



Computational Intelligence in Electrical Engineering
Vol. 16, No. 4, 2026
pp. 29-42
[Research Paper](#)

Clustering with Apollonius kernel Based on Laplacian Matrix and k-means

Shahin Pourbahrami

¹ 1Assistant Professor, Department of Computer Engineering, National University of Skills, Tehran, Iran

Abstract:

In this study, a novel clustering algorithm based on the Apollonius kernel is presented. By utilizing the graph Laplacian matrix and the k-means algorithm, this method outperforms traditional approaches in separating clusters. In the first step, the data are normalized to standardize the feature scales. Then, by computing the Euclidean distance between points and applying the Apollonius kernel, the pairwise similarity matrix is constructed. Using the k-nearest neighbor graph, only the connections of the nearest neighbors are retained. Next, the graph Laplacian matrix is obtained as the difference between the degree matrix and the similarity matrix. This matrix represents the relationships between data points. By performing eigenvalue decomposition on the Laplacian matrix, the eigenvectors corresponding to the smallest eigenvalues are selected. The data are then projected onto a reduced space where the clusters are more clearly distinguishable. Then, the k-means algorithm is applied to the reduced data to determine the clusters. Finally, accuracy and other standard clustering validation metrics are used to evaluate the quality of the clustering. Experimental results show that this algorithm performs better than traditional methods, especially on datasets with nonlinear structures. Furthermore, the use of the Apollonius kernel in constructing the similarity graph improves the accuracy of identifying local structures in the data.

Keywords: Clustering, Apollonius kernel, Laplacian Matrix, k-Means, Similarity Matrix, Eigenvalues.



This is an open access article under the CC BY-NC-ND/4.0/ License (<https://creativecommons.org/licenses/by-nc-nd/4.0/>).



<https://doi.org/10.22108/isee.2026.145741.1748>

خوشه‌بندی با هسته آپولونیوس بر اساس ماتریس لاپلاسین و k میانگین

شهین پوربهرامی *

استادیار گروه مهندسی کامپیوتر، دانشگاه ملی مهارت، تهران، ایران

shpourbahrani@tvu.ac.ir

چکیده: در این پژوهش، یک الگوریتم خوشه‌بندی نوین مبتنی بر هسته آپولونیوس ارائه شده است که با بهره‌گیری از ماتریس لاپلاسین گراف و الگوریتم k میانگین، عملکردی بهتر در تفکیک خوشه‌ها نسبت به روش‌های سنتی دارد. در مرحله نخست، داده‌ها نرمال‌سازی می‌شوند تا مقیاس ویژگی‌ها یکسان شود. سپس، با محاسبه فاصله اقلیدسی بین نقاط و اعمال هسته آپولونیوس، ماتریس شباهت بین نمونه‌ها محاسبه می‌شود. با استفاده از گراف k نزدیک‌ترین همسایگان، فقط اتصالات همسایگان نزدیک حفظ می‌شوند. در ادامه، ماتریس لاپلاسین گراف از تفاضل ماتریس درجه و ماتریس شباهت به دست می‌آید که روابط بین نقاط داده را نمایان می‌کند. با انجام تجزیه مقادیر ویژه بر روی ماتریس لاپلاسین، بردارهای ویژه متناظر با کوچک‌ترین مقادیر ویژه انتخاب و داده‌ها به فضای کاهش‌یافته منتقل می‌شوند که در آن، خوشه‌ها واضح‌تر قابل تفکیک هستند. سپس، الگوریتم k میانگین بر روی داده‌های کاهش‌یافته اعمال می‌شود تا خوشه‌ها تعیین شوند. در نهایت، برای ارزیابی کیفیت خوشه‌بندی، از شاخص دقت و سایر شاخص‌های معتبر استفاده می‌شود. نتایج تجربی نشان می‌دهد این الگوریتم در مقایسه با روش‌های سنتی، به ویژه در داده‌های با ساختار غیرخطی، عملکردی بهتر دارد. همچنین، استفاده از هسته آپولونیوس در ساخت گراف شباهت باعث افزایش دقت در شناسایی ساختارهای محلی داده‌ها می‌شود.

واژه‌های کلیدی: خوشه‌بندی، هسته آپولونیوس، ماتریس لاپلاسین، k میانگین، ماتریس شباهت، مقادیر ویژه.

۱- مقدمه

خوشه‌بندی طیفی، با وجود توانایی در شناسایی ساختارهای پیچیده در داده‌ها، دارای معایبی است. یکی از مهم‌ترین چالش‌ها پیچیدگی محاسباتی زیاد است. همچنین، خوشه‌بندی طیفی ممکن است در شناسایی خوشه‌هایی با اندازه‌ها یا چگالی‌های متفاوت دچار مشکل شود و به حضور نویز و داده‌های پرت حساس باشد که ممکن است بر کیفیت نتایج تأثیر منفی بگذارد. در نهایت، تعیین تعداد مناسب خوشه‌ها برای خوشه‌بندی طیفی دشوار است و نیاز به دانش قبلی یا آزمایش‌های متعدد دارد.

الگوریتم k میانگین^۲ استاندارد توسط جیمز مک‌کوئین در مقاله‌ای با عنوان «برخی روش‌ها برای طبقه‌بندی و تحلیل مشاهدات چندمتغیره» معرفی شد که در سال ۱۹۶۷ در مجموعه مقالات پنجمین سمپوزیوم برکلی درباره آمار و احتمال ریاضی منتشر شد [۲]. در این روش، مک‌کوئین

در دهه‌های اخیر، الگوریتم‌های خوشه‌بندی طیفی^۱ به عنوان ابزارهای قدرتمند برای شناسایی ساختارهای پیچیده و غیرخطی در داده‌ها مورد توجه قرار گرفته‌اند. این الگوریتم‌ها با استفاده از ماتریس لاپلاسین گراف و تجزیه مقادیر ویژه، داده‌ها را به فضای ویژگی‌های جدیدی منتقل می‌کنند که در آن، خوشه‌ها واضح‌تر قابل تفکیک هستند. با این حال، محاسبه مقادیر ویژه برای ماتریس‌های بزرگ ممکن است از نظر محاسباتی پرهزینه باشد.

یک الگوریتم ساده برای خوشه‌بندی طیفی ارائه شده است که از بردارهای ویژه ماتریس لاپلاسین نرمال‌شده گراف شباهت داده‌ها استفاده می‌کند [۱]. این الگوریتم با بهره‌گیری از نظریه اختلال ماتریس، شرایطی را تحلیل می‌کند که تحت آن می‌تواند خوشه‌بندی مؤثری انجام دهد. الگوریتم

نشانی نویسنده مسئول: ایران، تهران، دانشگاه ملی مهارت، گروه مهندسی کامپیوتر

^۱ تاریخ ارسال مقاله: ۱۴۰۴/۰۴/۰۸

تاریخ پذیرش مقاله: ۱۴۰۵/۰۲/۰۵

نام نویسنده مسئول: شهین پوربهرامی

ارزیابی تمایل به خوشه‌بندی استفاده می‌شود، در حالی که مفهوم جفت نزدیک‌ترین همسایه در فرایند ادغام زیرمجموعه‌ها به کار می‌رود [۴].

در این راستا، الگوریتم پیشنهادی ما با بهره‌گیری از هسته آپولونیوس برای ساخت گراف شباهت، ماتریس لاپلاسی را تشکیل و سپس با استفاده از تجزیه طیفی و الگوریتم k میانگین، خوشه‌بندی داده‌ها را انجام می‌دهد. این رویکرد با تمرکز بر ساختارهای محلی داده‌ها و کاهش هزینه‌های محاسباتی، می‌تواند عملکردی بهتر در شناسایی خوشه‌های پیچیده و غیرخطی ارائه دهد.

ساختار این مقاله به شرح زیر سازمان‌دهی شده است. بخش ۲ مروری بر الگوریتم‌های خوشه‌بندی مبتنی بر چگالی و هندسه دارد و نقاط قوت و محدودیت‌های آنها را برجسته می‌کند. در بخش ۳، الگوریتم پیشنهادی را معرفی می‌کنیم که هسته آپولونیوس را با ماتریس لاپلاسی و k میانگین برای خوشه‌بندی داده‌ها ادغام می‌کند. بخش ۴ نتایج تجربی را ارائه و مورد بحث قرار می‌دهد و توانایی بهبود یافته روش پیشنهادی را در شناسایی ساختارهای پیچیده و غیرخطی در مقایسه با تکنیک‌های خوشه‌بندی سنتی و پیشرفته نشان می‌دهد. در نهایت، بخش ۵ نتیجه‌گیری مقاله را نشان می‌دهد.

