

مدل ترکیبی تحلیل مؤلفه اصلی احتمالاتی با نظارت

در چارچوب کاهش بُعد بدون اتلاف برای

شناسایی چهره

سمیه احمدخانی و پیمان ادیبی

گروه هوش مصنوعی، دانشکده مهندسی کامپیوتر، دانشگاه اصفهان، اصفهان، ایران

چکیده

در این مقاله ابتدا مدل بانظارت روش ترکیبی تحلیل مؤلفه اصلی احتمالاتی^۱ (SPPCAMM) ارائه شده است؛ سپس با در نظر گرفتن جریمه نگاشت در یادگیری مدل پیش‌گو، روشی برای شناسایی چهره با استفاده از یک رویکرد کاهش بُعد بدون اتلاف ارائه شده است. در روش پیشنهادی ابتدا یک خمینه زیربنایی محلی خطی با استفاده از مدل ترکیبی تحلیل مؤلفه اصلی احتمالاتی بانظارت از نمونه داده‌ها به دست می‌آید؛ سپس دسته‌بند ماشین بردار پشتیبان با اعمال جریمه نگاشت به عنوان مدل پیش‌گوی مذکور با استفاده از این خمینه محلی خطی آموزش داده می‌شود. بدین ترتیب از مزایای کاهش بُعد در مدل پیش‌گو استفاده می‌شود؛ و در عین حال جلوی از دست رفتن اطلاعات مفید گرفته می‌شود. برای آموزش و ارزیابی روش پیشنهادی، از پایگاه داده‌های چهره شناخته شده استفاده شده است. روش استخراج ویژگی گابور بر روی تصاویر به کار گرفته شده است. نتایج آزمایش‌ها نشان می‌دهد که روش پیشنهادی نسبت به بسیاری از روش‌های معمول که کاهش بُعد را انجام داده و سپس دسته‌بند را آموزش می‌دهند، و همچنین نسبت به روش جریمه نگاشت مبتنی بر مدل‌های کاهش بُعد خطی و غیرخطی، دقت بیشتری دارد.

واژگان کلیدی: کاهش بُعد بدون اتلاف، مدل ترکیبی، تحلیل مؤلفه اصلی احتمالی، بانظارت، جریمه نگاشت.

۱- مقدمه

یادگیری یک مدل دسته‌بندی در فضایی با بُعد بالا و با وجود محدودیت در تعداد نمونه‌های آموزشی، کار دشواری است. از این رو روش‌های کاهش بُعد فضای ویژگی به منظور به‌سازی یادگیری در این مدل‌ها مورد توجه پژوهش‌گران قرار گرفته‌اند. روش‌های کاهش بُعد علاوه بر کوچک کردن بردار ویژگی، به نوعی باعث افزایش دقت دسته‌بندی نیز می‌شوند. از آنجا که بُعد فضای ورودی در بسیاری از مسائل بیش از حد لازم، بزرگ است، استفاده از روش‌های کاهش بُعد علاوه بر کم کردن طول بردار ویژگی و کاهش هزینه محاسباتی مراحل بعدی، بالابردن کارایی از طریق مقابله با معضل بُعد^۱ را نیز در پی خواهد داشت. به همین دلیل این روش‌ها در عمل دارای اهمیت بوده و به‌طور وسیع مورد مطالعه قرار گرفته‌اند (بیشاپ، ۲۰۰۶).

یکی از روش‌های کاهش بُعد محبوب روش تحلیل مؤلفه اصلی^۲ (PCA) است، که تجزیه مقادیر منفرد^۳ (SVD) را روی ماتریس داده اجرا کرده و زیرفضایی با استفاده از مقادیر ویژه بزرگ ایجاد می‌کند (ون در ماتن و همکاران، ۲۰۰۹). بسیاری از روش‌های کاهش بُعد بدون نظارت هستند؛ به‌طوری که تنها روی داده‌های بدون برچسب مشاهده شده ورودی تمرکز می‌کنند؛ اما در مواقعی که مقادیر برچسب‌ها یا خروجی در دسترس است، وارد کردن این اطلاعات در نگاشت و رسیدن به یک افکنش بانظارت برای داده ورودی در کاربردهایی مانند دسته‌بندی می‌تواند مفیدتر باشد. بنابراین روش‌های کاهش بُعد بانظارت به‌منظور به کارگیری اطلاعات خروجی ارائه شده‌اند (یو و همکاران، ۲۰۰۶).

² Principal Component Analysis

³ Singular Value Decomposition (SVD)

¹ Curse of Dimensionality

روش تحلیل تمایز خطی^۱ یا LDA روشی بانظارت است که به کمک یک بهینه‌سازی ساده بردار افکنش را طوری به دست می‌آورد که فاصله درون‌دسته‌ای نمونه‌های آموزشی کمینه و فاصله بین دسته‌ای آن‌ها بیشینه شود (فیشر، ۱۹۳۸). روش تحلیل تمایز خطی وقتی به خوبی کار می‌کند که نمونه داده‌ها جدایی‌پذیر خطی باشند. برای غلبه بر این محدودیت روش تحلیل تمایز هسته^۲ (KDA) ارائه شده است (میکا و همکاران، ۱۹۹۹). فرض دیگری در روش تحلیل تمایز نخستین وجود دارد که توزیع احتمال هر دسته را محدود به توزیع گاسی تک‌قله‌ای می‌کند (فیشر، ۱۹۳۸). روش تحلیل فیشر حاشیه‌ای^۳ (MFA) با استفاده از ساختار جانشانی گراف^۴ (یان و همکاران، ۲۰۰۷)، سعی در غلبه بر محدودیت‌های روش LDA دارد. در مقایسه با LDA، روش MFA مزایای زیر را دارد: (۱) جهت‌های افکنش در دسترس خیلی بیش‌تر از LDA است، و (۲) هیچ فرضی روی توزیع داده هر دسته وجود ندارد (یان و همکاران، ۲۰۰۷). به‌تازگی روش حفظ همسایگی و جانشانی تمایز حاشیه‌ای^۵ (NP-MDE) بر اساس جانشانی خطی گراف و روش MFA پیشنهاد شده است؛ که می‌تواند پراکندگی درون‌دسته‌ای را کمینه و هم‌زمان حاشیه بین دسته‌ای را بیشینه کرده، و به‌علاوه ساختار همسایگی با هر دسته را نیز حفظ کند (لان و همکاران، ۲۰۱۲).

