

الگوریتم تعمیم یافته یادگیری دیکشنری حداقل مربعات بازگشتی با وزن وفقی و سازگار با برچسب

محدثه یوسفی^۱، دانشجوی کارشناسی ارشد، یاشار نادراحمیدیان^۲، استادیار

۱- گروه برق - دانشکده فنی - دانشگاه گیلان - رشت - ایران - mohadesehyousef@msc.guilan.ac.ir

۲- گروه برق - دانشکده فنی - دانشگاه گیلان - رشت - ایران - yna@guilan.ac.ir

چکیده: در میان روش‌های یادگیری دیکشنری بدون نظارت، الگوریتم تعمیم یافته حداقل مربعات بازگشتی با وزن وفقی (GAW-RLS) عملکرد قابل توجهی را در نمایش داده‌ها نشان داده است. در این مقاله از این الگوریتم برای به روزرسانی دیکشنری به صورت نظارت شده استفاده می‌شود. در تابع هزینه روش پیشنهادی، شرط جدیدی به عنوان خطای نمایش تُنک متمایزکننده به خطای نمایش و خطای طبقه‌بندی اضافه شده و دیکشنری، پارامتر نمایش تُنک متمایزکننده و پارامتر طبقه‌بندی کننده به طور هم‌زمان برای کاربرد طبقه‌بندی تصویر آموزش داده می‌شوند. روش جدید به عنوان الگوریتم تعمیم یافته حداقل مربعات بازگشتی با وزن وفقی و سازگار با برچسب (LC-GAWRLS) معرفی می‌شود. روش پیشنهادی از یک وزن تصحیح اضافی برای کنترل تاثیر داده‌های آموزشی جدید در فرآیند به‌روزرسانی دیکشنری استفاده می‌کند و در مقایسه با روش‌های یادگیری دیکشنری قبلی، مقاومت مدل را در برابر تغییرات داده‌های آموزشی افزایش می‌دهد. نتایج شبیه‌سازی نشان می‌دهد که روش پیشنهادی به‌ویژه در شرایطی که تعداد داده‌های آموزشی محدود است، دقت طبقه‌بندی بالاتری را ارائه می‌دهد.

واژه‌های کلیدی: یادگیری دیکشنری، نمایش تُنک، الگوریتم وفقی، طبقه‌بندی سازگار با برچسب

Label-Consistent Dictionary Learning via Generalized Adaptive Weighted Recursive Least Squares Algorithm

Mohadeseh Yousefi, Master Student¹, Yashar Naderahmadian, Assistant Professor²

1- Engineering Department, University of Guilan, Rasht, Iran, Email: mohadesehyousef@msc.guilan.ac.ir

2- Engineering Department, University of Guilan, Rasht, Iran, Email: yna@guilan.ac.ir

Abstract: The generalized adaptive weighted recursive least squares (GAW-RLS) algorithm has demonstrated exceptional performance in data representation within unsupervised dictionary learning methods. This paper extends the GAW-RLS algorithm to supervised settings for dictionary updating. In the proposed method, the cost function incorporates a novel condition termed discriminant sparse representation error, which is added to the representation error and classification error. The dictionary, discriminant sparse representation parameter, and classifier parameter are jointly trained to enhance image classification performance. This approach introduces the Label-Consistent Generalized Adaptive Weighted Recursive Least Squares (LC-GAWRLS) algorithm, which leverages an additional correction weight to regulate the influence of new training data during the dictionary updating process. This enhancement improves robustness against variations in training data compared to existing dictionary learning methods. Simulation results confirm that the LC-GAWRLS algorithm achieves superior classification accuracy, particularly in scenarios with limited training data, demonstrating its potential for effective supervised learning in challenging conditions.

Keywords: Dictionary Learning, Sparse representation, Adaptive algorithm, Label consistent classification

تاریخ ارسال مقاله: ۱۴۰۳/۱۰/۱۴

تاریخ اصلاح مقاله: ۱۴۰۴/۰۵/۰۲

تاریخ پذیرش مقاله: ۱۴۰۴/۰۶/۱۵

نام نویسنده مسئول: یاشار نادراحمیدیان

نشانی نویسنده مسئول: ایران - رشت - دانشگاه گیلان - دانشکده فنی - گروه برق.

۱- مقدمه

یادگیری دیکشنری برای نمایش تئیک داده نتایج چشمگیری در کاربردهای مختلف مانند تشخیص چهره [۱]، [۲]، طبقه‌بندی تصویر [۳]، حذف نویز تصویر [۴] و بازسازی تصویر [۵] به دست آورده است. این روش داده را به صورت ترکیب خطی از ستون‌های دیکشنری (اتم) نمایش می‌دهد. الگوریتم یادگیری دیکشنری شامل دو مرحله تخمین نمایش تئیک و به‌روزرسانی دیکشنری می‌باشد و بین این دو مرحله متناوب است. در مرحله نمایش تئیک، دیکشنری معلوم فرض می‌شود و نمایش تئیک به کمک الگوریتم‌های OMP، ISTA و FISTA [۶] - [۹] تخمین زده می‌شود. سپس برای به‌روزرسانی دیکشنری با ثابت در نظر گرفتن نمایش تئیک از الگوریتم‌های KSVD، MOD، RLS و GAW-RLS [۱۰] - [۱۲] استفاده می‌شود. الگوریتم‌های یادگیری دیکشنری را می‌توان بر اساس استفاده آن‌ها از برچسب مجموعه داده آموزشی به دو دسته یادگیری دیکشنری بدون نظارت و نظارت‌شده تقسیم‌بندی کرد.

یادگیری دیکشنری بدون نظارت: الگوریتم‌هایی مانند MOD، RLS و GAW-RLS بدون استفاده از برچسب داده‌های مجموعه آموزشی، دیکشنری را تعلیم می‌دهند. اگرچه این روش‌ها در طبقه‌بندی تصویر مورد استفاده قرار گرفته‌اند اما عملکرد آن‌ها محدود بوده است [۱۱]، [۱۳].

یادگیری دیکشنری نظارت‌شده: الگوریتم‌های یادگیری دیکشنری نظارت‌شده با استفاده از برچسب داده‌های مجموعه آموزشی برای کاربردهای خاصی تنظیم شده‌اند که از جمله کاربردهای طبقه‌بندی، تشخیص و... را می‌توان نام برد.

در سال‌های اخیر، الگوریتم‌های یادگیری دیکشنری نظارت‌شده به‌ویژه در حوزه‌هایی همچون طبقه‌بندی و تشخیص الگو، پیشرفت قابل توجهی داشته‌اند. روش‌هایی نظیر LC-KSVD^۲ [۱۴]، LC-RLS^۱ [۱۵] و UDLA^۱ [۱۶] از جمله الگوریتم‌های رایج در این زمینه محسوب می‌شوند. الگوریتم LC-KSVD^۲ به دلیل نیاز مکرر به دسترسی کامل به مجموعه داده‌های برچسب‌دار از مصرف حافظه بالایی برخوردار است و اجرای آن در مقیاس‌های بزرگ با مشکلاتی همراه است [۱۴]. در مقابل، الگوریتم LC-RLS برای محیط‌هایی با منابع حافظه محدود مناسب‌تر بوده و قادر است دیکشنری را به صورت بازگشتی و با ورود هر داده جدید به‌روزرسانی کند [۱۵]. در روش UDLA نیز برای هر کلاس یک دیکشنری مجزا آموزش داده می‌شود و تمرکز اصلی، برخلاف سایر روش‌های مبتنی بر نمایش تئیک بر کاهش خطای طبقه‌بندی است [۱۶]. با این حال، به دلیل استفاده از چندین دیکشنری و به کارگیری الگوریتم SVD در فرآیند به‌روزرسانی، این روش دارای پیچیدگی محاسباتی بالا و مصرف حافظه زیاد است. از طرفی هیچ کدام از روش‌های ذکر شده فاقد راهکار مؤثری برای مقابله با داده‌های نویزی یا برچسب‌گذاری نادرست داده‌های آموزشی می‌باشند.

در سوی دیگر، رویکردهای اخیر به سمت یادگیری دیکشنری عمیق (چندلایه) گرایش یافته‌اند. الگوریتم‌هایی مانند HILADLE^{۱۱} [۱۷]، DDL^{۱۲} [۱۸]، MDDL^{۱۳} [۱۹]، DDDL^{۱۴} [۲۰]، DSDL^{۱۵} [۲۱] و M-LCLEDL^{۱۶} [۲۲] با هدف افزایش دقت طبقه‌بندی از چندین لایه دیکشنری برای استخراج ویژگی‌های سطح بالاتر از داده‌ها بهره گرفته‌اند. با وجود دقت قابل قبول این روش‌ها، چالش‌هایی نظیر هزینه محاسباتی بالا، نیاز به داده‌های آموزشی بیشتر و پیچیدگی در طراحی و پیاده‌سازی به‌ویژه در محیط‌های با منابع محدود، کاربرد عملی آن‌ها را با محدودیت مواجه ساخته است.

در بسیاری از کاربردهای عملی یادگیری ماشین، داده‌های آموزشی معمولاً با سه چالش اساسی روبرو هستند: وجود نویز، برچسب‌گذاری نادرست و محدودیت در حجم داده‌های آموزشی. این عوامل می‌توانند تأثیر مخربی بر عملکرد مدل‌های طبقه‌بندی داشته و دقت پیش‌بینی را به میزان قابل توجهی کاهش دهند. از این رو، توسعه الگوریتم‌هایی که در برابر نویز و برچسب‌های اشتباه مقاوم باشند و در شرایط محدودیت منابع پردازشی و حافظه و یا محدود بودن نمونه‌های آموزشی نیز عملکرد مناسبی ارائه دهند اهمیت زیادی دارد.

در این مقاله، روشی آنلاین برای یادگیری دیکشنری نظارت‌شده ارائه می‌کنیم که به طور هم‌زمان چندین چالش کلیدی را مورد توجه قرار می‌دهد. روش پیشنهادی با کمک یک مکانیزم وزن‌دهی وقفی، مطابقت داده آموزشی جدید با مدل موجود را ارزیابی کرده و سهم هر داده آموزشی را در فرآیند یادگیری تنظیم می‌کند. این رویکرد چند مزیت مهم دارد: (۱) اثر نمونه‌های پرت یا دارای برچسب نادرست را در به‌روزرسانی مدل کاهش می‌دهد و موجب افزایش پایداری مدل در مواجهه با داده‌های پرت و برچسب‌گذاری نادرست می‌شود. (۲) داده‌ها را به صورت آنلاین و نمونه به نمونه پردازش می‌کند و نیازی به ذخیره‌سازی مجموعه داده در حافظه ندارد که این ویژگی آن را به گزینه‌ای مناسب برای سناریوهای با محدودیت منابع محاسباتی و حافظه‌ای تبدیل می‌سازد. (۳) روش پیشنهادی در دسته‌ی الگوریتم‌های یادگیری دیکشنری نظارت‌شده تک‌لایه قرار می‌گیرد. از این رو پیچیدگی محاسباتی کمتری نسبت به روش‌های یادگیری دیکشنری عمیق داشته و قدرت پردازشی کمتری نیاز دارد. (۴) روش پیشنهادی قادر است با حجم کمتری از داده‌های آموزشی، دقتی قابل مقایسه با روش‌های عمیق ارائه دهد. مجموعه این مزیت‌ها، روش پیشنهادی را به انتخابی بهینه و مناسب برای محیط‌های با منابع محدود (حافظه، داده و قدرت پردازشی) تبدیل می‌کند که قابل رقابت با رویکردهای پیشین است.

