSP 2007. IEEE International Conference on* (2007), volume 1, IEEE, pp. I–629.
- [47] Yang, H., and Wang, Y. A lbp-based face recognition method with hamming distance constraint. in *Image and Graphics, 2007. ICIG 2007. Fourth International Conference on* (2007), IEEE, pp. 645–649.
- [48] Zhang, L., Chu, R., Xiang, S., Liao, S., and Li, S. Z. Face detection based on multi-block lbp representation. in *Advances in biometrics*. Springer, 2007, pp. 11–18.
- [49] Zhang, W., Shan, S., Gao, W., Chen, X., and Zhang, H. Local gabor binary pattern histogram sequence (lgbphs): A novel non-statistical model for face representation and recognition. in *Computer Vision, 2005. ICCV 2005. Tenth IEEE International Conference on* (2005), volume 1, IEEE, pp. 786–791.
- [50] Zhao, G., and Pietikainen, M. Dynamic texture recognition using local binary patterns with an application to facial expressions. *Pattern Analysis and Machine Intelligence, IEEE Transactions on* 29, 6 (2007), 915–928.
- [51] Zhao, G., and Pietikäinen, M. Experiments with facial expression recognition using spatiotemporal local binary patterns. in *Multimedia and Expo, 2007 IEEE International Conference on* (2007), IEEE, pp. 1091–1094.
- [52] Zhao, H., Wang, Z., and Liu, P. The ordinal relation preserving binary codes. *Pattern Recognition* (2015).

واژه‌نامه انگلیسی به فارسی

Concept based image retrieval	بازیابی تصویر بر اساس متن
Content based image retrieval	بازیابی تصویر بر اساس محتوا
Key word	کلمه کلیدی
Caption	زیرنویس
Description text	نوشته توصیفی
Analyzed	تجزیه
Polysemy	چندمعنایی
Synonymy	ترادف
Query image	تصویر درخواستی
Feature extractor	استخراج ویژگی ها
Features vector	پایگاه داده ویژگی ها
Similarity search	جست و جو شباهت
Indexing and retrieval	اندیس گذاری و بازیابی
www	صفحات وب
Medical image	تصاویر پزشکی
Brute force matching	تطابق ناشیانه
Practical	عملی
Object recognition	شناسایی اشیاء
Finding partial image duplicate on the web	کپی های غیر مجاز تصاویر در صفحات وب
Searching indivisual video frame	جست و جو کردن فریم های ویدئو
Image classification	دسته بندی کردن تصاویر
Image Representation	نمایش تصویر
Traditional	قدیمی
Image key point descriptor	نقاط کلیدی تصویر
Interest point descriptors	توصیفگرهای نقاط مورد علاقه
Well defined position	خوش تعریف

Scaling	مقیاس
Rotate	چرخش
Reproducibility	تکرارپذیری
Numerical representation	نمایش عددی
Keypoint descriptor vecotr	توصیفگر نقاط کلیدی
Feature vector	بردار ویژگی
Features	ویژگی‌ها
Local keypoint of interest	نقاط کلیدی محلی مورد علاقه
Transformation	جابه جایی‌ها
Local interest region	نواحی مورد علاقه
Distinctive	جداساز
Nearest neighbor search	جستجوی نزدیکترین همسایه
Proximity problems	مسائل مجاورت
Exact nearest neighbor search	نزدیکترین همسایه ی دقیق
Linear scan	بررسی خطی
Curse of dimensionality	مصیبت بعد
Approximation	تقریب زدن
NP-hard	ان پی هارد
Collision	برخورد
Divide and conquer approach	رویکرد تقسیم و تخییر
Kd-trees	درخت های کی دی
Multidimensional spaces	فضاهای چندبعدي
Balanced binary search tree	درخت جستجو دودویی متوازن
Original LBP operator	عملگر الگوی دودویی محلی اصلی
Threshold	آستانه سازی
Basic	پایه
Feature Detector	آشکار ساز ویژگی
Micro pattern	ریزالگوها
Robust	تحمل
Hash functions	تابع های هش
Bounded distortion	تغییرات محدود
Embedding	جاسازی
Random projections	نگاشت های تصادفی
Hash bins	جعبه ی هش

Sparse.....	تَنک
Shift invariant kernels.....	کرنل های ثابت به جابه جایی
Iterative quantization.....	چندی سازی تکرار شونده
Random Fourier features	ویژگی های تصادفی فوریه
Ordinary Correlation Analysis (OCA)	تحلیل همبستگی معمولی
Coordinating System.....	سیستم مختصاتی
Redundancy Coefficients	ضرایب افزونگی
Redundant.....	زائد
Bi-diagonalization	قطری سازی
Principal component analysis	تکنیک آنالیز اجزای اصلی
Canonical correlation analysis.....	تحلیل فاصله ی کانونی
encoding	رمزنگاری
Quantization loss	فقدان چندی سازی
Reranking based on bit importance weight	مرتب سازی نتایج بر اساس اهمیت هر بیت
Feedback technique.....	تکنیک بازخور

Abstract:

Hashing technique is an efficient algorithm for approximating nearest neighbor in large-scale image dataset. Learning binary code is a key step to improve its performance and still an ongoing challenge. The inputs of hashing affects its performance. Proposed method in this thesis improve the efficiency of learning binary code by using suitability input. In this method, image features are extracted by local binary pattern then, they are used as input vector of iterative quantization (ITQ) that leads to compact codes, less memory and computational requirement. The reasons behind these achievements are the binary nature and high efficiency in feature generation of local binary pattern. Finally in order to rerank results based on most similarity to query image, bit importance reranking method is used and then big weights are assigned to important bits and small weights are assigned to minor bits and weighted hamming distance is used to compute similarity between query image and top results. To evaluate performance of proposed method we use CIFAR-10 and MNIST as dataset and precision vs recall as evaluation criteria. The simulations compare the new algorithm with three state of the art and along the line algorithms from three points of view; the hashing code size, memory space and computational cost, and the results demonstrate the effectiveness of the new approach.

Keywords: large-scale image retrieval, Local Binary Pattern, iterative quantization, image hashing algorithm



Kharazmi University
Computer Engineering Department

Large-scale content based image retrieval using local binary patterns and iterative quantization

**A Thesis Submitted in Partial Fulfillment of the Requirement for the Degree
of Master of Science in Computer Engineering**

By:

Mona shakerdonyavi

Supervisor:

Dr.Jamshid Shanbehzadeh

August 2015