روش TRIMAP^۶ روش بانظارت دیگری برای کاهش بعد است (چن و همکاران، ۲۰۱۰). این روش که ترکیبی از روش Isomap (ژو و مارتینز، ۲۰۰۶) و تحلیل تمایز خطی است، ابتدا از Isomap برای پیدا کردن زیرفضای مناسب استفاده می‌کند؛ اما این روش تابع فاصله اقلیدسی معمولی را برای تعیین فواصل به کار نمی‌برد؛ بلکه از تابع فاصله‌ای استفاده می‌کند که فاصله نمونه‌های متعلق به دسته‌های متفاوت را بیش‌تر از حد واقعی تخمین می‌زند. این کار باعث می‌شود تمایز بیش‌تری بین دسته‌های متفاوت ایجاد شود؛ سپس این مدل از تحلیل تمایز خطی برای ایجاد تمایز بیش‌تر میان دسته‌ها استفاده می‌کند.

برای رفع هم‌زمان مشکلات منتج از فرض‌های روش تحلیل تمایز خطی، روش تحلیل تمایز زیردسته مبتنی بر

هسته^۷ ارائه شده است (ژو و مارتینز، ۲۰۰۶). در این روش با استفاده از تقسیم هر دسته به تعدادی زیردسته سعی می‌شود تا بر فرض گاسی تک‌قله‌ای بودن تمام دسته‌ها غلبه شود؛ اما این روش مجبور است هر دسته را به کمک روش‌های خوشه‌بندی به زیرگروه‌هایی تقسیم کند. تحلیل تمایز زیردسته ترکیبی با گام کسری^۸ روش بانظارت دیگری برای کاهش بعد محسوب می‌شود (گاللیس و همکاران، ۲۰۱۳)، که در آن مانند تحلیل تمایز زیردسته هر دسته به زیردسته‌هایی تقسیم می‌شود. ویژگی مهم این روش حل مشکل جدایی زیردسته‌ها است که زمانی رخ می‌دهد که رتبه ماتریس فاصله داخلی درون زیردسته‌ای از تعداد بعد کاهش‌یافته بیش‌تر باشد. نسخه مبتنی بر هسته این روش در (گاللیس و همکاران، ۲۰۱۳) ارائه شده است. در (هان و همکاران، ۲۰۱۳) نیز روشی برای حل مسئله کاهش بعد بانظارت ارائه شده است.

در میان روش‌های کاهش بعد ارائه‌شده تاکنون روش‌های بدون نظارت بسیاری وجود دارد؛ و به دلیل این که بانظارت کردن یک روش می‌تواند نتایج را در کاربردهایی مانند دسته‌بندی بهبود دهد، پژوهش‌هایی در زمینه بانظارت کردن برخی روش‌های کاهش بعد بدون نظارت صورت گرفته است. در (لی و گو، ۲۰۰۶) نسخه‌ای بانظارت برای روش Isomap به نام SE-Isomap پیشنهاد شده است. در این روش ماتریس فاصله ژئودسیک^۹ هر دسته با توجه به اطلاعات مربوط به آن دسته به دست می‌آید؛ و سپس روش مقیاس‌سازی چندبعدی^{۱۰} (MDS) بر روی گرافی که با توجه به این ماتریس‌های فاصله به دست آمده، اعمال می‌شود. در (چن و لیو، ۲۰۱۱) نیز پیشنهادهایی برای بانظارت کردن روش جانشانی محلی خطی^{۱۱} (LLE) مطرح شده است.

در این مقاله ابتدا مدل بانظارت روش ترکیبی تحلیل مؤلفه اصلی احتمالاتی^{۱۲} (PPCMM) ارائه شده که SPPCMM نامیده می‌شود. این مدل توسعه‌ای از مدل ترکیبی تحلیل مؤلفه اصلی احتمالاتی است که اطلاعات مربوط به برجسب داده را در نگاشت یادگیری‌شده منظور می‌کند؛ سپس این رویکرد بانظارت کاهش بعد محلی خطی SPPCMM در چارچوبی به نام جریمه نگاشت^{۱۳} (ژانگ و

⁷ Kernel Sub-class Discriminant Analysis (KSDA)

⁸ Fractional step Mixture Subclass Discriminant Analysis (FMSDA)

⁹ Geodesic

¹⁰ Multidimensional Scaling

¹¹ Locally Linear Embedding

¹² Probabilistic Principal Component Analysis

¹³ Projection penalties

¹ Linear Discriminant Analysis (LDA)

² Kernel Discriminant Analysis (KDA)

³ Marginal Fisher Analysis (MFA)

⁴ Graph Embedding

⁵ Neighborhood Preserving and Marginal Discriminant Embedding

⁶ Tensor-based Riemannian Manifold Distance-Approximating Projection (TRIMAP)

ابعاد $N \times N$ و $M \times M$ بوده و D ماتریسی قطری متشکل از مقادیر منفرد می‌باشد، که به‌ترتیب نزولی مرتب شده و در امتداد قطر قرار گرفته‌اند. ماتریس متشکل از K ستون نخست از ماتریس U که با U_K نشان داده می‌شود، به‌عنوان نگاشت Ψ در نظر گرفته می‌شود (پیرسون، ۱۹۰۱).