۲- کارهای مرتبط

الگوریتم ساده چندفیلتری مبتنی بر k میانگین با استفاده از معیار هم‌ترازی هسته در محیط بدون نظارت، ترکیب فیلترها برای خوشه‌بندی را بهینه‌سازی می‌کند. با فرموله کردن مسئله به صورت بهینه‌سازی نرم، الگوریتم قادر به حل مؤثر آن با استفاده از روش نزول گرادین کاهش یافته است [۵].

روش یادگیری ساختار متقابل برای خوشه‌بندی‌های چندگانه^۴ با یادگیری هم‌زمان ساختارهای زوجی و خوشه‌ای، ماتریس تقسیم‌بندی برای خوشه‌بندی را بهبود می‌دهد. با استفاده از گراف شباهت تبعیض‌آمیز در فضای هسته و هم‌ترازی با تقسیم‌بندی‌های پایه، ساختارهای مختلف به صورت یکپارچه یاد گرفته می‌شوند [۶]. الگوریتم ابتدا یک گراف شباهت تبعیض‌آمیز در فضای هسته می‌سازد که با

الگوریتمی را برای تقسیم داده‌ها به k خوشه ارائه داد که هدف آن کمینه‌سازی واریانس درون خوشه‌ای است. الگوریتم با انتخاب تصادفی مراکز اولیه خوشه‌ها آغاز می‌شود و سپس، به صورت تکراری، هر داده به نزدیک‌ترین مرکز خوشه اختصاص می‌یابد و مراکز خوشه‌ها بر اساس میانگین داده‌های اختصاص یافته به آن‌ها به‌روز می‌شوند. این فرایند تا زمانی ادامه می‌یابد که مراکز خوشه‌ها تغییر نکنند یا تغییرات آنها به حداقل برسد. مک‌کوئین در این مقاله ویژگی‌های آماری الگوریتم و کاربردهای آن در تحلیل داده‌های چندمتغیره را بررسی کرد.

خوشه‌بندی k میانگین هسته پراکنده چالش‌های خوشه‌بندی داده‌های پیچیده و غیرخطی را بررسی می‌کند. روش‌های سنتی مانند k میانگین در شناسایی ساختارهای غیرخطی محدودیت دارند؛ بنابراین، تکنیک‌های جایگزین مانند k میانگین هسته و خوشه‌بندی طیفی توسعه یافته‌اند. با این حال، حضور متغیرهای نامربوط در داده‌ها می‌تواند عملکرد این الگوریتم‌ها را کاهش دهد [۳]. برای مقابله با این مشکل، روش‌های انتخاب متغیر مانند فیلتر، پوششی و جاسازی شده معرفی شده‌اند. الگوریتم k میانگین هسته یک روش انتخاب متغیر جاسازی‌شده با استفاده از فضای ضرب تانسوری و هسته تحلیل واریانس عمومی برای خوشه‌بندی غیرخطی است. این رویکرد با استفاده از فضای ضرب تانسوری، تعاملات پیچیده بین ویژگی‌ها را مدل‌سازی می‌کند و با اعمال هسته تحلیل واریانس، ساختارهای غیرخطی در داده‌ها را بهتر شناسایی می‌کند [۳].

الگوریتم خوشه‌بندی بر اساس آمار هاپکینز و k میانگین، محدودیت‌های الگوریتم سنتی k میانگین در خوشه‌بندی داده‌هایی با ساختارهای نامنظم را بررسی می‌کند و با استفاده از آمار هاپکینز^۳ و مفهوم جفت نزدیک‌ترین همسایه طراحی شده است. در مرحله اول، داده‌ها به زیرمجموعه‌های گروهی و نسبتاً یکنواخت تقسیم می‌شوند. سپس، با تعیین روابط مجاورت بین این زیرمجموعه‌ها، زیرمجموعه‌های مجاور و به هم پیوسته تا رسیدن به تعداد خوشه‌های مدنظر k ادغام می‌شوند. آمار هاپکینز در فرایند تقسیم‌بندی داده‌ها برای

از الگوریتم جست و جوی شاخک سوسک با ضریب میرایی استفاده می شود تا بر مشکل ناپایداری نتایج خوشه بندی غلبه کند.

خوشه بندی طیفی به نام SC-NR⁷ روشی برای بهبود ماتریس شباهت است که از ترکیب فاصله اقلیدسی وزن دار با تابع کرنل گاوسی استفاده می کند. وزن ها بر اساس ترتیب همسایگان نزدیک تنظیم می شوند تا روابط واقعی بین نقاط بهتر بازتاب یابند. با این حال، تابع کرنل گاوسی با استفاده از فاصله اقلیدسی و یک پارامتر ثابت سیگما، شباهت بین نقاط داده را اندازه گیری می کند. استفاده از یک سیگما ثابت می تواند به نتایج نادرست در خوشه بندی منجر شود، زیرا فاصله اقلیدسی به تنهایی قادر به بازتاب ساختار واقعی داده ها نیست. برای مثال، ممکن است نقاطی از خوشه های مختلف دارای فاصله اقلیدسی یکسانی باشند که در نتیجه، تابع کرنل گاوسی شباهتی یکسان برای آنها محاسبه می کند و این امر ممکن است به خوشه بندی نادرست منجر شود [۹].

در الگوریتم «K میانگین استاندارد هسته پراکنده برای انتخاب ویژگی در خوشه بندی غیرخطی»^۸ [۳]. نویسندگان چارچوبی برای خوشه بندی غیرخطی همراه با انتخاب متغیر پراکنده ارائه کرده اند که بر پایه گسترش الگوریتم K میانگین استاندارد کرنل با استفاده از کرنل ANOVA بنا شده است. این روش با تجزیه کرنل به مؤلفه های مربوط به اثرات اصلی و تعامل متغیرها و اعمال قیدهای ترکیبی بر ضرایب وزنی، امکان حذف ویژگی های بی اهمیت را فراهم می کند و هم زمان، ساختارهای غیرخطی پیچیده در داده را مدل می سازد. الگوریتم SKKM از بهینه سازی تناوبی برای به روزرسانی خوشه ها و وزن های ویژگی استفاده می کند. این رویکرد یکی از نخستین تلاش ها برای تلفیق یادگیری غیرخطی و انتخاب متغیر در چارچوبی یکپارچه است.

۳- روش پیشنهادی

مرحله اول: ساخت گراف شباهت با استفاده از هسته آپولونیوس

در این مرحله، ابتدا داده ها را نرمال سازی می کنیم تا مقیاس

استفاده از هسته های پایه مختلف، شباهت های بین نمونه ها را به صورت دقیق تر مدل می کند. سپس، با هم ترازی این گراف با تقسیم بندی های پایه، ساختارهای مختلف به صورت یکپارچه یاد گرفته می شوند. یکی از ویژگی های این الگوریتم استفاده از منظم سازی لاپلاسین^۹ است که به حفظ ساختار زوجی در ماتریس تقسیم بندی کمک می کند و در عین حال، اطلاعات ساختار خوشه ای را در گراف شباهت تزریق می کند. این روش به ویژه در کاربردهایی مانند تحلیل داده های چندمنظوره، پردازش تصاویر و شناسایی الگوهای پیچیده در داده های بزرگ، پتانسیل زیادی دارد.

در روش جدید مبتنی بر تابع آپولونیوس و خوشه بندی قله چگالی، کرنل جدید برای ماشین بردار پشتیبان^۶ معرفی شده است که بر پایه تابع آپولونیوس است. این رویکرد با هدف بهبود عملکرد طبقه بندی در داده های پیچیده و غیرخطی ارائه شده است [۷]. در این رویکرد، تابع آپولونیوس برای مدل سازی دقیق تر روابط هندسی بین نقاط داده استفاده می شود، در حالی که خوشه بندی قله های چگالی به شناسایی ساختارهای محلی و نقاط مرکزی در داده ها کمک می کند. این ترکیب باعث می شود کرنل آپولونیوس به صورت تطبیقی با ساختارهای مختلف داده ها سازگار شود و عملکردی بهتر در طبقه بندی داده های پیچیده نسبت به کرنل های سنتی مانند کرنل گاوسی ارائه دهد. در روش پیشنهادی، برای خوشه بندی از این کرنل برای افزایش دقت استفاده شده است.