ایده‌ی اولیه‌ی روش پیشنهادی نخست در مقاله [۲۳] معرفی شد. در این مقاله، با توسعه‌ی چارچوب نظری، استخراج روابط ریاضی دقیق، تحلیل پیچیدگی محاسباتی و ارائه‌ی نتایج تجربی گسترده‌تر، نسخه‌ی کامل‌تر و تحلیلی عمیق‌تر از آن روش ارائه شده است.

۲- طرح مساله

۲-۱- الگوریتم‌های یادگیری دیکشنری بدون نظارت

با در نظر گرفتن مجموعه‌ای از داده‌های آموزشی به صورت $\mathbf{Y} = [\mathbf{Y}_1, \mathbf{Y}_2, \dots, \mathbf{Y}_L] \in \mathbb{R}^{n \times L}$ می‌توان این مجموعه با L داده آموزشی را به صورت زیر نمایش داد:

$$\mathbf{Y} = \mathbf{D}\mathbf{X} \quad (۱)$$

در رابطه (۱)، ماتریس $\mathbf{X} = [\mathbf{X}_1, \mathbf{X}_2, \dots, \mathbf{X}_L] \in \mathbb{R}^{K \times L}$ نمایش تُنک مجموعه داده‌های آموزشی $\mathbf{Y}$ و ماتریس $\mathbf{D} = [\mathbf{D}_1, \mathbf{D}_2, \dots, \mathbf{D}_K] \in \mathbb{R}^{n \times K}$ تعداد ستون‌ها (اتم‌های دیکشنری) بیشتر از تعداد سطرها (بعد داده) می‌باشد. در این حالت دیکشنری $\mathbf{D}$ یک دیکشنری بیش کامل نامیده می‌شود و با در نظر گرفتن این دیکشنری بیش کامل، نمایش یکتایی برای داده $\mathbf{Y}$ وجود ندارد. برای یکتا بودن نمایش، نمایش تُنک $\mathbf{X}$ با اضافه کردن قید تُنکی می‌کند [۱۲]. با این فرض هدف اصلی الگوریتم‌های یادگیری دیکشنری یافتن دیکشنری و نمایش تُنک بهینه برای کمینه کردن خطای نمایش زیر است:

$$\langle \mathbf{D}, \mathbf{X} \rangle = \arg \min_{\mathbf{D}, \mathbf{X}} \|\mathbf{Y} - \mathbf{D}\mathbf{X}\|_F^2, \quad s.t. \quad \forall i \in \{1:L\} \|\mathbf{X}_i\|_0 \leq S, \quad \forall k \in \{1:K\} \|\mathbf{D}_k\|_2 = 1 \quad (۲)$$

در رابطه (۲) $\|\cdot\|_F$ نرم فروبنیوس می‌باشد و برای ماتریس نمونه $\mathbf{A}$ به صورت زیر محاسبه می‌شود:

$$\|\mathbf{A}\|_F = \sqrt{\sum_{i,j} a_{i,j}^2} \quad (۳)$$

نرم صفر به صورت $\|\cdot\|_0$ نشان داده می‌شود. این نرم در رابطه (۲) تعداد عناصر غیر صفر بردار $\mathbf{X}_i$ را نشان می‌دهد و S به عنوان شرط تُنکی، تعداد ضرایب غیر صفر را در $\mathbf{X}_i$ مشخص می‌کند. نرم دو با $\|\cdot\|_2$ مشخص شده و برای بردار نمونه $\mathbf{A}$ به صورت زیر محاسبه می‌شود:

$$\|\mathbf{A}\|_2 = \sqrt{\sum_i a_i^2} \quad (۴)$$

در رابطه (۲) $\|\mathbf{D}_k\|_2 = 1$ قید نرمالیزه کردن برای اتم‌های دیکشنری است. بدون نرمالیزه کردن اتم‌های دیکشنری این اتم‌ها در طول فرآیند یادگیری می‌توانند دارای مقادیر بسیار بزرگ باشند، در نتیجه ضرایب تُنک مرتبط با این اتم‌های بزرگ نزدیک به صفر شده و مشارکت آن‌ها در نمایش داده کاهش می‌یابد. تابع هزینه (۲) غیر محدب است به دلیل اینکه هر دو متغیر $\mathbf{D}$ و $\mathbf{X}$ در آن مجهول است. برای تبدیل آن به یک تابع هزینه محدب، الگوریتم یادگیری دیکشنری بین دو گام به صورت متناوب در حال تکرار است [۲۴]. در گام اول، دیکشنری $\mathbf{D}$ ثابت فرض می‌شود و نمایش تُنک بهینه برای هر داده $\mathbf{Y}_i \in \mathbb{R}^{n \times 1}$ از مجموعه داده $\mathbf{Y}$ با حل تابع هزینه زیر محاسبه می‌شود:

$$\mathbf{X}_i = \arg \min_{\mathbf{X}} \|\mathbf{Y}_i - \mathbf{D}\mathbf{X}\|_2^2, \quad s.t. \quad \forall i, \quad \|\mathbf{X}_i\|_0 \leq S \quad (۵)$$

برای یافتن نمایش تُنک $\mathbf{X}_i$ در تابع هزینه (۵)، الگوریتم‌هایی مانند OMP و ℓ_1 LASSO [۲۵]، [۲۶] ارائه شده‌اند. از این نمایش تُنک برای به‌روزرسانی اتم‌های دیکشنری در گام دوم استفاده می‌شود. در این گام، تابع هزینه یادگیری دیکشنری به صورت زیر در نظر گرفته می‌شود:

$$\hat{\mathbf{D}} = \arg \min_{\mathbf{D}} \|\mathbf{Y} - \mathbf{D}\mathbf{X}\|_F^2, \quad s.t. \quad \forall k \in \{1:K\}, \|\mathbf{D}_k\|_2 = 1 \quad (۶)$$

الگوریتم‌هایی مانند MOD، KSVD، RLS و GAWRLS برای یافتن پاسخ بهینه (۶) مورد استفاده قرار می‌گیرند.

۲-۲- الگوریتم‌های یادگیری دیکشنری نظارت‌شده

الگوریتم‌های یادگیری دیکشنری نظارت‌شده از داده‌های برچسب‌دار در فرآیند یادگیری دیکشنری استفاده می‌کنند. به کار بردن برچسب‌های کلاس مجموعه آموزشی، عملکرد روش‌های یادگیری دیکشنری نظارت شده را در مقایسه با روش‌های بدون نظارت در طبقه‌بندی تصویر بهبود می‌بخشد [۲۰]، [۲۷]. استفاده از این برچسب‌ها به الگوریتم‌های نظارت‌شده این امکان را می‌دهد تا ویژگی‌های متمایزتری را از تصویر اصلی استخراج کنند. در نتیجه دقت طبقه‌بندی روش‌های نظارت شده در مقایسه با روش‌های بدون نظارت در طبقه‌بندی تصویر افزایش می‌یابد. تابع هزینه الگوریتم یادگیری دیکشنری نظارت‌شده به صورت رابطه (۷) است [۱۴]:

$$\langle \mathbf{D}, \mathbf{A}, \mathbf{G}, \mathbf{X} \rangle = \arg \min_{\mathbf{D}, \mathbf{A}, \mathbf{G}, \mathbf{X}} \|\mathbf{Y} - \mathbf{D}\mathbf{X}\|_F^2 + \alpha \|\mathbf{Q} - \mathbf{A}\mathbf{X}\|_F^2 + \beta \|\mathbf{H} - \mathbf{G}\mathbf{X}\|_F^2 \quad (۷)$$

جمله اول در رابطه (۷) خطای نمایش تُنک است که در این خطا، مجموعه‌ای شامل L داده آموزشی به صورت $\mathbf{Y} \in \mathbb{R}^{n \times L}$ مشخص می‌شود. n بعد داده آموزشی را در این مجموعه داده نشان می‌دهد. $\mathbf{X} \in \mathbb{R}^{K \times L}$ نیز نمایش تُنک مربوط به مجموعه داده آموزشی $\mathbf{Y}$ است. در جمله اول رابطه (۷) دیکشنری $\mathbf{D} \in \mathbb{R}^{n \times K}$ را می‌توان دیکشنری بیش کامل یا کم کامل در نظر گرفت. در دیکشنری بیش کامل، تعداد اتم‌های دیکشنری (تعداد ستون‌های $\mathbf{D}$) بیشتر از بعد داده آموزشی (تعداد سطرها $\mathbf{D}$) می‌باشد. در دیکشنری کم کامل تعداد ستون‌ها کمتر از تعداد سطرها است و در مقایسه با دیکشنری بیش کامل تعداد پارامترهای کمتری دارد، در نتیجه از نظر محاسباتی کارآمدتر است [۱۶]، [۲۷].

جمله دوم رابطه (۷) خطای نمایش تُنک متمایزکننده است. این خطا تمایز پذیری را بین نمایش تُنک کلاس‌های مختلف افزایش می‌دهد. ماتریس $\mathbf{Q} \in \mathbb{R}^{K \times L}$ ماتریس سازگاری برچسب است و هر ستون از این ماتریس به صورت بردار $\mathbf{Q}_i = [0, 0, \dots, 1, 1, 0, \dots, 0]^T \in \mathbb{R}^{K \times 1}$ زمانی مشاهده می‌شوند که اتم $\mathbf{D}_k$ و داده آموزشی $\mathbf{Y}_i$ به یک کلاس تعلق داشته باشند. ماتریس $\mathbf{A} \in \mathbb{R}^{K \times K}$ ماتریس تبدیل خطی است. این ماتریس، نمایش تُنک $\mathbf{X}$ را به فضای ویژگی تُنک متمایز در ماتریس $\mathbf{Q}$ نگاشت می‌کند [۱۴].