روش تحلیل مؤلفه اصلی احتمالاتی (PPCA)، یک بیان احتمالاتی از روش PCA را ارائه می‌کند (تپینگ و بیشاپ، ۱۹۹۹-ب). روش PPCA یک مدل متغیر نهان^۵ است و فرآیند تولید هر ورودی x را به‌صورت $x = W_x z + \mu_x + \varepsilon_x$ مدل می‌کند، که $z \in R^K$ متغیرهای نهان و W_x ماتریسی $M \times K$ است، که factor loading نامیده می‌شود. در این مدل احتمالاتی به‌طور معمول فرض می‌شود که بردار متغیرهای نهان z دارای توزیع گاسین با میانگین صفر و ماتریس کواریانس واحد است (یعنی $z \sim N(0, I)$)، و ε_x نوفه می‌باشد که دارای توزیع گاسین متقارن شعاعی به‌صورت $\varepsilon_x \sim N(0, \delta_x^2 I)$ فرض شده است. پارامترهای مدل با یک روش تکراری به نام الگوریتم میانگین‌گیری-بیشینه‌سازی^۶ EM تعیین می‌شوند. لگاریتم شباهت^۷ برای نمونه داده‌ها در این مدل به‌صورت زیر تعریف می‌شود (تپینگ و بیشاپ، ۱۹۹۹-الف):

$$l = \sum_{n=1}^N \ln \{p(x_n)\} \quad (1)$$

که در آن $p(x_n)$ توزیع حاشیه‌ای x است، و به‌صورت $p(x_n) = \int p(x|z)p(z)dz$ تعریف می‌شود. نتایج مرتبط می‌تواند در قالب قضیه ۱ به‌صورت زیر خلاصه شود. جزییات اثبات این قضیه در مرجع (یو و همکاران، ۲۰۰۶) و (تپینگ و بیشاپ، ۱۹۹۹-الف) قابل دستیابی است:

قضیه ۱: اگر $S_x = \frac{1}{N} \sum_{n=1}^N (x_n - \mu_x)(x_n - \mu_x)^T$ ماتریس کواریانس نمونه داده‌های $\{x_n\}_{n=1}^N$ باشد و $\lambda_1 \geq \dots \geq \lambda_M$ مقادیر ویژه متناظر با بردارهای ویژه $u_1, \dots, u_M$ ماتریس S_x باشد، و فضای نهان در مدل PPCA، فضایی K بعدی باشد، آن‌گاه (یو و همکاران، ۲۰۰۶):

(الف) تخمین بیش‌ترین شباهت برای W_x به‌صورت زیر به‌دست می‌آید:

$$W_x = U_K (\Lambda_K - \delta_x^2 I)^{\frac{1}{2}} R \quad (2)$$

اشنایدر، ۲۰۱۰) به‌کار گرفته می‌شود. جریمه نگاشت این امکان را فراهم می‌آورد که از کاهش بعد برای بهبود دقت مدل پیش‌گو بدون خطر از دست‌دادن اطلاعات مرتبط استفاده کنیم (ژانگ و اشنایدر، ۲۰۱۰). چنان‌که در بخش ۴ خواهیم دید، نتایج آزمایش‌ها بر روی پایگاه داده‌های تصاویر چهره، نشان‌دهنده کارایی بالاتر روش پیشنهادی SPPCamm در چهارچوب جریمه نگاشت نسبت به سایر روش‌های بررسی شده است.

ساختار مقاله به‌صورت زیر است: در بخش ۲ مروری بر کارهای مرتبط خواهیم داشت، و چند روش کاهش بعد بررسی می‌شود. در بخش ۳ روش پیشنهادی یعنی مدل با نظارت روش ترکیبی تحلیل مؤلفه اصلی احتمالاتی و به‌کارگیری آن در چارچوب جریمه نگاشت معرفی می‌شود. نتایج پیاده‌سازی‌ها در بخش ۴ گزارش می‌شوند. در نهایت تحلیل نتایج در بخش ۵، و جمع‌بندی مقاله در بخش ۶ ارائه خواهد شد.

۲- مروری بر کارهای مرتبط

در این بخش روش‌های کاهش بعد مرتبط با روش پیشنهادی را مورد بررسی قرار می‌دهیم. این روش‌ها عبارت از روش تحلیل مؤلفه اصلی^۱ (PCA)، تحلیل مؤلفه اصلی احتمالاتی^۲ (PPCA)، تحلیل مؤلفه اصلی بانظارت^۳ (SPPCA)، و مدل ترکیبی تحلیل مؤلفه اصلی احتمالاتی^۴ (PPCamm) (تپینگ و بیشاپ، ۱۹۹۹-الف) هستند. در ادامه فرض می‌کنیم مجموعه‌ای از N نمونه آموزشی در اختیار است. ورودی m ام که با $x_n \in X$ نمایش داده می‌شود، توسط یک بردار ویژگی M بعدی مشخص می‌شود ($x_n \in X \subseteq R^M$). برای کاهش بعد، هدف یافتن نگاشتی به‌صورت $\Psi: X \rightarrow Z$ است که هر ورودی را به یک فضای K بعدی با شرط $K < M$ نگاشت می‌کند ($Z \subseteq R^K$). روش PCA یک روش کاهش بعد بدون نظارت است که هدف آن یافتن مؤلفه‌های اصلی، که جهت‌هایی با بیش‌ترین واریانس را برای داده مشخص می‌کنند، است. اگر فرض کنیم $X = [x_1, \dots, x_N]^T$ ماتریس ورودی متمرکز شده را مشخص کند، که در آن میانگین نمونه‌ها از هر نمونه کم شده است، $X = VDU^T$ تجزیه مقدار منفرد X را مشخص می‌کند که V و U ماتریس‌های متعامد ستونی با

¹ Principal Component Analysis

² Probabilistic Principal Component Analysis

³ Supervised Probabilistic Principal Component Analysis

⁴ PPCA Mixture Model

⁵ Latent Variable

⁶ Expectation Maximization

⁷ Log-Likelihood

که $\Lambda_K = \text{diag}(\lambda_1, \dots, \lambda_K)$ و $U_K = [u_1, \dots, u_K]$ و R یک ماتریس متعامد $K \times K$ است.