الگوریتم خوشه بندی طیفی بر اساس عملکرد کرنل گاوسی بهبود یافته و جست و جوی آنتن سوسک با ضریب میرایی با بهبود تابع کرنل گاوسی و استفاده از الگوریتم جست و جوی شاخک سوسک با عامل میرایی، مشکلات مربوط به پارامترهای مقیاس و پایداری نتایج در خوشه بندی طیفی را حل می کند [۸]. در این الگوریتم، برای ساخت ماتریس شباهت، از یک تابع کرنل گاوسی بهبود یافته بهره برداری می شود که با استفاده از اطلاعات فاصله برخی از نزدیک ترین همسایگان، به طور تطبیقی پارامتر مقیاس را انتخاب می کند. این رویکرد به حل مشکل انتخاب پارامتر مقیاس در کرنل گاوسی کمک می کند. در مرحله خوشه بندی،

خوشه‌بندی با هسته آپولونیوس بر اساس ماتریس لاپلاسین و k میانگین

فراهم می‌کند که در فضای کاهش یافته قرار دارند و خوشه‌ها در آن واضح‌تر قابل تفکیک هستند.

مرحله چهارم: خوشه‌بندی با استفاده از k میانگین

در این مرحله، الگوریتم k میانگین را بر روی ماتریس ویژگی‌های به‌دست‌آمده از مرحله قبلی اعمال می‌کنیم. در این الگوریتم، ابتدا مراکز اولیه خوشه‌ها به صورت تصادفی انتخاب می‌شوند. سپس، هر داده به نزدیک‌ترین مرکز خوشه اختصاص می‌یابد و مراکز خوشه‌ها بر اساس میانگین داده‌های اختصاص یافته به آنها به‌روز می‌شوند. این فرایند تا زمانی ادامه می‌یابد که مراکز خوشه‌ها تغییر نکنند یا تغییرات آنها به حداقل برسد [۲].

۴- نتایج

ارزیابی عملکرد روش پیشنهادی

برای ارزیابی کیفیت خوشه‌بندی از شاخص‌هایی مختلف همچون ARI، Accuracy (دقت) و NMI استفاده می‌کنیم. شاخص ARI میزان شباهت بین خوشه‌بندی به‌دست‌آمده و برچسب‌های واقعی داده‌ها را اندازه‌گیری می‌کند. مقدار ARI بین ۱- و ۱+ متغیر است که مقدار ۱ نشان‌دهنده تطابق کامل بین خوشه‌بندی و برچسب‌های واقعی است. معیار اطلاعات متقابل نرمال‌شده NMI بر پایه نظریه اطلاعات است و میزان اشتراک اطلاعات بین برچسب‌های واقعی و برچسب‌های پیش‌بینی‌شده توسط الگوریتم خوشه‌بندی را اندازه‌گیری می‌کند. NMI بین ۰ و ۱ مقدار می‌گیرد؛ مقدار ۱ نشان‌دهنده تطابق کامل بین خوشه‌بندی و برچسب‌های واقعی است، در حالی که مقدار ۰ نشان‌دهنده عدم تطابق کامل است. یکی از ویژگی‌های مهم NMI این است که مستقل از تعداد خوشه‌ها و برچسب‌ها عمل می‌کند و به همین دلیل، برای مقایسه الگوریتم‌های مختلف خوشه‌بندی مناسب است.

معیار دقت^۱ خوشه‌بندی بر اساس فرمول (۳)، نسبت تعداد نمونه‌هایی را که به‌درستی خوشه‌بندی شده‌اند به کل نمونه‌ها محاسبه می‌کند (برای محاسبه معنادار دقت در مسائل خوشه‌بندی لازم است ابتدا یک نگاشت یک‌به‌یک بین

ویژگی‌ها یکسان شود. سپس، با ضرب ویژگی‌ها در وزن‌های مشخص‌شده، اهمیت هر ویژگی را تنظیم می‌کنیم. در ادامه، یک گراف k نزدیک‌ترین همسایگان با استفاده از فاصله اقلیدسی بین نقاط ساخته می‌شود. برای محاسبه شباهت بین نقاط [۱۰، ۱۱]. از هسته آپولونیوس که بر اساس ساخت دایره آپولونیوس است [۱۲، ۱۳] استفاده می‌کنیم که به صورت فرمول (۱) تعریف می‌شود:

$$k(x, y) = \frac{1}{d(x, y) + \varepsilon} \quad (1)$$

که در آن، اپسیلون یک مقدار ثابت کوچک است که از تقسیم بر صفر جلوگیری می‌کند (پارامتر اپسیلون نقش فیلتر نرم‌کننده^۹ را دارد تا از تقسیم بر صفر جلوگیری و حساسیت شباهت به فواصل خیلی کوچک را کنترل کند). شباهت محاسبه‌شده با ماتریس k نزدیک‌ترین همسایگان ضرب می‌شود تا فقط اتصالات همسایگان نزدیک حفظ شوند.

مرحله دوم: ساخت ماتریس لاپلاسین گراف

در این مرحله، ماتریس درجه را محاسبه می‌کنیم که هر قطر آن نشان‌دهنده مجموع وزن‌های یال‌های متصل به گره مربوط است. سپس، ماتریس لاپلاسین گراف را با استفاده از فرمول (۲) محاسبه می‌کنیم:

$$L = W - D \quad (2)$$

که در آن، D ماتریس درجه و W ماتریس شباهت است. این ماتریس نمایانگر ساختار گراف و روابط بین نقاط داده است.

مرحله سوم: تجزیه طیفی و کاهش ابعاد

در این مرحله، تجزیه مقادیر ویژه ماتریس لاپلاسین انجام می‌شود تا مقادیر ویژه و بردارهای ویژه آن به دست آیند. سپس، بردارهای ویژه متناظر با کوچک‌ترین مقادیر ویژه را انتخاب می‌کنیم و آنها را به صورت ماتریسی با ابعاد $n \times k$ (که در آن k تعداد خوشه‌ها و n تعداد نمونه‌های ذخیره‌شده است) نشان می‌دهیم. این ماتریس ویژگی‌هایی جدید را برای داده‌ها

می‌کنند.

ارزیابی عملکرد الگوریتم پیشنهادی و سایر الگوریتم‌ها با مجموعه داده‌های واقعی

مجموعه داده‌های واقعی در جدول (۱)^۲، مانند Iris، Olivetti، Cancer، Breast، Digits، Wine و Faces، به طور گسترده برای ارزیابی و آموزش الگوریتم‌های مختلف استفاده می‌شوند. مجموعه داده Iris شامل ۱۵۰ نمونه از سه گونه گل زنبق است که هر نمونه دارای ۴ ویژگی است. مجموعه داده Wine شامل ۱۷۸ نمونه از سه نوع نوشیدنی با ۱۳ ویژگی شیمیایی است. مجموعه داده Digits شامل ۱۷۹۷ تصویر ۸×۸ از ارقام دست‌نویس ۰ تا ۹ است که هر تصویر به یک بردار ۶۴ بُعدی تبدیل شده است. مجموعه داده Cancer شامل ۵۶۹ نمونه با ۳۰ ویژگی عددی برای تشخیص سرطان سینه است. مجموعه داده Diabetes شامل ۷۶۸ نمونه با ۸ ویژگی برای پیش‌بینی ابتلا به دیابت است. مجموعه داده Olivetti Faces شامل ۴۰۰ تصویر چهره از ۴۰ فرد مختلف است که هر تصویر به صورت یک بردار ۴۰۹۶ بُعدی (۶۴×۶۴ پیکسل) نمایش داده می‌شود. این مجموعه داده‌ها به عنوان معیارهای استاندارد در پژوهش‌های علمی شناخته می‌شوند و برای مسائل مختلف مانند طبقه‌بندی، خوشه‌بندی و کاهش ابعاد داده‌ها کاربرد دارند.

برچسب‌های خوشه‌ای و برچسب‌های واقعی کلاس پیدا شود). از آنجا که در خوشه‌بندی، برچسب‌های پیش‌بینی شده ممکن است با برچسب‌های واقعی از نظر نام‌گذاری متفاوت باشند، قبل از محاسبه دقت، معمولاً یک نگاشت بین برچسب‌های پیش‌بینی شده و واقعی انجام می‌شود تا بیشترین تطابق حاصل شود. دقت خوشه‌بندی نیز بین ۰ و ۱ مقدار می‌گیرد؛ مقدار ۱ نشان‌دهنده تطابق کامل است.