عملکرد و مقاومت روش پیشنهادی LC-GAWRLS در برابر داده‌های آموزشی ناسازگار در مقایسه با روش LC-RLS بهبود می‌یابد. روش پیشنهادی LC-GAWRLS در دو مرحله یادگیری و ارزیابی معرفی می‌شود.

۳-۱- مرحله آموزش و یادگیری روش نظارت‌شده LC-GAWRLS

برای بهینه‌سازی تابع هزینه الحاقی یادگیری دیکشنری نظارت‌شده (۱۰) از دو مرحله استفاده می‌شود. در مرحله اول با فرض ثابت بودن $\mathbf{D}_{aug}$ ، نمایش $\mathbf{X}$ با استفاده از الگوریتم‌های نمایش $\mathbf{X}$ مانند OMP تخمین زده می‌شود و سپس در مرحله دوم این نمایش $\mathbf{X}$ برای به‌روزرسانی دیکشنری به‌کار می‌رود. با مشتق گرفتن از رابطه (۱۰) نسبت به $\mathbf{D}_{aug}$ به‌روزرسانی دیکشنری به‌صورت زیر محاسبه می‌شود:

$$\mathbf{D}_{aug} = \mathbf{Y}_{aug} \mathbf{X}^T (\mathbf{X} \mathbf{X}^T)^{-1} \quad (11)$$

رابطه (۱۱) به‌دلیل محاسبه ماتریس معکوس پیچیدگی محاسباتی بالایی دارد. برای خودداری از محاسبه ماتریس معکوس در رابطه (۱۱) و به‌روزرسانی دیکشنری به‌صورت بازگشتی، در روش RLS تابع هزینه زیر در نظر گرفته می‌شود [۱۲]:

$$(\mathbf{D}_{aug})_i = \arg \min_{\mathbf{D}_{aug}} \lambda \left\| (\mathbf{Y}_{aug})_{i-1} - \mathbf{D}_{aug} \mathbf{X}_{i-1} \right\|_F^2 + \left\| (\mathbf{Y}_{aug})_i - \mathbf{D}_{aug} \mathbf{X}_i \right\|_2^2 \quad (12)$$

در رابطه (۱۲) ضریب فراموشی λ عددی بین صفر و یک است و برای کنترل حافظه و تاثیر داده‌های آموزشی قبلی در مرحله به‌روزرسانی دیکشنری استفاده می‌شود. در رابطه (۱۲) تمام داده‌های آموزشی جدید دارای تاثیر یکسانی در به‌روزرسانی دیکشنری می‌باشند. در نتیجه روش RLS در مواجهه با ناسازگاری در داده‌های آموزشی جدید مقاوم نمی‌باشد [۱۲].

در روش پیشنهادی برای به‌روزرسانی $\mathbf{D}_{aug}$ به‌صورت بازگشتی از روش GAW-RLS استفاده می‌شود. در روش GAW-RLS وزن تصحیحی به رابطه (۱۲) اضافه شده و تابع هزینه به‌روزرسانی دیکشنری به‌صورت زیر در نظر گرفته می‌شود:

$$(\mathbf{D}_{aug})_i = \arg \min_{\mathbf{D}_{aug}} \lambda \left\| (\mathbf{Y}_{aug})_{i-1} - \mathbf{D}_{aug} \mathbf{X}_{i-1} \right\|_F^2 + w_i \left\| (\mathbf{Y}_{aug})_i - \mathbf{D}_{aug} \mathbf{X}_i \right\|_2^2 \quad (13)$$

در رابطه (۱۳) w_i وزن تصحیحی است که در روش پیشنهادی، سازگاری داده‌های جدید را با دیکشنری فعلی کنترل می‌کند [۱۲]. در رابطه (۱۳) ماتریس‌های فعلی $\mathbf{X}_i$ و $(\mathbf{Y}_{aug})_i$ به ماتریس‌های متناظر خود در مرحله قبل یعنی $\mathbf{X}_{i-1}$ و $(\mathbf{Y}_{aug})_{i-1}$ به‌صورت رابطه (۱۴) مرتبط می‌باشند:

جمله سوم در رابطه (۷) خطای طبقه‌بندی است. در این خطا ماتریس برچسب کلاس به‌صورت $\mathbf{H} \in \mathbb{R}^{m \times L}$ مشخص می‌شود که m تعداد کلاس‌ها را در مجموعه داده نشان می‌دهد. هر ستون از $\mathbf{H}$ بردار برچسب داده Y_i می‌باشد که به‌صورت $\mathbf{H}_i = [0, 0, \dots, 1, \dots, 0]^T \in \mathbb{R}^{m \times 1}$ نشان داده می‌شود. مکان مقدار غیر صفر در این بردار، کلاس داده Y_i را مشخص می‌کند. ماتریس $\mathbf{G} \in \mathbb{R}^{m \times K}$ ماتریس طبقه‌بندی پیش‌بینی کننده خطی است که در طول فرآیند یادگیری محاسبه می‌شود [۱۴]. ضرایب α و β نیز به‌ترتیب تاثیر خطای نمایش $\mathbf{X}$ متمایزکننده و خطای طبقه‌بندی را در رابطه (۷) کنترل می‌کنند. رابطه (۷) را می‌توان به‌صورت ماتریس الحاقی زیر در نظر گرفت [۱۴]:

$$\langle \mathbf{D}, \mathbf{A}, \mathbf{G}, \mathbf{X} \rangle = \arg \min_{\mathbf{D}, \mathbf{A}, \mathbf{G}, \mathbf{X}} \left\| \begin{bmatrix} \mathbf{Y} \\ \sqrt{\alpha} \mathbf{Q} \end{bmatrix} - \begin{bmatrix} \mathbf{D} \\ \sqrt{\alpha} \mathbf{A} \end{bmatrix} \begin{bmatrix} \mathbf{X} \\ \sqrt{\beta} \mathbf{G} \end{bmatrix} \right\|_F^2 \quad (8)$$

با در نظر گرفتن ماتریس‌های الحاقی $\mathbf{Y}_{aug}$ و $\mathbf{D}_{aug}$ به‌صورت زیر:

$$\mathbf{Y}_{aug} = \begin{bmatrix} \mathbf{Y}^T, \sqrt{\alpha} \mathbf{Q}^T, \sqrt{\beta} \mathbf{H}^T \end{bmatrix}^T \in \mathbb{R}^{(n+K+m) \times L}, \quad (9)$$

$$\mathbf{D}_{aug} = \begin{bmatrix} \mathbf{D}^T, \sqrt{\alpha} \mathbf{A}^T, \sqrt{\beta} \mathbf{G}^T \end{bmatrix}^T \in \mathbb{R}^{(n+K+m) \times K}$$

و جایگذاری آن‌ها در رابطه (۸)، رابطه (۷) به‌صورت زیر تغییر پیدا می‌کند:

$$\langle \mathbf{D}_{aug}, \mathbf{X} \rangle = \arg \min_{\mathbf{D}_{aug}, \mathbf{X}} \left\| \mathbf{Y}_{aug} - \mathbf{D}_{aug} \mathbf{X} \right\|_F^2 \quad (10)$$

تابع هزینه به‌دست آمده در رابطه (۱۰) را می‌توان با روش‌های یادگیری دیکشنری کمینه کرد. در روش‌های یادگیری دیکشنری نظارت‌شده مانند LC-KSVD2 و LC-RLS، تابع هزینه (۱۰) به‌ترتیب با روش‌های KSVD و RLS کمینه می‌شود [۱۴]، [۱۵].

۳- راه حل پیشنهادی

الگوریتم‌های یادگیری دیکشنری نظارت‌شده مانند LC-KSVD2 و LC-RLS از اطلاعات برچسب داده‌های آموزشی برای تعلیم دیکشنری متمایز استفاده می‌کنند. به‌دلیل ماهیت پردازش دسته‌ای روش LC-KSVD2، این روش در هر تکرار باید به کل مجموعه داده آموزشی دسترسی داشته باشد. در نتیجه چالشی اساسی در سناریوهایی با منابع حافظه محدود ایجاد می‌کند [۱۴]. روش LC-RLS با استفاده از روش آنلاین و بازگشتی RLS این مشکل را برطرف می‌کند [۱۵]. در روش RLS از یک ضریب فراموشی برای تنظیم حافظه داده‌های گذشته استفاده می‌شود. اما با تمام داده‌های ورودی جدید به‌طور برابر رفتار می‌کند، در نتیجه روش LC-RLS در برابر تغییرات و ناسازگاری داده‌های آموزشی مقاوم نمی‌باشد.

روش پیشنهادی LC-GAWRLS با در نظر گرفتن یک وزن تصحیح برای داده‌های آموزشی در طول فرآیند یادگیری، به‌طور وقتی تاثیر داده آموزشی جدید را بر به‌روزرسانی دیکشنری بر اساس سازگاری نسبی آن با تخمین فعلی مدل تنظیم می‌کند. در نتیجه

وزن تصحیحی پیشنهادی w_i را می‌توان مانند روش GAW-RLS با کمک بردار خطای نمایش به‌دست آمده از رابطه (۲۲) برای هر داده ورودی جدید به‌صورت زیر محاسبه کرد:

$$w_i = \left(\xi + \sqrt{R_i^T R_i} \right)^{-1} \quad (27)$$

در رابطه بالا ξ مقدار کوچک مثبتی برای نرمالیزه کردن وزن تصحیح پیشنهادی w_i است. در روش پیشنهادی LC-GAWRLS مقدار ξ برابر با یک در نظر گرفته می‌شود. پس از به‌روزرسانی $\mathbf{D}_{aug}$ از رابطه (۲۵)، دیکشنری $\mathbf{D}$ ماتریس تبدیل خطی $\mathbf{A}$ و ماتریس طبقه‌بندی کننده $\mathbf{G}$ با استفاده از رابطه (۹) از $\mathbf{D}_{aug}$ استخراج شده و سپس به‌صورت زیر نرمالیزه می‌شوند:

$$\hat{\mathbf{D}} = \left\{ \frac{D_1}{\|D_1\|_2}, \dots, \frac{D_K}{\|D_K\|_2} \right\}, \quad \hat{\mathbf{A}} = \left\{ \frac{A_1}{\|A_1\|_2}, \dots, \frac{A_K}{\|A_K\|_2} \right\} \quad (28)$$

$$\hat{\mathbf{G}} = \left\{ \frac{G_1}{\|G_1\|_2}, \dots, \frac{G_K}{\|G_K\|_2} \right\}$$

مرحله آموزش و یادگیری روش LC-GAWRLS در الگوریتم ۱ خلاصه شده است.