(ب) نتیجه نگاشت برای ورودی جدید x^* ، بردار z^* است که به صورت زیر به دست می آید:

$$z^* = R^T (\Lambda_K - \delta_x^2 I)^{-\frac{1}{2}} \Lambda_K^{-1} U_K^T (x^* - \mu) \quad (3)$$

مدل PPCA به دلیل بیان در چارچوب احتمالاتی، مزایایی نسبت به PCA دارد، که از آن جمله می توان رویه یادگیری سریع الگوریتم EM، ارائه روشی اصولی برای مدیریت ورودی های مفقود، و امکان در نظر گرفتن مدل ترکیبی برای PCA را بر شمرد (تپینگ و بیشاپ، ۱۹۹۹-الف) (تپینگ و بیشاپ، ۱۹۹۹-ب).

در مسائل یادگیری بانظارت هر داده x با یک مقدار خروجی $y \in Y$ متناظر است. مدل SPPCA نسخه بانظارت مدل PPCA است، که در آن برای مسایل رگرسیون $y \in R$ و برای مسایل دسته بندی $y \in \{+1, -1\}$ در نظر گرفته می شود (یو و همکاران، ۲۰۰۶). مدل SPPCA بر این باور است که بین فضای ورودی X و فضای خروجی Y کواریانس وجود دارد و این امکان وجود دارد که PPCA را به مدلی که این کواریانس را به خوبی نشان می دهد تعمیم داد. بیان ریاضی روش SPPCA به صورت زیر است: اگر بعد فضای خروجی را L فرض کنیم، هر ورودی x متناظر با یک بردار خروجی $y = [y_1, \dots, y_L]^T \in Y \subseteq R^L$ است. در SPPCA داده مشاهده شده (x, y) توسط مدل متغیر نهان زیر تولید می شود:

$$\begin{aligned} x &= W_x z + \mu_x + \epsilon_x \\ y &= W_y z + \mu_y + \epsilon_y \end{aligned} \quad (4)$$

متغیرهای نهان $z \sim N(0, I)$ بین x و y مشترک می باشند، و دو مدل نوفه که از یکدیگر مستقل هستند و هر دو دارای توزیع ایزوتوپیک گاسین هستند، به صورت $\epsilon_x \sim N(0, \delta_x^2 I)$ و $\epsilon_y \sim N(0, \delta_y^2 I)$ در این روش اگر قرار دهیم:

$$\Phi = \begin{pmatrix} \delta_x^2 I & 0 \\ 0 & \delta_y^2 I \end{pmatrix}, \quad \mu = \begin{pmatrix} \mu_x \\ \mu_y \end{pmatrix}, \quad W = \begin{pmatrix} W_x \\ W_y \end{pmatrix} \quad (5)$$

بر اساس فرضیاتی که برای این روش در نظر گرفته شده است، به راحتی دیده می شود که (x, y) دارای توزیع توأم گاسین با میانگین μ و ماتریس کواریانس $\Phi + WW^T$ هستند (یو و همکاران، ۲۰۰۶). نتایج مرتبط با این روش را می توان به طور خلاصه در قالب قضیه ۲ به شکل زیر بیان

کرد (جزئیات این قضیه در مرجع (یو و همکاران، ۲۰۰۶) موجود است):

قضیه ۲: ماتریس کواریانس نمونه داده های هنجار سازی شده $\{(x_n, y_n)\}_{n=1}^N$ را به صورت زیر تعریف می کنیم:

$$S = \frac{1}{N} \sum_{n=1}^N \Phi^{-\frac{1}{2}} \begin{pmatrix} x_n \\ y_n \end{pmatrix} \begin{pmatrix} x_n \\ y_n \end{pmatrix}^T \quad \Phi^{-\frac{1}{2}} = \begin{pmatrix} \frac{1}{\delta_x^2} S_x & \frac{1}{\delta_x \delta_y} S_{xy} \\ \frac{1}{\delta_x \delta_y} S_{yx} & \frac{1}{\delta_y^2} S_y \end{pmatrix} \quad (6)$$

و $\lambda_1 \geq \dots \geq \lambda_{(M+L)}$ مقادیر ویژه متناظر با بردارهای ویژه $u_1, \dots, u_{(M+L)}$ از ماتریس S می باشند. ماتریس S_x ماتریس کواریانس بین ورودی ها، S_{xy} و S_{yx} ماتریس کواریانس بین خروجی ها، و S_y ماتریس کواریانس بین ورودی ها است. اگر فضای نهان را K بعدی در نظر بگیریم، نتایج زیر را خواهیم داشت (یو و همکاران، ۲۰۰۶):

(الف) در SPPCA، ماتریس های W_x و W_y به صورت زیر به دست می آیند:

$$\begin{aligned} W_x &= \delta_x U_x (\Lambda_K - I)^{\frac{1}{2}} R \\ W_y &= \delta_y U_y (\Lambda_K - I)^{\frac{1}{2}} R \end{aligned} \quad (7)$$

که $\Lambda_K = \text{diag}(\lambda_1, \dots, \lambda_K)$ و U_x شامل M سطر نخست از ماتریس $[u_1, \dots, u_K]$ بوده و U_y حاوی L سطر باقیمانده از آن است. ماتریس R یک ماتریس متعامد دلخواه $K \times K$ است.