بر اساس فرمول (۳)، y_i نشان‌دهنده برچسب واقعی کلاس است و C_i شماره برچسب خوشه‌ای است که از فرایند خوشه‌بندی به دست آمده است. تابع $\delta(x, y) = \begin{cases} 1 & x = y \\ 0 & x \neq y \end{cases}$ یک تابع شاخص^{۱۱} است که در صورت برابری دو برچسب، مقدار ۱ و در غیر این صورت، مقدار ۰ می‌گیرد. بهترین نگاشت بین برچسب‌های خوشه و برچسب‌های واقعی کلاس با استفاده از الگوریتم Hungarian یا Munkres تعیین می‌شود تا بیشترین تطابق ممکن بین خوشه‌ها و کلاس‌های واقعی حاصل شود.

$$ACC = \frac{1}{n} \sum_{i=1}^n \delta(y_i, \text{map}(c_i)) \quad (3)$$

استفاده هم‌زمان از NMI و دقت می‌تواند دیدی جامع‌تر از عملکرد الگوریتم خوشه‌بندی ارائه دهد، زیرا هر یک از این معیارها جنبه‌هایی متفاوت از کیفیت خوشه‌بندی را بررسی

جدول (۱): مجموعه داده‌های واقعی

تعداد داده‌ها	تعداد ابعاد	تعداد خوشه‌ها	مجموعه داده
۱۵۰	۴	۳	Iris
۱۷۸	۱۳	۳	Wine
۱۷۹۷	۶۴	۱۰	Digits
۵۶۹	۳۰	۲	Breast Cancer (Wisconsin Breast Cancer Diagnostic (WBCD))
۷۶۸	۸	۲	Diabetes
۴۰۰	۴۰۹۶	۴۰	Olivetti Faces

جدول (۲): ارزیابی الگوریتم پیشنهادی در مجموعه داده‌های واقعی

نکته کلیدی عملکرد	نوع لاپلاسین	پارامتر k	دقت الگوریتم پیشنهادی	NMI الگوریتم پیشنهادی	مجموعه داده
شکاف طیفی بین λ_3 و λ_4 نشانگر سه خوشه است.	L غیر نرمال شده	۱۰	۰/۶۸۷	۰/۴۶۰	Iris
حساس به انتخاب k ، عملکرد پایدار با k برابر ۱۲ است.	L_{sym} نرمال شده مقارن	۱۲	۰/۹۵۵	۰/۸۵۳	Wine
تعبیه با $d=k$ عملکردی بهتر از شکاف طیفی دارد.	L_{rw} نرمال شده تصادفی	۱۵	۰/۸۲۹	۰/۸۹۲	Digits
خوشه‌بندی پایدار، اما حساس به نویز.	L غیر نرمال شده	۱۰	۰/۹۳۵	۰/۶۴۵	Breast Cancer
نتیجه بعد از اعمال کاهش ابعاد با PCA	L غیر نرمال شده	۱۰	۰/۹۳۸	۰/۶۵۸	Breast Cancer (PCA)
داده با هم‌پوشانی زیاد، دقت کمتر ولی NMI زیاد.	L_{sym}	۸	۰/۳۷۹	۰/۷۹۳	Diabetes
افزایش k باعث بهبود پایداری تعبیه شد.	L_{rw}	۱۵	۰/۵۵۲	۰/۷۹۵	Olivetti Faces
نتیجه بعد از اعمال کاهش ابعاد با PCA	L_{rw}	۱۵	۰/۶۱۸	۰/۸۰۹	Olivetti Faces (PCA)
—	—	—	۰/۲۱۴±۰/۷۳۳	۰/۲۲۲±۰/۷۳۹	SD ± میانگین عملکرد الگوریتم پیشنهادی

جدول (۳): ارزیابی الگوریتم پیشنهادی، الگوریتم SKKM و k میانگین استاندارد در مجموعه داده‌های واقعی

مجموعه داده‌ها	ARI	ARI	دقت	NMI
	الگوریتم پیشنهادی	الگوریتم SKKM	استاندارد میانگین k	استاندارد میانگین k
Iris	۰/۷۰	۰/۷۴	۰/۷۸	۰/۷۴
Wine	۰/۸۷	۰/۸۶	۰/۶۴	۰/۳۱
Digits	۰/۸۳	۰/۸۱	۰/۳۷	۰/۵۰
Breast Cancer	۰/۷۵	۰/۷۲	۰/۸۳	۰/۷۳
Diabetes	-	-	۰/۶۵	۰/۰۰۵
Olivetti Faces	-	-	۰/۵۴	۰/۵۵
SD ± میانگین	۰/۷۸۸±۰/۰۷۱	۰/۷۸۳±۰/۰۶۱	۰/۶۳۵±۰/۱۵۸	۰/۴۷۲±۰/۲۵۵

ساختارهای درونی داده‌ها را به شکلی دقیق‌تر شناسایی کند. در مجموعه داده Iris، مقدار NMI برابر ۰/۴۶۰ و دقت برابر ۰/۶۸۷ به دست آمده است که با استفاده از لاپلاسین غیر نرمال شده (L) حاصل شده است. بررسی طیف ویژه لاپلاسین نشان داد شکاف قابل ملاحظه‌ای میان λ_3 و λ_4 وجود دارد که وجود سه خوشه متمایز را در این داده تأیید می‌کند.

جدول (۲) نتایج عملکرد الگوریتم پیشنهادی مبتنی بر گراف آپولونیوسی را در چند مجموعه داده واقعی به همراه پارامترهای بهینه و نوع لاپلاسین به کار رفته نشان می‌دهد. همان‌گونه که در ستون‌های NMI و دقت مشاهده می‌شود، الگوریتم پیشنهادی در بیشتر مجموعه داده‌ها نتایجی پایدار داشته و در مقایسه با روش‌های کلاسیک، توانسته است

طور کلی، این نتایج اثبات می‌کنند ترکیب کرنل آپولونیوس با روش تعبیه طیفی می‌تولند یک چارچوب انعطاف‌پذیر و دقیق برای خوشه‌بندی داده‌های پیچیده فراهم کند.

جدول (۳) عملکرد الگوریتم پیشنهادی را در مقایسه با روش‌های پایه شامل k میانگین استاندارد بر اساس معیارهای NMI و دقت و الگوریتم SKKM بر اساس ARI روی مجموعه داده‌های مختلف نشان می‌دهد. نتایج NMI و دقت k میانگین استاندارد با میانگین‌های نسبتاً کم (به ترتیب $0/472$ و $0/635$) و انحراف معیارهای زیاد بیانگر این است که k میانگین استاندارد در برابر ساختارهای پیچیده و داده‌های توزیع غیرخطی کارایی محدودی دارد. در مقابل، مقدار ARI الگوریتم پیشنهادی در چهار مجموعه داده مشترک از جمله Wine، Iris، Digits و Breast Cancer میانگین $0/788 \pm 0/071$ را نشان می‌دهد که بیانگر عملکرد باثبات و دقیق در بازیابی ساختار واقعی خوشه‌هاست. این مقدار اندکی بیشتر از الگوریتم SKKM با میانگین $0/783 \pm 0/061$ است که نشان می‌دهد روش پیشنهادی نه فقط در بیشتر مجموعه داده‌ها عملکرد مشابه یا بهتر نسبت به SKKM دارد، بلکه توانایی استخراج ویژگی‌های ساختاری عمیق‌تری از داده‌ها را نیز ارائه می‌دهد. همچنین، نبود مقدار ARI برای مجموعه داده‌های Diabetes و Olivetti Faces به دلیل ابعاد یا ویژگی‌های خاص آنهاست که در این آزمایش‌ها لحاظ نشده‌اند. در مجموع، جدول (۳) نشان می‌دهد الگوریتم پیشنهادی از نظر پایداری، دقت خوشه‌بندی و توانایی تفکیک ساختارهای درهم‌تنیده، نسبت به روش‌های پایه و حتی SKKM برتری نسبی دارد.