الگوریتم ۱
ورودی: $\mathbf{Y} \in \mathbb{R}^{n \times L}$, $\mathbf{Q} \in \mathbb{R}^{K \times L}$, $\mathbf{H} \in \mathbb{R}^{m \times L}$, سطح اسپارسیته K ، پارامترهای α , β و تعداد تکرار T
مقداردهی اولیه: ماتریس‌های $\mathbf{D}^{(0)}$, $\mathbf{A}^{(0)}$ و $\mathbf{G}^{(0)}$ را با روش پیشنهادی در [۱۴] و ماتریس $\mathbf{C}^{(0)}$ را با استفاده از ماتریس همانی مقدار دهی کنید.
شروع الگوریتم:
۱) برای تعداد تکرار T تا ۱ انجام دهید.
۲) برای نمونه $i = 1, \dots, L$ انجام دهید.
۳) نمونه جدید Y_i را دریافت کنید.
۴) ماتریس‌های الحاقی را با استفاده از (۹) و (۱۵) تشکیل دهید.
۵) نمایش $\mathbf{X}_i$ را محاسبه کنید:
$\mathbf{X}_i = \text{OMP}(\mathbf{Y}_{aug}, (\mathbf{D}_{aug})_{i-1}, S)$
۶) بردار خطای نمایش را محاسبه کنید:
$\mathbf{R}_i = (\mathbf{Y}_{aug})_i - (\mathbf{D}_{aug})_{i-1} \mathbf{X}_i$
۷) ضریب تصحیح w_i را از رابطه (۲۷) محاسبه کنید.
۸) بردار U_i را حساب کنید: $U_i = \lambda^{-1} \mathbf{C}_{i-1}^{-1} \mathbf{X}_i$
۹) اسکالر δ_i را حساب کنید: $\delta_i = \frac{w_i}{1 + w_i \mathbf{X}_i^T U_i}$
۱۰) دیکشنری را با استفاده از رابطه (۲۵) به‌روزرسانی کنید.
۱۱) ماتریس $\mathbf{C}_i^{-1}$ از رابطه (۲۶) به‌روزرسانی کنید.
۱۲) پایان
۱۳) $\mathbf{D}$, $\mathbf{A}$ و $\mathbf{G}$ را از $\mathbf{D}_{aug}$ استخراج کنید.
۱۴) $\mathbf{D}$, $\mathbf{A}$ و $\mathbf{G}$ را نرمالیزه کنید.
۱۵) پایان تعداد تکرار
خروجی: $\mathbf{D}$, $\mathbf{A}$ و $\mathbf{G}$

$$\begin{aligned} (\mathbf{Y}_{aug})_i &= \left[\sqrt{\lambda} (\mathbf{Y}_{aug})_{i-1}, \sqrt{w_i} (Y_{aug})_i \right] \\ \mathbf{X}_i &= \left[\sqrt{\lambda} \mathbf{X}_{i-1}, \sqrt{w_i} \mathbf{X}_i \right] \end{aligned} \quad (14)$$

در رابطه (۱۴) ماتریس الحاقی $\mathbf{Y}_{aug} \in \mathbb{R}^{n \times 1}$ برای نمونه آموزشی جدید به‌صورت زیر تشکیل می‌شود:

$$\mathbf{Y}_{aug} = \left[\mathbf{Y}_i^T, \sqrt{\alpha} \mathbf{Q}_i^T, \sqrt{\beta} \mathbf{H}_i^T \right]^T \in \mathbb{R}^{(n+K+m) \times 1} \quad (15)$$

در رابطه (۱۵) بردارهای $\mathbf{H}_i \in \mathbb{R}^{m \times 1}$ و $\mathbf{Q}_i \in \mathbb{R}^{K \times 1}$ به‌ترتیب بردار برچسب و بردار سازگاری برچسب داده ورودی Y_i می‌باشند. با در نظر گرفتن رابطه (۱۱) و (۱۴)، برای به‌روزرسانی دیکشنری به‌صورت بازگشتی ابتدا ماتریس‌های زیر در نظر گرفته می‌شوند:

$$\mathbf{B}_i = (\mathbf{Y}_{aug})_i \mathbf{X}_i^T = \lambda \mathbf{B}_{i-1} + w_i (\mathbf{Y}_{aug})_i \mathbf{X}_i^T \quad (16)$$

$$\mathbf{C}_i = \mathbf{X}_i \mathbf{X}_i^T = \lambda \mathbf{C}_{i-1} + w_i \mathbf{X}_i \mathbf{X}_i^T \quad (17)$$

در رابطه (۱۷) ماتریس $\mathbf{C}_i \in \mathbb{R}^{K \times K}$ ماتریسی متقارن است. با استفاده از لم معکوس‌سازی ماتریس به‌صورت زیر:

$$(\mathbf{E} + \mathbf{U} \mathbf{V})^{-1} = \mathbf{E}^{-1} - \mathbf{E}^{-1} \mathbf{U} (\mathbf{I}^{-1} + \mathbf{V} \mathbf{E}^{-1} \mathbf{U})^{-1} \mathbf{V} \mathbf{E}^{-1} \quad (18)$$

و با در نظر گرفتن تغییر متغیرهای زیر:

$$\mathbf{E} = \lambda \mathbf{C}_{i-1}, \quad \mathbf{U} = w_i \mathbf{X}_i, \quad \mathbf{I} = \mathbf{I}, \quad \mathbf{V} = \mathbf{X}_i^T \quad (19)$$

معکوس ماتریس $\mathbf{C}_i$ به‌صورت زیر محاسبه می‌شود:

$$\mathbf{C}_i^{-1} = \lambda^{-1} \mathbf{C}_{i-1}^{-1} - \frac{\lambda^{-1} \mathbf{C}_{i-1}^{-1} w_i \mathbf{X}_i \mathbf{X}_i^T \lambda^{-1} \mathbf{C}_{i-1}^{-1}}{1 + \mathbf{X}_i^T \lambda^{-1} \mathbf{C}_{i-1}^{-1} w_i \mathbf{X}_i} \quad (20)$$

با ضرب دو طرف رابطه (۲۰) در $\mathbf{B}_i$ از رابطه (۱۶)، رابطه بازگشتی برای $(\mathbf{D}_{aug})_i$ به‌صورت زیر به‌دست می‌آید:

$$\begin{aligned} (\mathbf{D}_{aug})_i &= \mathbf{B}_i \mathbf{C}_i^{-1} = \mathbf{B}_{i-1} \mathbf{C}_{i-1}^{-1} \\ &\quad - \frac{\lambda \mathbf{B}_{i-1} \lambda^{-1} \mathbf{C}_{i-1}^{-1} w_i \mathbf{X}_i \mathbf{X}_i^T \lambda^{-1} \mathbf{C}_{i-1}^{-1}}{1 + \mathbf{X}_i^T \lambda^{-1} \mathbf{C}_{i-1}^{-1} w_i \mathbf{X}_i} \\ &\quad + w_i (\mathbf{Y}_{aug})_i \mathbf{X}_i^T \lambda^{-1} \mathbf{C}_{i-1}^{-1} \\ &\quad - \frac{w_i (\mathbf{Y}_{aug})_i \mathbf{X}_i^T \lambda^{-1} \mathbf{C}_{i-1}^{-1} w_i \mathbf{X}_i \mathbf{X}_i^T \lambda^{-1} \mathbf{C}_{i-1}^{-1}}{1 + \mathbf{X}_i^T \lambda^{-1} \mathbf{C}_{i-1}^{-1} w_i \mathbf{X}_i} \end{aligned} \quad (21)$$

سپس با در نظر گرفتن تغییر متغیرهای زیر:

$$\mathbf{R}_i = (\mathbf{Y}_{aug})_i - \mathbf{B}_{i-1} \mathbf{C}_{i-1}^{-1} \mathbf{X}_i = (\mathbf{Y}_{aug})_i - (\mathbf{D}_{aug})_{i-1} \mathbf{X}_i \quad (22)$$

$$\mathbf{U}_i = \lambda^{-1} \mathbf{C}_{i-1}^{-1} \mathbf{X}_i \quad (23)$$

$$\delta_i = \frac{w_i}{1 + w_i \mathbf{X}_i^T \mathbf{U}_i} \quad (24)$$

رابطه بازگشتی به‌روزرسانی دیکشنری (۲۱) به‌صورت زیر ساده می‌شود:

$$(\mathbf{D}_{aug})_i = (\mathbf{D}_{aug})_{i-1} + \delta_i \mathbf{R}_i \mathbf{U}_i^T \quad (25)$$

به‌روزرسانی بازگشتی ماتریس $\mathbf{C}_i^{-1} \in \mathbb{R}^{K \times K}$ در رابطه (۲۳) برای تکرار بعدی نیز به‌صورت زیر محاسبه می‌شود:

$$\mathbf{C}_i^{-1} = \lambda^{-1} \mathbf{C}_{i-1}^{-1} - \delta_i \mathbf{U}_i \mathbf{U}_i^T \quad (26)$$

۳-۲- مرحله ارزیابی LC-GAWRLS

با در نظر گرفتن داده تست Y_j ابتدا نمایش تُنک آن با حل تابع هزینه زیر محاسبه می‌شود:

$$X_j = \arg \min_x \|Y_j - \hat{D}X\|_2^2 \text{ s.t. } \forall j, \|X_j\|_0 \leq S \quad (29)$$

در رابطه (۲۹) $\hat{D}$ دیکشنری تعلیم یافته به دست آمده در رابطه (۲۸) است. با استفاده از این دیکشنری، نمایش تُنک در تابع هزینه (۲۹) با الگوریتم OMP تخمین زده می‌شود. سپس این نمایش تُنک در ماتریس $\hat{G}$ به دست آمده از رابطه (۲۸) به صورت زیر ضرب می‌شود:

$$\tilde{H} = \hat{G}X_j \quad (30)$$

شاخص مربوط به بزرگترین درایه در بردار $\tilde{H} \in \mathbb{R}^{m \times 1}$ کلاس داده تست را مشخص می‌کند. طبقه‌بندی داده تست بدون برچسب با استفاده از روش LC-GAWRLS در الگوریتم ۲ خلاصه شده است.