(ب) نتیجه نگاشت برای ورودی متمرکز شده x^* ، بردار z^* است که به صورت زیر به دست می آید:

$$z^* = \frac{1}{\delta_x} R^T (\Lambda_K - I)^{-\frac{1}{2}} [U_x^T U_x + (\Lambda_K - I)^{-1}]^{-1} U_x^T x^* \quad (8)$$

در حالت خاص اگر $L=0$ باشد، آن گاه مدل بدون نظارت است و $S = \frac{1}{\delta_x^2} S_x$ خواهد بود، که مشابه ماتریس مربوط به PPCA خواهد شد (یو و همکاران، ۲۰۰۶).

در مدل ترکیبی تحلیل مؤلفه های اصلی احتمالاتی (PPCMM) از ترکیب خطی این مدل ها بر روی مجموعه داده ها استفاده می شود. در واقع روش PPCMM ترکیب چند PPCA محلی است که در این روش بخش بندی مناسب نمونه داده ها و تعیین مؤلفه های اصلی مدل برای هر بخش به طور هم زمان با استفاده از تخمین بیش ترین شباهت تعیین می شود. لگاریتم شباهت نمونه داده ها برای این مدل ترکیبی به صورت زیر تعیین می شود:

همان‌طور که ذکر شد، تیپینگ و بیشاپ روشی مطرح کردند که در آن PCA در چارچوب بیش‌ترین شباهت^۲ و براساس صورت خاصی از مدل متغیر نهان گاسی فرمول‌بندی شده است. این روش که به‌عنوان روش تحلیل مؤلفه اصلی احتمالاتی (PPCA) معرفی شده است، منجر به ایجاد یک مدل ترکیبی خوش‌تعریف برای تحلیل مؤلفه اصلی احتمالاتی محلی خطی می‌شود، که پارامترهای آن می‌تواند با استفاده از الگوریتم تکراری EM تعیین شوند (تیپینگ و بیشاپ، ۱۹۹۹-الف)؛ اما این روش هنگامی که خروجی‌های مربوط به نمونه‌داده‌های ورودی در اختیار باشد، مانند مسائل دسته‌بندی و رگرسیون، قادر به استفاده مفید از این اطلاعات نیست.

با توجه به اهمیت مدل ترکیبی تحلیل مؤلفه اصلی احتمالاتی در فضاهای غیر خطی و پیچیده و نیز غلبه بر محدودیت‌های روش PCA و نیز با توجه به این موضوع که روش‌های کاهش بُعد بانظارت در مسائل مختلف به‌خصوص دسته‌بندی و پیش‌بینی، به‌دلیل استفاده از اطلاعات مسأله از کارایی بهتری برخوردار هستند، ارائه روش بانظارت مدل ترکیبی تحلیل مؤلفه اصلی احتمالاتی در این مقاله مورد توجه قرار گرفته است. روش ارائه‌شده تحت عنوان SPPCMM معرفی می‌شود. در روش پیشنهادی از ترکیب خطی مدل‌ها بر روی مجموعه داده‌ها با در نظر گرفتن برجسب دسته‌ها استفاده می‌شود.

روش SPPCMM را می‌توان ترکیب چند SPPCA محلی دانست، که در آن بخش‌بندی مناسب نمونه‌داده‌ها و تعیین مؤلفه‌های اصلی مدل برای هر بخش به‌طور هم‌زمان با استفاده از تخمین بیش‌ترین شباهت انجام می‌شود. لگاریتم شباهت نمونه‌داده‌ها برای این مدل مشابه PPCAMM (تیپینگ و بیشاپ، ۱۹۹۹-الف) بوده و به‌صورت زیر تعریف می‌شود، با این تفاوت که در این حالت هر نمونه به‌صورت $(\mathbf{x}_n, \mathbf{y}_n)$ است:

$$l = \sum_{n=1}^N \ln\{p(\mathbf{x}_n, \mathbf{y}_n)\} = \sum_{n=1}^N \ln\{\sum_{i=1}^M \pi_i p(\mathbf{x}_n, \mathbf{y}_n | i)\} \quad (14)$$

که M تعداد مدل‌های ترکیبی، N تعداد نمونه داده‌ها و $p(\mathbf{x}_n, \mathbf{y}_n | i)$ یک مدل SPPCA است. اگر پاسخ احتمال پسین مدل λ ام برای تولید نقطه داده $(\mathbf{x}_n, \mathbf{y}_n)$ را به‌صورت $\mathbf{R}_{ni} = p(i | \mathbf{x}_n, \mathbf{y}_n)$ بیان کنیم، از رابطه زیر می‌توان این احتمال را محاسبه کرد:

$$\mathbf{R}_{ni} = \frac{p(\mathbf{x}_n, \mathbf{y}_n | i) \pi_i}{p(\mathbf{x}_n, \mathbf{y}_n)} \quad (15)$$

² Maximum likelihood

$l = \sum_{n=1}^N \ln\{p(\mathbf{x}_n)\} = \sum_{n=1}^N \ln\{\sum_{i=1}^M \pi_i p(\mathbf{x}_n | i)\}$ (۹) که M تعداد مدل‌های مشارکت‌کننده در ترکیب، N تعداد نمونه داده‌ها، و $p(\mathbf{x}_n | i)$ یک مدل PPCA است. ضریب π_i نسبت مشارکت مؤلفه مدل λ ام در ترکیب مورد نظر می‌باشد، که برای آن محدودیت‌های احتمالاتی $\pi_i \geq 0$ و $\sum_i \pi_i = 1$ برقرار است. هر مؤلفه از ترکیب بردار میانگین μ_i و همچنین پارامترهای $\mathbf{W}_{xi}$ و δ_{xi} مربوط به خود را دارد (تیپینگ و بیشاپ، ۱۹۹۹-الف).