این یافته نشان می‌دهد الگوریتم پیشنهادی قادر است ساختار خوشه‌ای داده‌های کم‌بعد و نسبتاً منظم را به خوبی بازیابی کند. در مجموعه داده Wine، نمایش مقادیر حاصل $NMI=0/853$ و دقت $0/955$ نشان‌دهنده عملکرد بسیار مطلوب الگوریتم در داده‌های با ویژگی‌های پیوسته و خوشه‌های هم‌پوشان است. استفاده از لاپلاسی‌ن نرمال‌شده متقارن موجب پایداری نتایج شده است، زیرا این نوع لاپلاسی‌ن توانسته اثر تفاوت در اندازه خوشه‌ها را تعدیل کند. در مقابل، در مجموعه داده Digits که دارای ساختار غیرخطی و چندخوشه‌ای است، الگوریتم با انتخاب لاپلاسی‌ن نرمال‌شده تصادفی مقدار $NMI=0/892$ و دقت $0/892$ را به دست آورده است.

در مجموعه داده Breast Cancer، الگوریتم پیشنهادی با $NMI=0/645$ و دقت $0/935$ توانسته است عملکردی پایدار ارائه دهد، هرچند به نوبه حساس است. این داده نشان می‌دهد در حضور ویژگی‌های همبسته و نویزی، استفاده از لاپلاسی‌ن غیرنرمال‌شده می‌تواند خوشه‌بندی معناداری ایجاد کند. در مجموعه داده Diabetes، مقدار NMI زیاد ($0/793$) در کنار دقت کمتر ($0/379$) مشاهده می‌شود که نشان می‌دهد الگوریتم توانسته است ساختارهای نسبی خوشه‌ها را تشخیص دهد، اما به دلیل هم‌پوشانی زیاد بین داده‌ها، برجسب‌های نهایی دقت کمتری دارند. در اینجا، استفاده از لاپلاسی‌ن متقارن باعث بهبود پایداری نسبت به حالت غیرنرمال‌شده شده است. در نهایت، در مجموعه داده Olivetti Faces که از تصاویر چهره تشکیل شده، مقادیر $NMI=0/793$ و دقت $0/552$ نشان‌دهنده توانایی الگوریتم در بازنمایی روابط چهره‌هاست. استفاده از لاپلاسی‌ن تصادفی و افزایش k تا ۱۵ موجب بهبود پایداری تعبیه شده است.

به طور کلی، میانگین نتایج در تمام مجموعه داده‌ها با NMI برابر $0/739 \pm 0/22$ و دقت برابر $0/733 \pm 0/214$ نشان می‌دهد الگوریتم پیشنهادی از نظر میانگین عملکرد و پایداری دارد. تحلیل پارامترها مشخص می‌کند بازه بهینه برای k حدود ۱۰ تا ۱۵ است. همچنین، نوع لاپلاسی‌ن باید با ساختار داده تطبیق یابد که برای داده‌های هم‌تراز و L ساده مناسب است. به

خوشه‌بندی با هسته آپولونیوس بر اساس ماتریس لاپلاسین و k میانگین

جدول (۴): ارزیابی خوشه‌بندی طیفی در مجموعه داده‌های واقعی

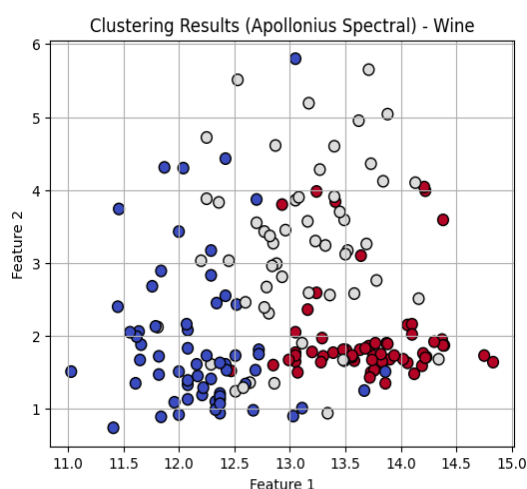
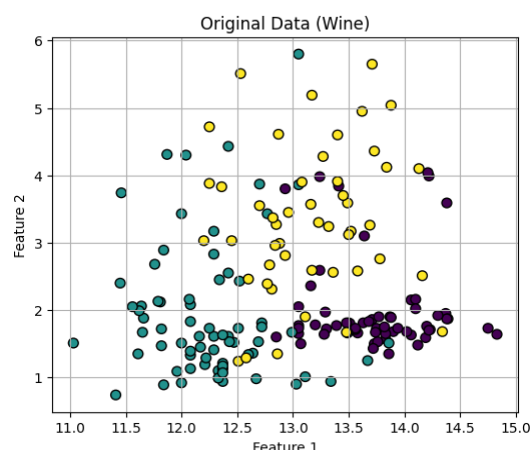
مجموعه داده‌ها	NMI (الگوریتم خوشه‌بندی طیفی)	دقت (الگوریتم خوشه‌بندی طیفی)
Iris	۰/۷۶	۰/۹۰
Wine	۰/۴۶	۰/۷۳
Digits	۰/۵۵	۰/۷۰
Breast Cancer	۰/۷۵	۰/۹۶
Diabetes	۰/۰۲	۰/۶۸
Olivetti Faces	۰/۵۸	۰/۶۰
$\pm$ میانگین SD	$۰/۲۶۶ \pm ۰/۵۱۰$	$۰/۱۳۵ \pm ۰/۷۶۲$

مقایسه عملکرد الگوریتم پیشنهادی با الگوریتم خوشه‌بندی طیفی بر روی مجموعه داده‌های مختلف در جدول (۲) و (۴) نشان می‌دهد الگوریتم پیشنهادی در برخی از مجموعه داده‌ها عملکردی بهتر دارد. برای مثال، در مجموعه داده‌های Wine و Digits، الگوریتم پیشنهادی با NMI برابر ۰/۸۵۳ و ۰/۸۹۲ و دقت‌های برابر ۰/۹۵۵ و ۰/۸۹۲، نسبت به خوشه‌بندی طیفی که NMI ۰/۴۰ و ۰/۵۵ و دقت‌های ۰/۷۳ و ۰/۷۰ دارد، عملکردی بهتر نشان داده است. این بهبود می‌تواند به دلیل استفاده از تکنیک‌های پیشرفته‌تر در الگوریتم پیشنهادی باشد که ساختارهایی پیچیده‌تر را در داده‌ها شناسایی می‌کند.

از سوی دیگر، در مجموعه داده‌هایی مانند Iris، Breast Cancer، الگوریتم خوشه‌بندی طیفی عملکردی بهتر نسبت به الگوریتم پیشنهادی دارد. برای مثال، در مجموعه داده Iris، خوشه‌بندی طیفی با $NMI=۰/۷۶$ و دقت ۰/۹۰، نسبت به الگوریتم پیشنهادی با $NMI=۰/۴۶۰$ و دقت ۰/۶۸۷، نتایجی بهتر ارائه داده است. این به دلیل سادگی ساختار مجموعه داده‌ها است که الگوریتم‌های مبتنی بر طیف قادر به شناسایی خوشه‌ها با دقت بیشتری باشند.

ارزیابی عملکرد الگوریتم پیشنهادی و سایر الگوریتم‌ها با مجموعه داده‌های مصنوعی

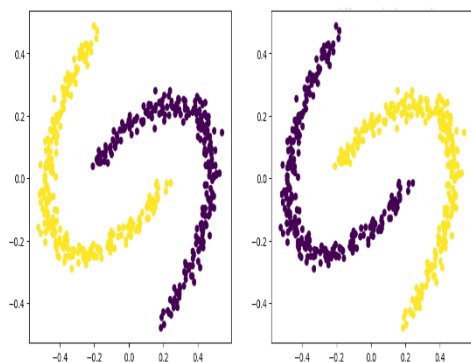
دیتاست Blobs مجموعه‌ای از داده‌هاست که به صورت خوشه‌های گوسی در فضای ویژگی‌ها توزیع شده‌اند. ابعاد قابل تنظیم به طور پیش فرض ۲ بعدی هستند. تعداد نمونه‌ها به طور پیش فرض ۱۰۰ نمونه است. تعداد خوشه‌ها به طور



شکل (۱): داده واقعی Wine در دو خروجی اصلی و خروجی با الگوریتم پیشنهادی

در شکل (۱)، نتایج اجرای الگوریتم بر روی داده‌ی واقعی Wine نشان داده شده است. در این شکل، توزیع اصلی داده در فضای دو ویژگی نخست نمایش یافته است و نمونه‌ها بر اساس برچسب‌های واقعی رنگ‌آمیزی شده‌اند. همچنین، خروجی خوشه‌بندی مبتنی بر گراف آپولونیوس و تعبیه طیفی ارائه شده است. همان‌گونه که مشاهده می‌شود، ساختار غیرخطی داده در فضای طیفی به صورت واضح‌تر نمایان شده است و خوشه‌ها تفکیک‌پذیری بیشتری یافته‌اند که بیانگر توانایی مدل در استخراج ساختار درونی داده و بهبود جداسازی نسبت به فضای ویژگی اولیه است.