الگوریتم ۲	
ورودی: $Y^{test} \in \mathbb{R}^{n \times L_{test}}$ (مجموعه داده‌های تست)، $\hat{G}$ و $\hat{D}$ و سطح تُنکی S	
شروع الگوریتم:	
(۱) برای نمونه $j = 1, \dots, L_{test}$ انجام دهید.	
(۲) نمونه جدیدی تست Y_j را دریافت کنید.	
(۳) نمایش تُنک X_j را حساب کنید:	
(۴) $X_j = OMP(Y_j, \hat{D}, S)$	
(۵) طبقه‌بندی کنید: $\tilde{H} = \hat{G}X_j$	
(۶) برچسب داده Y_j را مشخص کنید.	
(۷) پایان	

۳-۳- پیچیدگی محاسباتی

مهم‌ترین بخش پیچیدگی محاسباتی الگوریتم‌های یادگیری دیکشنری مربوط به مرحله نمایش تُنک و به‌روزرسانی دیکشنری می‌باشد که در هر مرحله به پارامترهایی مانند سایز دیکشنری K ، بعد داده تعلیمی m ، تعداد داده تعلیمی L ، تُنکی S وابسته است. در مقاله [۱۱]، پیچیدگی محاسباتی برای روش یادگیری دیکشنری بدون نظارت RLS برای یک تکرار کامل بر روی مجموعه آموزشی با L داده به‌صورت $O((S^2 + n + K)KL)$ محاسبه شده است. در روش یادگیری دیکشنری نظارت شده LC-RLS با اضافه شدن ماتریس الحاقی مانند رابطه (۹)، پیچیدگی محاسباتی به‌صورت $O((S^2 + (n + K + m) + K)KL)$ به دست می‌آید که در آن m تعداد کلاس‌ها است.

روش پیشنهادی LC-GAWRLS در مقایسه با روش LC-RLS فقط در گام (۷) الگوریتم ۱ با محاسبه وزن تصحیح برای هر داده آموزشی از رابطه (۲۷)، پیچیدگی $O(2(n + m + K) + 2)$ را به روش LC-RLS اضافه می‌کند که با ساده‌سازی به‌صورت $O(n + m + K)$ تبدیل می‌شود. در نتیجه برای یک تکرار الگوریتم بر روی مجموعه

آموزشی با L داده پیچیدگی به‌صورت $O((S^2 + (n + K + m) + K)KL + (n + K + m)L)$ برای روش LC-GAWRLS به دست می‌آید. با توجه به اینکه ضریب KL در جمله اول بسیار بزرگتر از ضریب L در جمله دوم است، می‌توان از جمله دوم صرف‌نظر کرد و بنابراین تاثیر پیچیدگی اضافه شده ناچیز بود و پیچیدگی $O((S^2 + (n + K + m) + K)KL)$ برای روش LC-GAWRLS به دست می‌آید. زمان لازم برای آموزش دو الگوریتم در جدول ۲ در بخش شبیه‌سازی باهم مقایسه شده است.

۴- شبیه‌سازی

۴-۱- داده‌ها، معیارهای ارزیابی و پارامترهای مورد استفاده

روش پیشنهادی با استفاده از سه مجموعه داده استاندارد شامل مجموعه داده ORL، AR، YaleB ارائه شده در شکل ۱، ارزیابی شده است [۲۸]. این مجموعه داده‌ها پیش‌تر به‌طور گسترده در ارزیابی روش‌های یادگیری دیکشنری مورد استفاده قرار گرفته‌اند و فرض می‌شود که فاقد نویز و دارای برچسب‌های صحیح و از پیش تعیین شده هستند [۲۹]. برای ارزیابی عملکرد الگوریتم‌ها، معیارهایی مانند دقت طبقه‌بندی، میانگین امتیاز F1 و حساسیت مورد استفاده قرار گرفته‌اند [۳۰]، [۳۱]. دقت طبقه‌بندی به‌صورت زیر محاسبه می‌شود:

$$CA = \frac{N_{Correct}}{L_{test}} \quad (31)$$

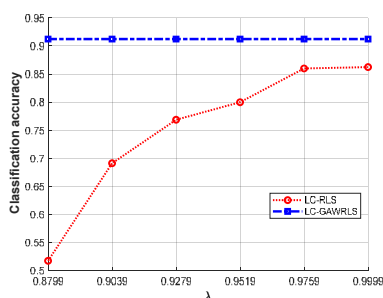
در رابطه بالا $N_{Correct}$ تعداد پیش‌بینی‌های صحیح و L_{test} تعداد کل نمونه‌های ورودی در مرحله ارزیابی الگوریتم‌ها می‌باشد. امتیاز F1 و حساسیت برای هر کلاس به‌ترتیب از رابطه (۳۲) و (۳۳) محاسبه شده است [۳۱]، [۳۲]:

$$F1 = \frac{2 \times TP}{2 \times TP + FP + FN} \quad (32)$$

$$SE = \frac{TP}{TP + FN} \quad (33)$$

که TP مثبت حقیقی، FP مثبت کاذب و FN منفی کاذب است. روش پیشنهادی با روش‌های نظارت شده LC-KSVD2 و LC-RLS مقایسه می‌شود.

برای ارزیابی پایداری روش پیشنهادی نسبت به تغییرات پارامترهای α و β ، آزمایش‌هایی روی سه مجموعه داده AR، YaleB و ORL انجام شده است. این بررسی با مقادیر مختلف داده‌های آموزشی و تنظیمات دیکشنری و تُنکی صورت گرفته است: در AR از ۲ نمونه در هر کلاس (سایز دیکشنری ۲۰۰، تُنکی ۳۸)، در YaleB از ۵ نمونه (سایز دیکشنری ۱۹۰، تُنکی ۳) و در ORL از ۸ نمونه (سایز دیکشنری ۸۰، تُنکی ۱۱) استفاده شده است. مطابق جدول ۱، دقت طبقه‌بندی روش پیشنهادی در برابر تغییرات α و β تغییر محسوسی نداشته و انحراف معیار آن در AR و YaleB برابر با ۰/۰۰۲ و در ORL برابر با



شکل ۲: تاثیر انتخاب مقادیر متفاوت λ بر دقت طبقه‌بندی در مجموعه داده ORL.

جدول ۲: مقایسه زمان آموزش روش‌ها بر روی مجموعه داده ORL.

الگوریتم	زمان آموزش
LC-RLS [۱۵]	۶/۵۷ ثانیه
LC-GAWRLS	۷/۵۱ ثانیه

۴-۲- نتایج شبیه‌سازی

در توضیحات شبیه‌سازی‌های انجام شده مقادیر بهینه برای سائز دیکشنری و مقدار λ انتخاب شده است. انتخاب این مقادیر پس از ارزیابی مقادیر مختلف سائز دیکشنری و مقدار λ با استفاده از دامنه مقادیر رایج در مقالات پیشین [۱۴]، [۱۵] به گونه‌ای انجام شده است که تعادل مناسبی میان دقت طبقه‌بندی و هزینه (سرعت) محاسباتی برقرار گردد.

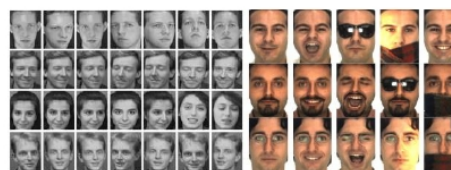
نتایج شبیه‌سازی بر روی مجموعه داده AR

مجموعه داده چهره AR شامل ۴۰۰۰ تصویر رنگی از ۱۲۶ فرد مختلف است. این تصاویر تغییرات مختلفی را در نورپردازی، حالت چهره و پوشش چهره (عینک آفتابی، شال گردن و ...) ثبت می‌کنند. زیر مجموعه‌ای شامل ۲۶۰۰ تصویر از ۵۰ مرد و ۵۰ زن (سائز کلاس ۱۰۰) انتخاب می‌شود. برای کاهش ابعاد تصویر و استخراج ویژگی‌های آن از یک ماتریس تصویر سازی که به‌طور تصادفی تولید شده است استفاده می‌شود [۳۵]. در این حالت، هر تصویر با سائز 120×165 پیکسل بر اساس وزن‌های تصادفی به برداری با بعد ۵۴۰ تبدیل می‌شود. باید در نظر داشت که این پیش‌پردازش قبلاً انجام شده است و ویژگی‌های استخراج شده حاصل از آن به‌عنوان داده ورودی برای شبیه‌سازی روش پیشنهادی استفاده می‌شود [۱۴].

به‌طور تصادفی ۵ تصویر از هر کلاس برای آموزش و بقیه برای تست انتخاب می‌شود. در مجموع ۵۰۰ داده برای آموزش و ۲۱۰۰ داده برای تست به کار می‌رود. سائز دیکشنری ۵۰۰ و سطح λ ۳۸ تنظیم شده است. نتایج جدول ۳ برای داده AR نشان می‌دهد که روش پیشنهادی LC-GAWRLS در مقایسه با روش‌های دیگر دقت طبقه‌بندی بالاتری دارد. در ادامه ارزیابی روش پیشنهادی، مدل با استفاده از نمونه‌های محدود شامل ۵، ۸، ۱۰، ۱۳، ۱۵ و ۲۰ نمونه در هر کلاس تعلیم داده می‌شود. شکل ۳ نشان می‌دهد که روش

صفر است. این نتایج نشان‌دهنده‌ی پایداری بالای الگوریتم پیشنهادی در برابر تغییرات پارامترهای α و β است. بر همین اساس، با استناد به نتایج ارزیابی ارائه شده در [۱۴]، [۱۵] مقادیر α و β به‌ترتیب برابر با ۴ و ۲ در نظر گرفته شده‌اند. همچنین تاثیر مقدار ضریب فراموشی λ بر دقت طبقه‌بندی دو روش LC-GAWRLS و LC-RLS برای مجموعه داده ORL بررسی شد و در آن مقادیر α و β به‌ترتیب برابر با ۴ و ۲ ثابت نگه داشته شدند. نتایج ارائه شده در شکل ۲ نشان می‌دهد که روش پیشنهادی LC-GAWRLS در مقایسه با LC-RLS پایداری بیشتری در برابر تغییرات λ دارد. بر این اساس، مقدار پارامتر λ در تمام آزمایش‌های بعدی برابر با ۰/۹۹۹۹ انتخاب شده است. همچنین برای مقایسه منصفانه، تمامی روش‌ها مطابق با روش ارائه شده در [۱۴] مقدار دهی اولیه می‌شوند.

در جدول ۲ زمان آموزش الگوریتم‌های LC-RLS و LC-GAWRLS بر روی مجموعه داده ORL با در نظر گرفتن ۶ داده آموزشی در هر کلاس، سائز دیکشنری ۸۰ و λ ۱۱ با یکدیگر مقایسه شده است. نتایج جدول ۲ نشان می‌دهد که زمان آموزش الگوریتم LC-GAWRLS با اضافه کردن وزن تصحیح در مقایسه با روش LC-RLS به‌طور قابل توجهی افزایش نمی‌یابد. نتیجه ارزیابی پیچیدگی محاسباتی و زمان آموزش روش LC-GAWRLS و نتایج دقت طبقه‌بندی در بخش نتایج نشان می‌دهد که روش پیشنهادی بدون اینکه پیچیدگی محاسباتی را به‌طور قابل توجهی افزایش دهد در مقایسه با روش LC-RLS دقت طبقه‌بندی را بهبود می‌دهد.



(الف) (ب)



(ب)

شکل ۱: تصاویر از (الف) مجموعه داده AR [۳۳] (ب) مجموعه داده چهره ORL [۳۴] (پ) مجموعه داده چهره YaleB [۳۳].