می‌توان برای بهینه‌کردن پارامترهای هر مؤلفه از ترکیب الگوریتم تکراری EM را مورد استفاده قرار داد. نشان داده می‌شود $\mathbf{R}_{ni} = p(i | \mathbf{x}_n)$ ، که پاسخ احتمال پسین مؤلفه λ ام برای تولید نقطه داده $\mathbf{x}_n$ را بیان می‌کند، به‌صورت زیر به‌دست می‌آید (تیپینگ و بیشاپ، ۱۹۹۹-الف):

$$\mathbf{R}_{ni} = \frac{p(\mathbf{x}_n | i) \pi_i}{p(\mathbf{x}_n)} \quad (10)$$

همچنین می‌توان نشان داد که π_i و μ_{xi} به‌صورت زیر به‌دست می‌آیند (تیپینگ و بیشاپ، ۱۹۹۹-الف):

$$\tilde{\pi}_i = \frac{1}{N} \sum_{n=1}^N \mathbf{R}_{ni} \quad (11)$$

$$\tilde{\mu}_{xi} = \frac{\sum_{n=1}^N \mathbf{R}_{ni} \mathbf{x}_n}{\sum_{n=1}^N \mathbf{R}_{ni}} \quad (12)$$

به‌علاوه، $\mathbf{W}_{xi}$ و δ_{xi}^2 با استفاده از بردارهای ویژه و مقادیر ویژه ماتریس کواریانس وزن‌دار محلی مدل λ ام که به‌صورت زیر تعریف می‌شود، به‌دست می‌آیند (تیپینگ و بیشاپ، ۱۹۹۹-الف):

$$\mathbf{S}_{xi} = \frac{1}{\tilde{\pi}_i N} \sum_{n=1}^N \mathbf{R}_{ni} (\mathbf{x}_n - \tilde{\mu}_{xi}) (\mathbf{x}_n - \tilde{\mu}_{xi})^T \quad (13)$$

۳- روش پیشنهادی

۳-۱- مدل ترکیبی تحلیل مؤلفه اصلی

احتمالاتی بانظارت

روش PCA با وجود محدودیتی که به‌خاطر خطی‌بودن دارد، یکی از پرکاربردترین روش‌های کاهش بُعد است. تعمیم‌های مختلفی از روش PCA برای غلبه بر محدودیت خطی‌بودن آن مطرح شده است. یک پیشنهاد مطرح این است که مجموعه داده‌های پیچیده به‌صورت ترکیبی خطی محلی از نگاشت‌های PCA بیان شود. با این وجود PCA متناظر با یک چگالی احتمالاتی نیست و لذا روش منحصره‌فردی برای ترکیب مدل‌های PCA وجود ندارد. پیشنهادهای قبلی برای توسعه یک مدل ترکیبی^۱ برای PCA تک‌منظوره بوده و عمومیت نداشته‌اند.

¹ Mixture model

که π_i نسبت ترکیب مدل i ام مورد نظر با محدودیت‌های $\pi_i \geq 0$ و $\sum_i \pi_i = 1$ است، و به صورت زیر به دست می‌آید:

$$\tilde{\pi}_i = \frac{1}{N} \sum_{n=1}^N R_{ni} \quad (16)$$

رابطه بالا در گام M از الگوریتم EM و با بیشینه کردن لگاریتم شباهت داده نسبت به π_i به دست می‌آید. همچنین با بیشینه کردن لگاریتم شباهت نسبت به پارامترهای میانگین x ها (μ_{xi})، میانگین y ها (μ_{yi})، δ_{xi}^2 ، δ_{yi}^2 ، W_{xi} و W_{yi} مقادیر جدید برای این پارامترها به دست می‌آید. به علاوه پارامترهای مجهول هر مؤلفه را می‌توان با تجزیه مقدار ویژه ماتریس کواریانس محلی داده برای آن مؤلفه به دست آورد. ماتریس کواریانس محلی برای هر مدل به صورت زیر قابل تعریف است:

$$S_i = \frac{1}{N} \sum_{n=1}^N \Phi_i \frac{1}{2} \begin{pmatrix} x_n - \mu_{xi} \\ y_n - \mu_{yi} \end{pmatrix} \begin{pmatrix} x_n - \mu_{xi} \\ y_n - \mu_{yi} \end{pmatrix}^T \quad (17)$$

$$\Phi_i \frac{1}{2} = \begin{pmatrix} \frac{1}{\delta_{xi}^2} S_{xi} & \frac{1}{\delta_{xi} \delta_{yi}} S_{xyi} \\ \frac{1}{\delta_{xi} \delta_{yi}} S_{xyi} & \frac{1}{\delta_{yi}^2} S_{yi} \end{pmatrix} \quad (18)$$

که در این رابطه S_{xi} به صورت زیر تعریف می‌شود:

$$S_{xi} = \frac{1}{\pi_i N} \sum_{n=1}^N R_{ni} (x_n - \tilde{\mu}_{xi})(x_n - \tilde{\mu}_{xi})^T \quad (19)$$

ماتریس‌های S_{xyi} ، S_{yi} و S_{yxi} نیز به صورت مشابه با رابطه بالا محاسبه می‌شوند.

با به دست آوردن مقادیر ویژه و بردارهای ویژه ماتریس کواریانس بالا می‌توان W_{xi} و W_{yi} برای هر مدل را با استفاده از روابط زیر به دست آورد:

$$W_{xi} = \delta_{xi} U_{xi} (\Lambda_K - I)^{\frac{1}{2}} R$$

$$W_{yi} = \delta_{yi} U_{yi} (\Lambda_K - I)^{\frac{1}{2}} R \quad (20)$$

که $\Lambda_K = \text{diag}(\lambda_1, \dots, \lambda_K)$ و U_{xi} شامل M سطر نخست از ماتریس $[u_{1i}, \dots, u_{Ki}]$ ، U_{yi} حاوی L سطر باقیمانده از آن است. ماتریس R یک ماتریس متعامد دلخواه $K \times K$ است.