پیش فرض ۳ خوشه است.



شکل (۳): داده مصنوعی Moons

جدول (۵): ARI در مجموعه داده های مصنوعی

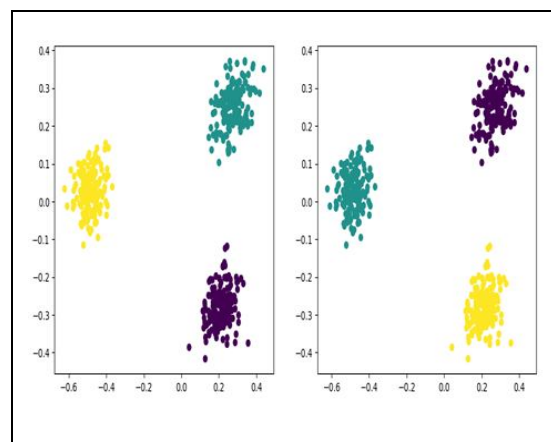
روش پیشنهادی	الگوریتم خوشه بندی طیفی	k میانگین استاندارد	مجموعه داده ها
۰/۸۸	۰/۶۴	۰/۵۷	R15
۰/۶۷۶	۱	۱	Blobs
۱	۱	۱	Moons
۰/۱۶۷±۰/۸۵۴	۰/۸۸۰ ۰/۱۶۶±	۰/۲۰۳±۰/۸۵۷	± SD میانگین

جدول (۶): NMI در مجموعه داده های مصنوعی

روش پیشنهادی	الگوریتم خوشه بندی طیفی	k میانگین استاندارد	مجموعه داده ها
۰/۹۴	۰/۷۱	۰/۵۷	R15
۰/۷	۱	۱	Blobs
۱	۰/۹۶	۰/۳۸	Moons
۰/۱۲۲±۰/۹۰	۰/۱۵۵±۰/۸۹۰	۰/۳۱۸±۰/۶۵۰	± SD میانگین

جدول (۷): دقت در مجموعه داده های مصنوعی

روش پیشنهادی	الگوریتم خوشه بندی طیفی	k میانگین استاندارد	مجموعه داده ها
۰/۹۲۸	۰/۹۱	۰/۸۷	R15
۰/۸۶۰	۱	۱	Blobs
۱	۰/۹۷	۰/۳۸	Moons
۰/۰۷۱±۰/۹۲۹	۰/۰۴۵±۰/۹۶۰	± ۰/۷۵۰ ۰/۲۹۶	± SD میانگین



شکل (۲): داده مصنوعی Blobs

دیتاست Moons شامل دو نیم دایره درهم تنیده است که به شکل دو هلال ماه در فضای دوبعدی قرار دارند. این ساختار غیرخطی برای آزمایش الگوریتم هایی مناسب است که توانایی شناسایی مرزهای پیچیده را دارند. این داده دوبعدی است. تعداد نمونه ها به طور پیش فرض ۱۰۰ نمونه است. تعداد خوشه ها ۲ خوشه است. نقاط داده در دو نیم دایره درهم تنیده با امکان افزودن نویز گوسی توزیع می شوند.

شکل های (۲) و (۳) اثربخشی الگوریتم خوشه بندی پیشنهادی را بیشتر نشان می دهند و عملکرد خوشه بندی بی نقصی را در دو مجموعه داده مصنوعی متنوع Blobs و Moons نشان می دهند. همه معیارهای ارزیابی به حداکثر مقادیر خود رسیده اند که این امر نشان می دهد الگوریتم با موفقیت ساختارهای خوشه بندی اساسی را حتی در توزیع های داده پیچیده و غیرمحدب بدون خطا شناسایی می کند. این امر سازگاری الگوریتم پیشنهادی را در مدیریت انواع اشکال داده برجسته می کند.

خوشه‌بندی با هسته آپولونیوس بر اساس ماتریس لاپلاسن و k میانگین

می‌توانند عملکردی مناسب داشته باشند، در حالی که در داده‌هایی با ساختار پیچیده‌تر، الگوریتم‌های پیشرفته‌تر مانند خوشه‌بندی طیفی و الگوریتم پیشنهادی می‌توانند عملکردی بهتر ارائه دهند.

تحلیل آماری کیفی (Wilcoxon & Friedman)

نتایج در جدول (۸) نشان می‌دهد الگوریتم پیشنهادی در مقایسه با دو روش دیگر، یعنی خوشه‌بندی طیفی و K میانگین استاندارد، عملکردی بسیار بهتر از خود نشان داده است. در بخش میانگین NMI ، مقدار به‌دست‌آمده برای روش پیشنهادی بیشتر از دو روش دیگر است که نشان می‌دهد الگوریتم پیشنهادی توانسته است ساختارهای درونی داده‌ها را بهتر شناسایی کند و خوشه‌هایی با شباهت درونی بیشتر ایجاد کند. همچنین، مقدار انحراف معیار کمتر بیانگر پایداری بهتر این روش در میان تکرارهای مختلف و بر روی مجموعه‌داده‌های گوناگون است. این نتایج نشان می‌دهد روش پیشنهادی علاوه بر دقت بیشتر، از نظر پایداری و سازگاری نیز برتر است.

در ادامه، آزمون Signed-Rank Wilcoxon برای مقایسه زوجی الگوریتم‌ها انجام شده است. بر اساس مقادیر p اختلاف میان روش پیشنهادی با هر دو الگوریتم دیگر در سطح اطمینان ۹۵ درصد معنادار است؛ به طوری که مقدار p بین روش پیشنهادی و خوشه‌بندی طیفی برابر $۰/۰۳۱$ و بین روش پیشنهادی و K میانگین استاندارد برابر $۰/۰۱۶$ بوده است. این بدان معناست که روش پیشنهادی به طور آماری در شاخص NMI عملکردی بهتر دارد. تنها تفاوت بین خوشه‌بندی طیفی و K میانگین از نظر آماری معنادار نیست که نشان می‌دهد عملکرد این دو روش در سطحی مشابه قرار دارد.

نتایج آزمون Friedman نیز این یافته‌ها را تأیید می‌کند. مقدار آماره χ^2 برابر $۱۰/۳۳$ است و همچنین p برابر $۰/۰۰۵۷$ است که نشان‌دهنده وجود تفاوت معنادار در عملکرد سه الگوریتم است. میانگین رتبه‌های فریدمن نشان می‌دهد الگوریتم پیشنهادی با رتبه $۱/۱۷$ بهترین عملکرد را دارد، در

برای مقایسه عملکرد الگوریتم‌های خوشه‌بندی بر روی مجموعه‌داده‌های مصنوعی R15، Blobs و Moons جدول‌های (۵)، (۶) و (۷)، سه معیار ارزیابی NMI ، ARI و دقت بررسی شده‌اند. نتایج نشان می‌دهد الگوریتم پیشنهادی، الگوریتم خوشه‌بندی طیفی و الگوریتم k میانگین استاندارد در هر یک از این مجموعه‌داده‌ها عملکردی متفاوت دارند.