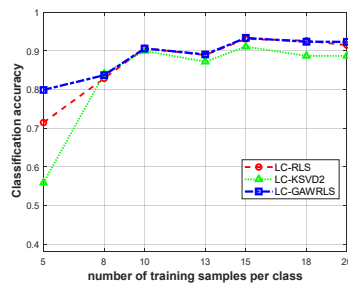
جدول ۱: تاثیر انتخاب مقادیر متفاوت α و β بر دقت طبقه‌بندی روش پیشنهادی در مجموعه داده AR، YaleB، و ORL.

انحراف	$\alpha=4$	$\alpha=4$	$\alpha=6$	$\alpha=4$	$\alpha=2$
معیار	$\beta=3$	$\beta=1$	$\beta=2$	$\beta=2$	$\beta=2$
AR	۰/۴۰۵۹	۰/۴۰۸۴	۰/۴۰۶۰	۰/۴۰۷۷	۰/۴۱۰۶
YaleB	۰/۷۴۶۶	۰/۷۴۳۷	۰/۷۴۱۳	۰/۷۴۲۴	۰/۷۴۳۲
ORL	۰/۹۳۰۰	۰/۹۳۰۰	۰/۹۳۰۰	۰/۹۳۰۰	۰/۹۳۰۰

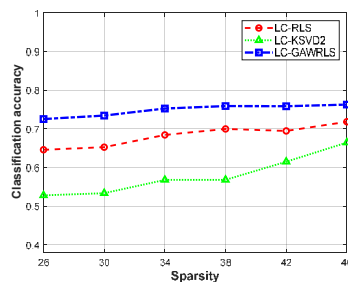
الگوریتم‌ها ۵۷۰ تنظیم شده است. مقدار تُنکی نیز برابر با ۳۰ تنظیم شده است. نتایج جدول ۳ برای داده YaleB نشان می‌دهد که الگوریتم LC-GAWRLS با دست یافتن به دقت طبقه‌بندی ۹۵/۵ درصد از دو الگوریتم دیگر بهتر عمل می‌کند.

جدول ۳: مقایسه درصدهای دقت طبقه‌بندی بین LC-GAWRLS و الگوریتم‌های موجود برای مجموعه داده‌های مختلف.

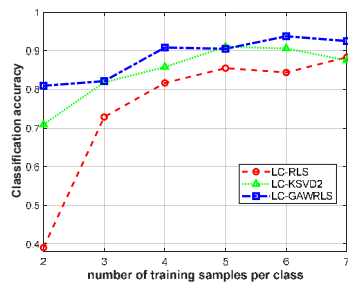
الگوریتم	AR	ORL	YaleB
LC-KSVD2 [۱۴]	۵۵/۸	۸۷/۵	۷۵/۵
LC-RLS [۱۵]	۷۱/۴	۸۸/۳	۸۶/۰
LC-GAWRLS	۷۹/۹	۹۲/۵	۹۵/۵



شکل ۳: دقت طبقه‌بندی در برابر تعداد نمونه‌های آموزشی در هر کلاس برای مجموعه داده AR.



شکل ۴: دقت طبقه‌بندی در برابر مقادیر مختلف تُنکی (S).



شکل ۵: دقت طبقه‌بندی در برابر تعداد نمونه‌های آموزشی در هر کلاس برای مجموعه داده ORL.

در شکل ۶ تاثیر تعداد داده‌های آموزشی بر دقت طبقه‌بندی بررسی شده است. همان‌طور که انتظار می‌رود با افزایش تعداد داده‌های آموزشی در هر کلاس، دقت طبقه‌بندی برای تمام روش‌های در نظر گرفته شده بهبود می‌یابد. نتایج شکل ۶ نشان می‌دهد که برای تعداد محدودی داده آموزشی در هر کلاس، روش پیشنهادی LC-GAWRLS

LC-GAWRLS به‌ویژه با تعداد داده آموزشی محدود، عملکرد و دقت طبقه‌بندی بالاتری نسبت به دو روش دیگر دارد.

در ادامه بررسی پارامترهای مورد استفاده، تاثیر تغییر مقدار تُنکی S بر دقت طبقه‌بندی برای مجموعه داده AR مورد بررسی قرار گرفته است. در این بررسی از ۵ نمونه در هر کلاس برای آموزش و از باقی مانده نمونه‌ها برای ارزیابی استفاده شده است. اندازه دیکشنری برابر با ۵۰۰ و برای تُنکی مقادیر مختلف در نظر گرفته شده است. همان‌طور که نتایج شکل ۴ نشان می‌دهد با افزایش مقدار تُنکی، دقت طبقه‌بندی در تمامی روش‌ها بهبود یافته است. در این میان، الگوریتم پیشنهادی نه تنها دقت طبقه‌بندی بالاتری نسبت به سایر روش‌ها در تمامی سطوح تُنکی به‌دست آورده، بلکه از پایداری بیشتری نیز در برابر تغییرات مقدار تُنکی برخوردار بوده است.

نتایج شبیه‌سازی بر روی مجموعه داده ORL

مجموعه داده ORL شامل ۴۰۰ تصویر از ۴۰ فرد مختلف است. هر فرد ۱۰ تصویر دارد و تصاویر تحت شرایط مختلف نورپردازی همراه با تغییرات مختلف در حالت چهره (لبخند زدن یا بستن چشم) و پوشش چهره (عینک زدن) گرفته شده‌اند. هر تصویر دارای سایز ۳۲×۳۲ پیکسل می‌باشد که این پیکسل‌ها به صورت برداری با بعد ۱۰۲۴ روی هم قرار می‌گیرند. در این مجموعه داده ۷ تصویر از هر کلاس برای آموزش و بقیه برای تست استفاده می‌شوند. در مجموع ۲۸۰ تصویر برای آموزش و ۱۲۰ تصویر برای تست به‌کار می‌رود. سایز دیکشنری ۸۰ و سطح تُنکی ۱۱ تنظیم شده است. نتایج جدول ۳ برای داده ORL نشان می‌دهد که روش پیشنهادی LC-GAWRLS در مقایسه با روش‌های دیگر به دقت طبقه‌بندی بالاتری دست می‌یابد.

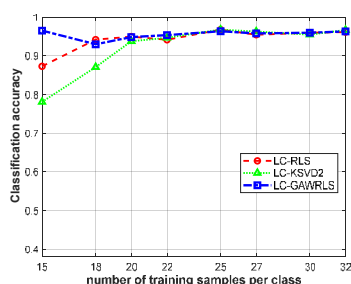
در ادامه ارزیابی، تاثیر تعداد داده‌های آموزشی بر دقت طبقه‌بندی بررسی شده است. همان‌طور که در شکل ۵ نشان داده شده است، الگوریتم LC-GAWRLS به‌ویژه با در نظر گرفتن داده‌های آموزشی محدود در هر کلاس، دقت طبقه‌بندی بالاتری دارد که برتری آن را در کاربردهایی با تعداد داده آموزشی محدود برجسته می‌کند.

نتایج شبیه‌سازی بر روی مجموعه داده توسعه یافته YaleB

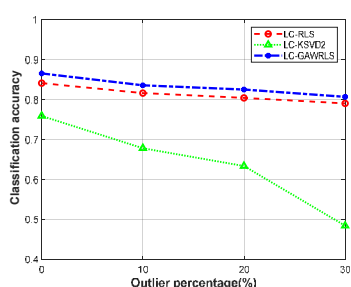
مجموعه داده توسعه یافته YaleB شامل ۲۴۱۴ تصویر چهره از ۳۸ فرد مختلف می‌باشد. این تصاویر تحت شرایط نوری متفاوت گرفته شده‌اند. هر فرد به‌طور میانگین ۶۴ تصویر دارد. این تصاویر در ابعاد ۱۶۸×۱۹۲ پیکسل برش داده و نرمالیزه شده‌اند. مانند مجموعه داده AR با به‌کار بردن ماتریس تصویرسازی تصادفی، هر تصویر با سایز ۱۶۸×۱۹۲ پیکسل بر روی برداری با بعد ۵۰۴ تصویر می‌شود. لازم به ذکر است که این پیش‌پردازش بر روی مجموعه داده YaleB قبلاً انجام شده است و نتایج آن به‌عنوان داده ورودی برای شبیه‌سازی روش پیشنهادی مورد استفاده قرار می‌گیرد [۱۴].

به‌طور تصادفی ۱۵ چهره از هر کلاس برای آموزش و بقیه برای تست الگوریتم‌ها انتخاب شده است. در مجموع ۵۷۰ داده برای آموزش و ۱۸۶۲ داده برای تست استفاده شده است. سایز دیکشنری در تمام

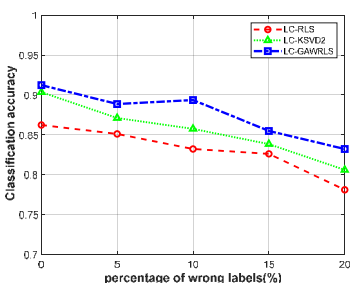
ORL نیز ارزیابی شده است. نتایج این ارزیابی در جدول ۴ ارائه شده و نشان می‌دهد که الگوریتم پیشنهادی در هر دو معیار یاد شده عملکرد بهتری نسبت به سایر روش‌ها داشته و امتیاز بالاتری کسب می‌کند.



شکل ۶: دقت طبقه‌بندی در برابر تعداد نمونه‌های آموزشی در هر کلاس برای مجموعه داده توسعه یافته YaleB.



شکل ۷: دقت طبقه‌بندی در مقابل درصد داده‌های پرت از داده‌های آموزشی هر کلاس در مجموعه داده توسعه یافته YaleB.



شکل ۸: دقت طبقه‌بندی الگوریتم‌ها در برابر درصد داده‌های دارای برچسب اشتباه در مجموعه داده ORL.

جدول ۴: مقایسه میانگین امتیاز F1 و حساسیت روی مجموعه داده

ORL		
الگوریتم	میانگین امتیاز F1	میانگین حساسیت
[۱۴] LC-KSVD2	۰/۹۰۱۰	۰/۹۰۶۲
[۱۵] LC-RLS	۰/۸۲۲۳	۰/۸۴۳۷
LC-GAWRLS	۰/۹۳۰۹	۰/۹۳۷۵

عملکرد الگوریتم پیشنهادی با دو روش یادگیری دیکشنری عمیق، شامل روش بدون نظارت HILADLE [۱۷] و روش نظارت شده MDDL [۱۹] بر روی مجموعه داده Extended YaleB [۲۹] نیز مورد ارزیابی قرار گرفته است. تمامی تصاویر این مجموعه داده به اندازه ۳۲×۳۲

در مقایسه با روش‌های دیگر دقت طبقه‌بندی بسیار بالاتری را به دست می‌آورد. این نتیجه، مزیت روش پیشنهادی را به‌ویژه در سناریو و کاربردهایی با تعداد داده محدود برجسته می‌کند.