نگاشت z^* برای ورودی جدید x^* در فضای کاهش بعد یافته به صورت زیر خواهد بود:

$$z^* = \frac{1}{\delta_{xi}} R^T (\Lambda_K - I)^{-\frac{1}{2}} [U_{xi}^T U_{xi} + (\Lambda_K - I)^{-1}]^{-1} U_{xi}^T (x^* - \mu_{xi}) \quad (21)$$

رابطه بالا بردار ویژگی کاهش بعدیافته برای ورودی x^* به کمک مؤلفه i ام از ترکیب را نشان می‌دهد. بردار ویژگی کاهش بعدیافته نهایی به صورت زیر به دست می‌آید. ابتدا بردار ویژگی کاهش بعدیافته با هر مؤلفه از ترکیب را با استفاده از رابطه (۲۱) به دست می‌آوریم؛ سپس فاصله اقلیدسی بازیابی بردار کاهش بعدیافته از بردار اصلی را به دست می‌آوریم. مؤلفه‌ای که کم‌ترین فاصله اقلیدسی را به دست دهد پیدا شده و بردار ویژگی کاهش بعدیافته نهایی با افکنش بردار اصلی بر آن مؤلفه به دست می‌آید. محاسبه مقادیر δ_{xi}^2 و δ_{yi}^2 با استفاده از روابط زیر انجام می‌شود (یو و همکاران، ۲۰۰۶) (تپینگ و بیشاپ، ۱۹۹۹-الف):

$$\delta_{xi}^2 = \frac{1}{M \sum_n R_{ni}} \left\{ \sum_{n=1}^N R_{ni} \|x_n - \mu_{xi}\|^2 + R_{ni} \text{tr}(W_x^T W_x C) - 2 R_{ni} \text{tr}(x W_x Z^T) \right\}$$

$$\delta_{yi}^2 = \frac{1}{M \sum_n R_{ni}} \left\{ \sum_{n=1}^N R_{ni} \|y_n - \mu_{yi}\|^2 + R_{ni} \text{tr}(W_y^T W_y C) - 2 R_{ni} \text{tr}(y W_y Z^T) \right\} \quad (22)$$

در روابط بالا $Z = [z_1, \dots, z_N]^T$ ، $X = [x_1, \dots, x_N]^T$ و $Y = [y_1, \dots, y_N]^T$ است، که $C = \sum_{n=1}^N z_n z_n^T$ و $\langle z_n z_n^T \rangle$ با تخمین توزیع z_n برای مشاهدات (x_n, y_n) به دست می‌آیند. وقتی که ورودی x و خروجی y هر دو در دست باشند، توزیع پسین z برای (x, y) داده شده به صورت زیر تعریف می‌شود (یو و همکاران، ۲۰۰۶):

$$p(z|x, y) \propto p(x, y|z) p(z) = p(x|z) p(y|z) p(z) \quad (23)$$

با توجه به این که سه جمله سمت راست دارای توزیع گاسین هستند، توزیع پسین z نیز گاسین $(N(\mu_z, \Sigma_z))$ خواهد بود، و داریم:

$$\mu_z = A^{-1} \left[\frac{1}{\delta_x^2} W_x^T (x - \mu_x) + \frac{1}{\delta_y^2} W_y^T (y - \mu_y) \right]$$

$$\Sigma_z = A^{-1}$$

$$A = \frac{1}{\delta_x^2} W_x^T W_x + \frac{1}{\delta_y^2} W_y^T W_y + I \quad (24)$$

بنابراین برای هر مدل در SPPCMM نیز مشابه بالا داریم:

$$A_i = \frac{1}{\delta_{xi}^2} W_{xi}^T W_{xi} + \frac{1}{\delta_{yi}^2} W_{yi}^T W_{yi} + I \quad (25)$$

و آماره‌های $\langle z_n \rangle$ و $\langle z_n z_n^T \rangle$ با استفاده از روابط بالا به صورت زیر تعریف می‌شوند:

دیدگاه مدل جستجو، همه نامزدهایی که در M_p نیستند از دامنه جستجو حذف می‌شوند.

بدین ترتیب کاهش بعد باعث از بین رفتن اطلاعات می‌شود، و هیچ تضمینی وجود ندارد که پارامترهای به‌دست آمده در فضای کاهش بعد یافته با پارامترهای بهینه در فضای اصلی یکسان باشند؛ لذا در روش جریمه نگاشت پیشنهاد شده است که جستجوی پارامترها در فضای اصلی انجام شود، و فاصله آن از زیرفضای به‌دست آمده به‌عنوان جریمه در نظر گرفته شود (ژانگ و اشنادیر، ۲۰۱۰). با این روش مدل پیش‌گو به‌صورت زیر در خواهد آمد:

$$\argmin_{\mathbf{w}} \sum_{i=1}^N L(y_i, \mathbf{w}^T \mathbf{x}_i + b) + \lambda J(\mathbf{w}) \quad (30)$$

$\mathbf{w} \in \mathbb{R}^p, \mathbf{v} \in \mathbb{R}^d, b$

که $\tilde{\mathbf{w}} = \mathbf{w} - \mathbf{P}^T \mathbf{v}$ و J تابع جریمه است و می‌تواند به‌صورت $\|\cdot\|_2$ یا $\|\cdot\|_1$ در نظر گرفته شود. متغیر λ نیز پارامتر تنظیم است. اگر مدل پیش‌گوی مورد نظر و عمل‌گر کاهش بعد طوری طراحی شوند که روی فضای هسته کار کنند، مدل پیش‌گو با جریمه نگاشت به‌صورت زیر خواهد بود:

$$\argmin_{\mathbf{w} \in \mathbb{R}^p, b} \sum_{i=1}^N L(y_i, \mathbf{w}^T \Phi(\mathbf{x}_i) + b) + \min_{\mathbf{v} \in \mathbb{R}^d} \lambda J(\mathbf{w} - \mathbf{P}^T \mathbf{v}) \quad (31)$$

