

بهبود خوشه بندی K-Means با الگوریتم مگس میوه

جواد حمیدزاده^۱، علی زمانی خلیل آباد^۲، علی آرچین^۳

^۱استادیار، دانشکده مهندسی کامپیوتر و فناوری اطلاعات، دانشگاه صنعتی سجاد، مشهد، ایران

ایمیل: j_hamidzadeh@sadjad.ac.ir

^۲دانشجوی کارشناسی ارشد، مهندسی کامپیوتر نرم افزار، دانشگاه آزاد اسلامی، فردوس، ایران

ایمیل: ali.zamani78@gmail.com

^۳دانشجوی کارشناسی ارشد، مهندسی کامپیوتر نرم افزار، دانشگاه آزاد اسلامی، فردوس، ایران

ایمیل: aliarchin54@gmail.com

آدرس پست الکترونیکی نویسنده رابط: ali.zamani78@gmail.com

نام ارائه دهنده: علی زمانی خلیل آباد

کدمقاله:

خلاصه

خوشه بندی داده ها به کلاس ها و یا دسته های متناسب، یکی از مباحث مهم و مطرح در تشخیص الگوست. آنچه در خوشه بندی حائز اهمیت است انجام این کار به گونه ایست که، داده هایی که درست طبقه بندی نشده اند به حداقل برسند و یا به عبارت دیگر در هر کلاس داده هایی قرار بگیرند که حداکثر نزدیکی و مشابهت را با هم داشته باشند. هدف این مقاله اینست که با کمک الگوریتم بهینه سازی مگس میوه، مدل پیشنهادی جدیدی که آن را FOA-Clustering نام نهاده ایم، جهت بهبود روش K-Means معرفی کنیم. در پایان، روش مزبور بر روی مجموعه ای از داده ها، آزمایش شده است. نتایج، نشان دهنده برتری روش پیشنهادی ما نسبت به سایر روش های مرز دانش است.

کلمات کلیدی: تشخیص الگو، خوشه بندی، K-Means، الگوریتم بهینه سازی مگس میوه، FOA

۱. مقدمه

یکی از روش های حیاتی کنترل و مدیریت داده ها، خوشه بندی داده های با خواص مشابه، درون مجموعه ای از دسته ها است. خوشه بندی فرآیندی است که در آن مجموعه ای از اشیاء داده به گروه های مجزایی از کلاس، خوشه، تقسیم می شوند، به طوری که اشیاء یک کلاس تا حد امکان به یکدیگر شبیه بوده و با اشیاء دیگر کلاس ها، متفاوت می باشند. خوشه بندی در زمینه های بسیاری از جمله در شناسایی الگو، یادگیری ماشین، داده کاوی، بازیابی اطلاعات و داده ورزی زیستی کاربرد دارد. این روش ابزاری برای طبقه بندی غیر نظارت شده داده ها است، در حقیقت خوشه بندی یک تکنیک جداسازی بدون ناظر است که معمولاً داده ها را به صورت بردارهایی در یک فضای چندبُعدی دریافت کرده و آن ها را در دسته ها یا خوشه هایی قرار می دهد. با این ویژگی که داده های یک خوشه از یک تشابهی سود می برند درحالی که داده های خوشه های مختلف، بالعکس، تشابهی باهم ندارند. بنابراین باید به دنبال معیاری برای مشابهت دسته ای از داده ها باشیم به گونه ای که این داده ها در خوشه یکسانی قرار گیرند. هدف اصلی خوشه بندی K-Means این است که مجموع عدم تشابه بین تمام اشیاء یک خوشه از مراکز خوشه های متناظرشان کمترین باشد. مهم ترین مشکل الگوریتم K-Means [۱] این است که نتایج

$$D = ||X - Z||$$

(۱)