برای ارزیابی میزان پایداری الگوریتم‌ها در برابر داده‌های پرت، نویز مصنوعی به داده‌های آموزشی مجموعه داده YaleB افزوده و تأثیر آن بر عملکرد طبقه‌بندی مورد بررسی قرار گرفته است. در این آزمایش، نویز گاوسی با میانگین صفر و واریانس ۸۰۰ به‌صورت تصادفی به درصدی از داده‌های آموزشی هر کلاس (در بازه ۰ تا ۳۰ درصد) اضافه شده است، به‌گونه‌ای که نسبت سیگنال به نویز (SNR) به‌طور متوسط برابر با ۱۲ دسی‌بل بوده است. نمونه‌های آلوده‌شده به نویز به‌عنوان داده‌های پرت یا نویزی در نظر گرفته شده‌اند. در این آزمایش، از ۱۰ نمونه در هر کلاس برای آموزش و از سایر نمونه‌ها برای ارزیابی عملکرد الگوریتم‌ها استفاده شده است. اندازه دیکشنری برابر با ۳۸۰ و سطح تنگی روی ۳۰ تنظیم شده است. نتایج ارائه‌شده در شکل ۷ نشان می‌دهد که با افزایش درصد نویز، دقت طبقه‌بندی در هر سه الگوریتم کاهش می‌یابد. با این حال الگوریتم پیشنهادی LC-GAWRLS در تمامی سطوح مختلف داده‌های پرت، عملکرد طبقه‌بندی بهتری نسبت به سایر روش‌ها از خود نشان می‌دهد. این بهبود ناشی از مکانیزم وزن‌دهی تطبیقی است که در ساختار الگوریتم لحاظ شده و نقش مؤثری در کاهش اثر نمونه‌های نویزی در فرآیند به‌روزرسانی مدل ایفا می‌کند. در نتیجه، مدل نهایی از پایداری بالاتر و دقت بیشتری در مواجهه با داده‌های پرت برخوردار است.

در ارزیابی بعدی تأثیر وجود برچسب‌های اشتباه در داده‌های آموزشی بر دقت طبقه‌بندی الگوریتم‌ها و میزان پایداری آن‌ها مورد بررسی قرار گرفته است. بدین منظور در مجموعه داده ORL، تعدادی از برچسب‌های آموزشی به‌صورت تصادفی در ماتریس برچسب H و ماتریس سازگاری برچسب Q به‌طور نادرست اختصاص داده شده‌اند. در این آزمایش بین ۰ تا ۲۰ درصد از کل داده‌های آموزشی به‌صورت تصادفی انتخاب و برچسب آن‌ها نادرست لحاظ شده است. برای آموزش از ۶ تصویر در هر کلاس و برای ارزیابی از باقی‌مانده تصاویر استفاده شده است. همچنین، اندازه دیکشنری برابر با ۸۰ و مقدار تنگی برابر با ۱۱ تنظیم شده است. نتایج ارائه شده در شکل ۸ نشان می‌دهد که با افزایش درصد برچسب‌های نادرست در داده‌های آموزشی، دقت طبقه‌بندی در تمامی روش‌ها کاهش می‌یابد. با این حال، الگوریتم پیشنهادی LC-GAWRLS در تمام سطوح نویز برچسب، عملکرد بهتری نسبت به سایر روش‌ها از خود نشان می‌دهد. این برتری ناشی از به‌کارگیری مکانیزم وزن‌دهی تطبیقی در الگوریتم پیشنهادی است که با تنظیم سهم هر نمونه در فرآیند یادگیری، تأثیر منفی برچسب‌های اشتباه را کاهش می‌دهد. در نتیجه، الگوریتم از پایداری و دقت بالاتری در مواجهه با نویز برچسب برخوردار است.

علاوه بر دقت طبقه‌بندی، عملکرد الگوریتم‌ها بر اساس معیارهای میانگین امتیاز F1 و حساسیت [۳۱] برای ۴۰ کلاس از مجموعه داده

۵- نتیجه‌گیری

در این مقاله، یک روش یادگیری دیکشنری آنلاین نظارت شده با عنوان LC-GAWRLS معرفی شد. در تابع هزینه این روش، علاوه بر خطای بازسازی تُنک، خطای طبقه‌بندی نیز لحاظ شده و سه مؤلفه‌ی کلیدی شامل دیکشنری، ماتریس تبدیل خطی و طبقه‌بندی کننده، به‌طور بازگشتی و هم‌زمان با استفاده از الگوریتم GAW-RLS به‌روزرسانی می‌شوند. این ساختار با بهره‌گیری از یادگیری آنلاین، داده‌های آموزشی را به‌صورت تک‌نمونه‌ای و ترتیبی پردازش می‌کند، بنابراین در مقایسه با روش‌های پیشین، حتی در مواجهه با مجموعه داده‌های بزرگ، نیاز به حافظه بالا ندارد و از محدودیت‌های ذخیره‌سازی رهایی می‌یابد.

از جمله ویژگی‌های مهم این روش، بهره‌گیری از مکانیزم وزن‌دهی تطبیقی است که با تخصیص وزن مناسب به هر نمونه آموزشی، اثر داده‌های پرت و برچسب‌گذاری نادرست را کاهش می‌دهد. همین عامل باعث شده تا الگوریتم در مواجهه با نویز در داده‌ها یا برچسب‌ها، پایداری و دقت بالاتری نسبت به سایر روش‌ها داشته باشد.

نتایج شبیه‌سازی روی مجموعه داده‌های مختلف نشان داد که روش پیشنهادی در سناریوهای با نمونه‌های آموزشی محدود، از نظر معیارهای دقت، حساسیت و امتیاز F1، عملکرد بهتری نسبت به روش‌های پیشین و قابل مقایسه با برخی روش‌های یادگیری عمیق از خود نشان داده است. در حالی که بسیاری از روش‌های مبتنی بر یادگیری عمیق به داده‌های آموزشی حجیم، منابع محاسباتی بالا و زمان آموزش زیاد وابسته‌اند، روش پیشنهادی LC-GAWRLS با ساختار ساده و به‌روزرسانی بازگشتی، هزینه محاسباتی کمتر و دقت رقابتی را فراهم می‌آورد. در مجموع، روش LC-GAWRLS با ترکیب مزایای یادگیری تُنک، به‌روزرسانی بازگشتی، وزن‌دهی تطبیقی و عملکرد بالا در شرایط نویزی، به‌عنوان جایگزینی مؤثر و سبک برای بسیاری از روش‌های سنگین یادگیری عمیق در مسائل طبقه‌بندی با داده‌های ناکامل یا دارای برچسب ناقص، قابل استفاده است.

مراجع

- [1] J. Wright, A. Y. Yang, A. Ganesh, S. S. Sastry, and Yi Ma, "Robust Face Recognition via Sparse Representation," *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 31, no. 2, pp. 210–227, Feb. 2009, doi: 10.1109/TPAMI.2008.79.
- [2] M. Yang, L. Zhang, J. Yang, and D. Zhang, "Robust sparse coding for face recognition," in *Proceedings of the IEEE Computer Society Conference on Computer Vision and Pattern Recognition*, Jun. 2011, no. 1, pp. 625–632. doi: 10.1109/CVPR.2011.5995393.
- [3] Jianchao Yang, Kai Yu, Yihong Gong, and T. Huang, "Linear spatial pyramid matching using sparse coding for image classification," in *2009 IEEE Conference on Computer Vision and Pattern Recognition*, Jun. 2009, pp. 1794–1801. doi: 10.1109/CVPR.2009.5206757.
- [4] H. Li and F. Liu, "Image Denoising Via Sparse and Redundant Representations Over Learned Dictionaries in Wavelet Domain," in *2009 Fifth International Conference on Image and Graphics*, Sep. 2009, vol. 15, no. 12, pp. 754–758. doi: 10.1109/ICIG.2009.101.
- [5] Y. Naderahmadian and S. Beheshti, "Generalized Adaptive Weighted Recursive Least Squares Dictionary Learning for Retinal Vessel Inpainting," in *2018 IEEE Statistical Signal*

پیکسل نرمال‌سازی شده‌اند. لازم به ذکر است که در روش‌های LC-RLS، LC-KSVD2، MDDL و همچنین روش پیشنهادی LC-GAWRLS هیچ‌گونه استخراج ویژگی بر روی تصاویر اعمال نشده و تصاویر ورودی مستقیماً در فرآیند آموزش مورد استفاده قرار گرفته‌اند. در مقابل، روش HILADLE از پیش‌پردازش PCA^{۱۸} برای کاهش ابعاد و استخراج ویژگی استفاده می‌کند [۱۷]. مطابق تنظیمات ارائه شده در [۱۹]، اندازه دیکشنری در لایه اول برای روش MDDL و همچنین سایر روش‌های تک‌لایه، برابر با ۴۵۶ در نظر گرفته شده و از ۳۲ تصویر در هر کلاس برای آموزش استفاده شده است. در حالی که در روش HILADLE اندازه دیکشنری در لایه اول برابر با ۸۰۰ انتخاب شده و از ۴۴ تصویر در هر کلاس برای آموزش استفاده شده است [۱۷]. سطح تُنکی برای تمامی روش‌های تک‌لایه برابر با ۳۰ در نظر گرفته شده و پارامترهای α و β نیز به ترتیب برابر با ۴ و ۲ تنظیم شده‌اند.

برای روش MDDL از نتایج گزارش شده در [۱۹] با دو و سه لایه و برای HILADLE از نتایج سه لایه گزارش شده در [۱۷] استفاده شده است. نتایج ارائه شده در جدول ۵ نشان می‌دهد که روش‌های یادگیری دیکشنری عمیق MDDL و HILADLE به دقت طبقه‌بندی بالاتری در حدود ۳ درصد نسبت به روش پیشنهادی LC-GAWRLS دست یافته‌اند. باید توجه داشت که این افزایش دقت عمدتاً ناشی از استفاده از ساختارهای چندلایه و استخراج ویژگی‌های سطح بالا در این روش‌ها است. با این حال، بهره‌گیری از لایه‌های بیشتر منجر به افزایش قابل توجه در پیچیدگی محاسباتی آن‌ها می‌شود. در مقابل، روش پیشنهادی LC-GAWRLS با دقت طبقه‌بندی ۹۵/۹۸ درصد عملکردی کاملاً قابل رقابت با روش‌های عمیق ارائه می‌دهد، در حالی که به دلیل ساختار تک‌لایه از پیچیدگی محاسباتی بسیار کمتری برخوردار است. این ویژگی، الگوریتم پیشنهادی را به گزینه‌ای مناسب و عملی برای کاربردهای واقعی با منابع محاسباتی محدود تبدیل می‌کند.