در مجموعه‌داده R15، الگوریتم پیشنهادی با ARI برابر $۰/۸۸$ ، NMI برابر $۰/۹۴$ و دقت برابر $۰/۹۲$ ، عملکردی بهتر نسبت به الگوریتم‌های دیگر داشته است. الگوریتم خوشه‌بندی طیفی با ARI برابر $۰/۶۴$ ، NMI برابر $۰/۷۱$ و دقت برابر $۰/۹۱$ در رتبه دوم قرار دارد، در حالی که الگوریتم k میانگین استاندارد با ARI برابر $۰/۵۷$ ، NMI برابر $۰/۵۷$ و دقت برابر $۰/۸۷$ کمترین عملکرد را نشان داده است. این نتایج نشان می‌دهد الگوریتم پیشنهادی در شناسایی خوشه‌های پیچیده‌تر موفق‌تر عمل کرده است. مجموعه‌داده Blobs در این مجموعه‌داده، الگوریتم خوشه‌بندی طیفی و k میانگین استاندارد هر دو با ARI و NMI برابر ۱ و دقت برابر ۱، عملکردی بسیار خوب داشته‌اند. الگوریتم پیشنهادی با ARI برابر $۰/۶۷۶$ ، NMI برابر $۰/۷۵۹$ و دقت برابر $۰/۸۶۰$ ، عملکردی کمتر نسبت به دو الگوریتم دیگر داشته است. این نتایج نشان می‌دهد در داده‌هایی با ساختار ساده‌تر، الگوریتم‌های سنتی‌تر مانند k میانگین می‌توانند عملکردی مناسب داشته باشند. مجموعه‌داده Moons در این مجموعه‌داده، الگوریتم پیشنهادی و الگوریتم خوشه‌بندی طیفی هر دو با ARI و NMI برابر ۱ و دقت برابر ۱، عملکردی بسیار خوب داشته‌اند. الگوریتم k میانگین استاندارد با ARI برابر ۱، NMI برابر $۰/۳۸$ و دقت برابر $۰/۷۸$ ، عملکردی با دقت کمتر نسبت به دو الگوریتم دیگر داشته است. این نتایج نشان می‌دهد در داده‌هایی با ساختار غیرخطی و پیچیده‌تر، الگوریتم‌های پیشرفته‌تر مانند خوشه‌بندی طیفی می‌توانند عملکردی بهتر داشته باشند.

در مجموع، این مقایسه‌ها نشان می‌دهند انتخاب الگوریتم خوشه‌بندی مناسب به ساختار داده‌ها بستگی دارد. در داده‌هایی با ساختار ساده‌تر، الگوریتم‌های سنتی‌تر مانند k میانگین

نسبت به K میانگین استاندارد برتری آشکار دارد و در مقایسه با خوشه‌بندی طیفی نیز عملکردی بهتر، ولی نه به صورت معنادار از نظر آماری نشان می‌دهد. در مجموع، می‌توان نتیجه گرفت الگوریتم پیشنهادی پایدارترین و دقیق‌ترین روش خوشه‌بندی در این آزمایش‌ها بوده است.

حالی که خوشه‌بندی طیفی و K میانگین استاندارد به ترتیب با رتبه‌های ۲/۰۰ و ۲/۸۳ در جایگاه دوم و سوم قرار گرفته‌اند. در آزمون Nemenyi Post-hoc نیز تفاوت بین روش پیشنهادی و K میانگین استاندارد از نظر آماری معنادار است، اما اختلاف بین روش پیشنهادی و خوشه‌بندی طیفی معنادار نیست. بنابراین، در سطح اطمینان ۹۵ درصد، روش پیشنهادی

جدول (۸): تحلیل آماری الگوریتم‌ها بر اساس شاخص NMI

نتیجه آماری/تفسیر	K میانگین استاندارد	خوشه‌بندی طیفی	الگوریتم پیشنهادی	معیار	شاخص/آزمون
روش پیشنهادی بیشترین مقدار میانگین را دارد	۰/۲۶±۰/۴۷	۰/۲۷±۰/۵۱	۰/۲۹±۰/۶۲	میانگین ± انحراف معیار	میانگین NMI واقعی
p برابر ۰/۰۳۱ → اختلاف معنادار به نفع الگوریتم پیشنهادی	—	—	—	p(Proposed vs Spectral)	Wilcoxon Signed-Rank Test
p برابر ۰/۰۱۶ → اختلاف معنادار به نفع الگوریتم پیشنهادی	—	—	—	p(Proposed vs K-Means)	
p برابر ۰/۳۱۳ → اختلاف معنادار نیست	—	—	—	p(Spectral vs K-Means)	
تفاوت معنادار بین الگوریتم‌ها → p کمتر از ۰/۰۱	—	—	—	χ^2 برابر ۰/۳۳ و p برابر ۰/۰۰۵۷	Friedman Test
الگوریتم پیشنهادی بهترین میانگین رتبه را دارد	۲/۸۳	۲/۰۰	۱/۱۷	رتبه (۱=بهتر)	میانگین رتبه فریدمن
CD کمتر از ۰/۸۵ → تفاوت معنادار نیست	—	—	۰/۸۳	Rank (Proposed–Spectral)	Nemenyi Post-hoc
CD بیشتر از ۰/۸۵ → تفاوت معنادار	—	—	۱/۶۶	Rank (Proposed–K-Means)	
CD برابر ۰/۸۵ در سطح اطمینان ۹۵ درصد	—	—	—	مقدار	Critical Difference (CD)
الگوریتم پیشنهادی از K-Means به طور معنادار بهتر و از Spectral کمی بهتر است	ضعیف‌تر در تمام داده‌ها	مشابه الگوریتم پیشنهادی در برخی از داده‌ها	الگوریتم پیشنهادی بهترین عملکرد کلی	تفسیر کلی	نتیجه نهایی

جدول (۹): مطالعه ابلیشن الگوریتم پیشنهادی در داده‌های واقعی

تغییر اعمال شده	توضیح مؤلفه جایگزین	NMI $\pm$ SD میانگین	دقت $\pm$ SD میانگین	تغییر عملکرد نسبت به مدل اصلی	تحلیل کیفی
مدل اصلی آپولونیوس $L + kNN$ + متقابل + تعبیه طیفی	پیکربندی پایه الگوریتم پیشنهادی	0.739 ± 0.22	0.733 ± 0.214	—	بهترین عملکرد، پایداری زیاد
جایگزینی کرنل آپولونیوس با کرنل گوسی	$S_{ij} = \exp\left(-\frac{\ x_i - x_j\ ^2}{2\sigma^2}\right)$	0.691 ± 0.19	0.684 ± 0.21	-۶۵%	افت جزئی در داده‌های پیچیده، کرنل آپولونیوس به فواصل بلند حساس‌تر است
تغییر لاپلاسی از L به L_{sym}	نرمال‌سازی متقارن	0.726 ± 0.20	0.724 ± 0.20	-۲/۱%	پایدارتر در داده‌های نامتوازن، اما بدون بهبود معنادار
تغییر گراف kNN متقابل $\leftrightarrow$ ساده	حذف شرط تقارن در همسایگی	0.705 ± 0.18	0.693 ± 0.19	-۷/۴%	کاهش پایداری به ویژه در داده‌های نویزی
اجرای K میانگین استاندارد در فضای اصلی (بدون تعبیه طیفی)	حذف گام طیف نگاری	0.552 ± 0.23	0.541 ± 0.21	-۱/۲۴%	افت شدید عملکرد، تأیید اهمیت تعبیه طیفی
اجرای K میانگین استاندارد پس از تعبیه طیفی (مدل کامل)	پیکربندی نهایی الگوریتم	0.739 ± 0.22	0.733 ± 0.214	—	بهینه‌ترین ترکیب از دید NMI و ACC

ساختار گراف از kNN متقابل به kNN ساده باعث افت محسوس عملکرد شد. گراف متقابل باعث حذف یال‌های غیرمعنادار می‌شود و در نتیجه، ساختار خوشه‌ای تمیزتری به دست می‌دهد.

در نهایت، مهم‌ترین بخش مطالعه ابلیشن مربوط به گام تعبیه طیفی است. اجرای مستقیم K میانگین استاندارد در فضای اصلی منجر به افت شدید (در NMI از 0.739 به 0.552) شد. این امر نشان می‌دهد فضای تعبیه طیفی با کاهش نویز و آشکارسازی ساختارهای غیرخطی، برای خوشه‌بندی ضروری است.

در مجموع، آزمایش‌ها نشان می‌دهند هر چهار مؤلفه کرنل آپولونیوس، لاپلاسی مناسب، گراف متقابل و تعبیه طیفی، در

نتایج جدول (۹) نشان می‌دهد هسته آپولونیوس نقشی کلیدی در عملکرد مدل دارد. جایگزینی آن با هسته گوسی باعث کاهش میانگین NMI از 0.739 به 0.691 شده است. این افت به ویژه در داده‌هایی با ساختار غیرخطی (مانند *Digits* و *Olivetti Faces*) بیشتر بوده است، زیرا کرنل آپولونیوس حساسیت کمتری به تغییرات محلی دارد و روابط فاصله‌ای را در بازه‌های طولانی‌تر به صورت دقیق‌تر مدل می‌کند.

تغییر نوع لاپلاسی از غیرنرمال‌شده به نرمال‌شده متقارن تأثیر کمی بر میانگین عملکرد گذاشت (حدود ۲ درصد کاهش)، اما در داده‌های نامتوازن مانند *Diabetes*، پایداری بیشتری در همگرایی خوشه‌ها ایجاد کرد. در مقابل، تغییر

ساختارهای محلی و غیرخطی داده‌ها را بهتر مدل‌سازی می‌کند. با این حال، در مجموعه داده‌هایی مانند Cancer، Iris و Breast، خوشه‌بندی طیفی عملکردی بهتر نسبت به الگوریتم پیشنهادی دارد. برای مثال، در مجموعه داده Iris، خوشه‌بندی طیفی با NMI (۰/۷۶) و دقت ۰/۹۰، نسبت به الگوریتم پیشنهادی با NMI (۰/۴۶۰) و دقت ۰/۶۸۷، نتایجی بهتر ارائه داده است. این امر ممکن است به دلیل سادگی ساختار داده‌های این مجموعه‌ها باشد که در آن، الگوریتم‌های مبتنی بر طیف قادر به شناسایی خوشه‌ها با دقت بیشتری هستند.

مراجع

- [1] A. Y. Ng, M. I. Jordan, Y. Weiss, "On spectral clustering: Analysis and an algorithm", in T. G. Dietterich, S. Becker, Z. Ghahramani (Eds.), *Advances in Neural Information Processing Systems 14 (NIPS 2001)*, Cambridge, MA: MIT Press, 2002. [Online]. Available: https://papers.nips.cc/paper_files/paper/2001/hash/801272ee79cfde7fa5960571fee36b9b-Abstract.html
- [2] J. MacQueen, "Some methods for classification and analysis of multivariate observations", in *Proc. 5th Berkeley Symp. Math. Statist. Probab.*, Vol. 1, pp. 281–297, Berkeley, CA: University of California Press, 1967. <https://www.semanticscholar.org/paper/ac8ab51a86f1a9ae74dd0e4576d1a019f5e654ed>
- [3] B. Park, C. Park, S. Hong, H. Choi, "Sparse kernel k-means clustering", *J. Appl. Stat.*, Vol. 52, No. 1, pp. 158–182, 2025. <https://doi.org/10.1080/02664763.2024.2362266>
- [4] H. He, "Clustering algorithm based on Hopkins statistics and K-means", *Multimedia Tools Appl.*, pp. 1–26, Feb. 2025. <https://doi.org/10.1007/s11042-025-20693-6>
- [5] X. Liu, "SimpleMKKM: Simple multiple kernel K-means", *IEEE Trans. Pattern Anal. Mach. Intell.*, Vol. 45, No. 4, pp. 5174–5186, Apr. 2023. <https://doi.org/10.1109/TPAMI.2022.3198638>
- [6] Z. Li, C. Tang, X. Zheng, Z. Wan, K. Sun, W. Zhang, X. Zhu, "Mutual structure learning for multiple kernel clustering", *Inf. Sci.*, Vol. 647, Art. No. 119445, Nov. 2023. <https://doi.org/10.1016/j.ins.2023.119445>
- [7] S. Pourbahrami, M. A. Balafar, L. M. Khanli, "ASVMK: A novel SVMs Kernel based on Apollonius function and density peak clustering", *Engineering Applications of*

کنار هم، موجب افزایش چشمگیر دقت و پایداری الگوریتم پیشنهادی شده‌اند.

تحلیل پیچیدگی زمانی الگوریتم پیشنهادی

در الگوریتم خوشه‌بندی طیفی بدون استفاده از PCA، ابتدا داده‌ها نرمال‌سازی می‌شوند که هزینه‌ای در حد $O(n*d)$ دارد. سپس، فاصله زوجی بین تمام نمونه‌ها در بُعد d محاسبه می‌شود که یکی از پرهزینه‌ترین مراحل است و زمان اجرای آن $O(n^2*d)$ است. ساخت گراف k -نزدیک‌ترین همسایه‌ها نیز معمولاً همان مرتبه پیچیدگی را دارد، مگر از روش‌های سریع‌تر استفاده شود. تشکیل ماتریس شباهت و محاسبه لاپلاسیان گراف نیز به ترتیب نیازمند زمان و حافظه‌ای در حدود $O(n*t*k*r)$ است. مرحله بعدی تجزیه ویژه لاپلاسیان است که اگر روی ماتریس چگال انجام شود، پیچیدگی زمانی $O(n^3)$ و حافظه‌ای $O(n^2)$ دارد. در پایان، الگوریتم K میانگین استلندارد روی embedding حاصل اجرا می‌شود که در هر تکرار هزینه‌ای در حد $O(t*n*k*r)$ (که t تکرار $r=k-1$) دارد. در مجموع، پیچیدگی کلی الگوریتم برابر $O(n*d) + O(n^2*d) + O(n^3) + O(n*t*k*r)$ است.

۵- نتیجه‌گیری

الگوریتم پیشنهادی با استفاده از هسته آپولونیوس و ساختار گراف، در برخی از مجموعه داده‌ها عملکردی بهتر نسبت به خوشه‌بندی طیفی دارد، اما در برخی دیگر، خوشه‌بندی طیفی نتایجی بهتر ارائه می‌دهد. بنابراین، انتخاب الگوریتم مناسب به ویژگی‌های داده‌ها و ساختار آنها بستگی دارد. بررسی عملکرد الگوریتم پیشنهادی مبتنی بر هسته آپولونیوس در مقایسه با خوشه‌بندی طیفی نشان می‌دهد این الگوریتم در برخی از مجموعه داده‌ها مانند Wine و Digits عملکردی بهتر دارد. برای مثال، در مجموعه داده Wine، الگوریتم پیشنهادی با NMI برابر ۰/۸۵۳ و دقت برابر ۰/۹۵۵، نسبت به خوشه‌بندی طیفی با NMI برابر ۰/۴۵ و دقت ۰/۷۳، نتایجی بهتر ارائه داده است. این بهبود ممکن است ناشی از استفاده از هسته آپولونیوس در ساخت گراف شباهت باشد که

- [11] S. Pourbahrami, M. A. Balafar, L. M. Khanli, Z. A. Kakarash, "A survey of neighborhood construction algorithms for clustering and classifying data points", *Computer Science Review*, Vol. 38, p. 100315, 2020.
<https://doi.org/10.1016/j.cosrev.2020.100315>
- [12] S. Pourbahrami, L. M. Khanli, S. Azimpour, "A novel and efficient data point neighborhood construction algorithm based on Apollonius circle", *Expert Systems with Applications*, Vol. 115, pp. 57-67, 2019.
<https://doi.org/10.1016/j.eswa.2018.07.066>
- [13] N. Abdolmaleki, L. M. Khanli, M. Hashemzadeh, S. Pourbahrami, "ACQC: Apollonius Circle-based Quantum Clustering", *Journal of Computational Science*, Vol. 64, p. 101877, 2022.
<https://doi.org/10.1016/j.jocs.2022.101877>
- Artificial Intelligence, Vol. 126, p. 106704, 2023.
<https://doi.org/10.1016/j.engappai.2023.106704>
- [8] Z. Zhang, X. Liu, L. Wang, "Spectral clustering algorithm based on improved Gaussian kernel function and beetle antennae search with damping factor", *Comput. Intell. Neurosci.*, Vol. 2020, Art. No. 1648573, 2020.
<https://doi.org/10.1155/2020/1648573>
- [9] H. Zhou, Z. Wang, H. Chen, X. Wang, "A novel spectral clustering algorithm based on neighbor relation and Gaussian kernel function with only one parameter", *Soft Comput.*, Vol. 28, No. 2, pp. 981-989, Jan. 2024.
<https://doi.org/10.1007/s00500-023-09309-z>
- [10] S. Pourbahrami, M. Hashemzadeh, "A geometric-based clustering method using natural neighbors", *Information Sciences*, Vol. 610, pp. 694-706, 2022.
<https://doi.org/10.1016/j.ins.2022.08.047>

¹ Spectral Clustering² k-Means³ Hopkins Statistic⁴ MSL-MKC⁵ Laplacian regularization⁶ Support Vector Machines (SVM)⁷ Spectral Clustering with Neighbor Relations⁸ Sparse Kernel k-Means for Feature Selection in Nonlinear Clustering (SKKM)⁹ smoothing factor¹⁰ Accuracy¹¹ Indicator Function¹² <https://archive.ics.uci.edu><https://www.kaggle.com>