جدول ۵: مقایسه دقت درصد طبقه‌بندی روش پیشنهادی با روش‌های یادگیری دیکشنری تک‌لایه و چندلایه بر روی مجموعه داده توسعه یافته YaleB.

تقسیم‌بندی	الگوریتم	دقت طبقه‌بندی
روش‌های یادگیری دیکشنری چند لایه	MDDL-2 [۱۹]	۹۸/۲۰
	MDDL-3 [۱۹]	۹۸/۳۰
	HILADLE-3 [۱۷]	۹۷/۴۱
روش‌های یادگیری دیکشنری تک لایه	LC-KSVD2 [۱۴]	۹۵/۳۳
	LC-RLS [۱۵]	۹۴/۸۲
	LC-GAWRLS	۹۵/۹۸

- IEEE Trans. Veh. Technol.*, pp. 1–11, 2025, doi: 10.1109/TVT.2025.3552082.
- [22] D. Zhao, P. Zhang, H. Yin, and J. Guo, “A novel multi-layer discriminative dictionary learning approach for image classification,” *Signal Processing*, vol. 226, p. 109670, Jan. 2025, doi: 10.1016/j.sigpro.2024.109670.
- [23] M. Yousefi, Y. Naderahmadian, and S. Beheshti, “Label Consistent Generalized Adaptive Weighted Recursive Least Squares Dictionary Learning,” in *2025 IEEE 6th International Conference on Image Processing, Applications and Systems (IPAS)*, Jan. 2025, pp. 1–6. doi: 10.1109/IPAS63548.2025.10924517.
- [24] Y. Naderahmadian, S. Beheshti, and M. A. Tinati, “Correlation based online dictionary learning algorithm,” *IEEE Trans. Signal Process.*, vol. 64, no. 3, pp. 592–602, 2016, doi: 10.1109/TSP.2015.2486743.
- [25] P. Henderson and I. M. Dale, “The partitioning of selected transition element ions between olivine and groundmass of oceanic basalts,” *Chem. Geol.*, vol. 5, no. 4, pp. 267–274, 1970, doi: 10.1016/0009-2541(70)90044-6.
- [26] M. Elad, *Sparse and redundant representations: From theory to applications in signal and image processing*, vol. 01. New York, NY: Springer New York, 2010. doi: 10.1007/978-1-4419-7011-4.
- [27] J. Mairal, F. Bach, and J. Ponce, “Task-driven dictionary learning,” *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 34, no. 4, pp. 791–804, 2012, doi: 10.1109/TPAMI.2011.156.
- [28] A. Golts and M. Elad, “Linearized Kernel Dictionary Learning,” *IEEE J. Sel. Top. Signal Process.*, vol. 10, no. 4, pp. 726–739, 2016, doi: 10.1109/JSTSP.2016.2555241.
- [29] “Popular Face Data Sets in Matlab Format.” <http://www.cad.zju.edu.cn/home/dengcai/Data/FaceData.html> (accessed Feb. 12, 2025).
- [30] Z. Zhang *et al.*, “Jointly learning structured analysis discriminative dictionary and analysis multiclass classifier,” *IEEE Trans. Neural Networks Learn. Syst.*, vol. 29, no. 8, pp. 3798–3814, 2018, doi: 10.1109/TNNLS.2017.2740224.
- [31] Y. Yang, F. Shao, Z. Fu, and R. Fu, “Discriminative dictionary learning for retinal vessel segmentation using fusion of multiple features,” *Signal, Image Video Process.*, vol. 13, no. 8, pp. 1529–1537, 2019, doi: 10.1007/s11760-019-01501-9.
- [32] V. Singhal, J. Maggu, and A. Majumdar, “Simultaneous Detection of Multiple Appliances from Smart-Meter Measurements via Multi-Label Consistent Deep Dictionary Learning and Deep Transform Learning,” *IEEE Trans. Smart Grid*, vol. 10, no. 3, pp. 2969–2978, 2019, doi: 10.1109/TSG.2018.2815763.
- [33] V. P. Vishwakarma and S. Dalal, “A novel non-linear modifier for adaptive illumination normalization for robust face recognition,” *Multimed. Tools Appl.*, vol. 79, no. 17–18, pp. 11503–11529, May 2020, doi: 10.1007/s11042-019-08537-6.
- [34] P. Ji, T. Zhang, H. Li, M. Salzmann, and I. Reid, “Deep Subspace Clustering Networks,” *Adv. Neural Inf. Process. Syst.*, vol. 2017-Decem, no. Nips, pp. 24–33, Sep. 2017, [Online]. Available: <https://doi.org/10.48550/arXiv.1709.02508>
- [35] Q. Zhang and B. Li, “Discriminative K-SVD for dictionary learning in face recognition,” in *Proceedings of the IEEE Computer Society Conference on Computer Vision and Pattern Recognition*, Jun. 2010, pp. 2691–2698. doi: 10.1109/CVPR.2010.5539989.
- [6] S. G. Mallat and Z. Zhang, “Matching Pursuits With Time-Frequency Dictionaries,” *IEEE Trans. Signal Process.*, vol. 41, no. 12, pp. 3397–3415, 1993, doi: 10.1109/78.258082.
- [7] G. Du, Y. Wang, L. Chen, and G. Wang, “Multiple Magnetic Targets Localization Using Multi-scale Cluster Matching Pursuit Method,” *IEEE Trans. Geosci. Remote Sens.*, vol. 62, pp. 1–16, 2024, doi: 10.1109/TGRS.2024.3478757.
- [8] J. M. Bioucas-Dias and M. A. T. Figueiredo, “A new TwIST: Two-step iterative shrinkage/thresholding algorithms for image restoration,” *IEEE Trans. Image Process.*, vol. 16, no. 12, pp. 2992–3004, 2007, doi: 10.1109/TIP.2007.909319.
- [9] A. Beck and M. Teboulle, “A fast iterative shrinkage-thresholding algorithm with application to wavelet-based image deblurring,” *ICASSP, IEEE Int. Conf. Acoust. Speech Signal Process. - Proc.*, no. 5, pp. 693–696, 2009, doi: 10.1109/icassp.2009.4959678.
- [10] M. Aharon, M. Elad, and A. Bruckstein, “K-SVD: An algorithm for designing overcomplete dictionaries for sparse representation,” *IEEE Trans. Signal Process.*, vol. 54, no. 11, pp. 4311–4322, 2006, doi: 10.1109/TSP.2006.881199.
- [11] K. Skretting and K. Engan, “Recursive least squares dictionary learning algorithm,” *IEEE Trans. Signal Process.*, vol. 58, no. 4, pp. 2121–2130, Apr. 2010, doi: 10.1109/TSP.2010.2040671.
- [12] Y. Naderahmadian, M. A. Tinati, and S. Beheshti, “Generalized adaptive weighted recursive least squares dictionary learning,” *Signal Processing*, vol. 118, pp. 89–96, Jan. 2016, doi: 10.1016/j.sigpro.2015.06.013.
- [13] J. Mairal, F. Bach, J. Ponce, and G. Sapiro, “Online dictionary learning for sparse coding,” in *Proceedings of the 26th Annual International Conference on Machine Learning*, Jun. 2009, vol. 11, pp. 689–696. doi: 10.1145/1553374.1553463.
- [14] Z. Jiang, Z. Lin, and L. S. Davis, “Label consistent K-SVD: Learning a discriminative dictionary for recognition,” *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 35, no. 11, pp. 2651–2664, 2013, doi: 10.1109/TPAMI.2013.88.
- [15] S. Matiz and K. E. Barner, “Label consistent recursive least squares dictionary learning for image classification,” in *Proceedings - International Conference on Image Processing, ICIP*, Sep. 2016, vol. 2016-Augus, pp. 1888–1892. doi: 10.1109/ICIP.2016.7532686.
- [16] S. Mohseni-Sehdeh and M. Babaie-Zadeh, “A fast dictionary-learning-based classification scheme using undercomplete dictionaries,” *Signal Processing*, vol. 212, p. 109124, 2023, doi: 10.1016/j.sigpro.2023.109124.
- [17] J. Gou, X. He, L. Du, B. Yu, W. Chen, and Z. Yi, “Hierarchical Locality-Aware Deep Dictionary Learning for Classification,” *IEEE Trans. Multimed.*, vol. 26, no. Ddl, pp. 447–461, 2024, doi: 10.1109/TMM.2023.3266610.
- [18] S. Tariyal, A. Majumdar, R. Singh, and M. Vatsa, “Deep Dictionary Learning,” *IEEE Access*, vol. 4, pp. 10096–10109, 2016, doi: 10.1109/ACCESS.2016.2611583.
- [19] J. Song, X. Xie, G. Shi, and W. Dong, “Multi-layer discriminative dictionary learning with locality constraint for image classification,” *Pattern Recognit.*, vol. 91, pp. 135–146, 2019, doi: 10.1016/j.patcog.2019.02.018.
- [20] J. Gou, X. Yuan, B. Yu, J. Yu, and Z. Yi, “Intra- and Inter-Class Induced Discriminative Deep Dictionary Learning for Visual Recognition,” *IEEE Trans. Multimed.*, vol. 25, pp. 1575–1583, 2023, doi: 10.1109/TMM.2023.3258141.
- [21] H. Huang, C. Wang, L. Zhao, W. Wang, S. Ding, and A. Vasilakos, “Wi-Fi Sensing Based on Deep Supervised Dictionary Learning for Robust Device-Free Localization,” *Processing Workshop, SSP 2018*, Jun. 2018, no. 1, pp. 26–29. doi: 10.1109/SSP.2018.8450765.

زیرنویس‌ها

[∗] Generalized adaptive weighted RLS

[^] Label Consistent KSVD2

[^] Label Consistent RLS

[^] Undercomplete Dictionary Learning Algorithm

[^] Hierarchical Locality-aware Deep Dictionary Learning

[^] Deep Dictionary Learning

[^] Orthogonal Matching Pursuit

[^] Iterative soft-thresholding algorithm

[^] Fast Iterative soft-thresholding algorithm

[^] K-means Singular Value Decomposition

[^] Method of Optimal Directions

[^] Recursive Least Square method

^{۱۳} Multi-Layer Discriminative Dictionary Learning

^{۱۴} Discriminative Deep Dictionary Learning

^{۱۵} Deep-Supervised Dictionary Learning

^{۱۶} Multi-layer local constraint and label embedding dictionary learning

^{۱۷} Least Absolute Shrinkage and Selection Operator

^{۱۸} Principal Component Analysis