که $\Phi(\mathbf{x}_i)$ بازنمایی نمونه $\mathbf{x}_i$ در فضای هسته است. از آن‌جا که مسأله، دسته‌بندی است، در رابطه بالا از SVM با متغیرهای سستی $\{\xi_i\}_{i=1}^N$ به‌عنوان تابع L استفاده می‌شود. به این ترتیب مسأله برنامه‌ریزی درجه دو ی زیر را داریم:

$$\argmin_{\mathbf{w} \in \mathbb{R}^p, \mathbf{v} \in \mathbb{R}^d, b, \xi_i} \frac{1}{2} \|\mathbf{w} - \mathbf{P}^T \mathbf{v}\|_2^2 + C \sum_{i=1}^N \xi_i$$

$s.t. \quad y_i(\mathbf{w}^T \Phi(\mathbf{x}_i) + b) \geq 1 - \xi_i, \forall i$

$\xi_i \geq 0 \quad \forall i \quad (32)$

که $C \propto \frac{1}{\lambda}$ به‌جای λ در رابطه (۳۱) قرار گرفته است. با قراردادن $\tilde{\mathbf{w}} = \mathbf{w} - \mathbf{P}^T \mathbf{v}$ داریم [۱۶]:

$$\argmin_{\tilde{\mathbf{w}} \in \mathbb{R}^p, \mathbf{v} \in \mathbb{R}^d, b, \xi_i} \frac{1}{2} \|\tilde{\mathbf{w}}\|_2^2 + C \sum_{i=1}^N \xi_i$$

$s.t. \quad y_i(\tilde{\mathbf{w}}^T \Phi(\mathbf{x}_i) + \mathbf{v}^T \mathbf{P} \Phi(\mathbf{x}_i) + b) \geq 1 - \xi_i, \forall i$

$\xi_i \geq 0 \quad \forall i \quad (33)$

$$\langle \mathbf{z}_{ni} \rangle = \mathbf{A}_i^{-1} \left[\frac{1}{\delta_{xi}^2} \mathbf{W}_{xi}^T (\mathbf{x}_n - \mu_{xi}) + \frac{1}{\delta_{yi}^2} \mathbf{W}_{yi}^T (y_n - \mu_{yi}) \right]$$

$\langle \mathbf{z}_{ni} \mathbf{z}_{ni}^T \rangle = \mathbf{A}_i^{-1} + \langle \mathbf{z}_{ni} \rangle \langle \mathbf{z}_{ni} \rangle^T \quad (26)$

در شکل (۲) کارنمای روش پیشنهادی نمایش داده شده است.

۳-۲- کاهش بعد محلی خطی بدون اتلاف با استفاده از جریمه نگاشت

همان‌طور که پیش‌تر گفته شد، در این مقاله روشی برای شناسایی چهره با استفاده از کاهش بعد بدون اتلاف ارائه می‌شود. یادگیری یک مدل پیش‌بینی خطی با پارامترهای $(\mathbf{w}, b)$ در یک فضای p بعدی با N نمونه آزمایشی $\{\mathbf{x}_i, y_i\}_{i=1}^N$ به‌صورت زیر قابل تعریف است:

$$\argmin_{\mathbf{w} \in \mathbb{R}^p} \sum_{i=1}^N L(y_i, \mathbf{w}^T \mathbf{x}_i + b) \quad (27)$$

که $\mathbf{w} \in \mathbb{R}^p$ بردار پارامتر و b پارامتر مدل پیش‌گو را بازنمایی می‌کند. متغیرهای $\mathbf{x}_i$ و y_i به‌ترتیب بردار ویژگی p بعدی و خروجی مربوط به نمونه i ام را مشخص می‌کنند. تابع L یک تابع اتلاف است، و به مدل پیش‌گو بستگی دارد که می‌تواند به‌صورت مربعات خطا در نظر گرفته شود.

اگر $\mathbf{P}$ یک ماتریس با ابعاد $d \times p$ در نظر بگیریم، که یک کاهش بعد خطی به d بعد بر روی فضای ورودی p بعدی انجام می‌دهد، بازنمایی کاهش‌بعدیافته ورودی $\mathbf{x}$ به‌صورت $\mathbf{P}\mathbf{x}$ خواهد بود. با اعمال کاهش بعد خطی مدل پیش‌گو در فضای کاهش بعد یافته به‌صورت زیر تعریف می‌شود:

$$\argmin_{\mathbf{v} \in \mathbb{R}^d, b} \sum_{i=1}^N L(y_i, \mathbf{v}^T (\mathbf{P}\mathbf{x}_i) + b) \quad (28)$$

که $\mathbf{v} \in \mathbb{R}^d$ بردار پارامتر یادگرفته‌شده در فضای کاهش‌بعدیافته است.

از مقایسه رابطه (۲۷) با (۲۸) ارتباط بین کاهش بعد فضای ویژگی و محدودیت روی فضای پارامتر را می‌بینیم. استفاده از ماتریس کاهش بعد $\mathbf{P}$ در فضای ورودی، موجب می‌شود که بردار پارامتر $\mathbf{w}$ با بعد p بر زیرفضای زیر واقع شود:

$$M_p = \{\mathbf{w} \in \mathbb{R}^p \mid \mathbf{w} = \mathbf{P}^T \mathbf{v}, \exists \mathbf{v} \in \mathbb{R}^d\} \quad (29)$$

بنابراین اعمال عمل‌گر کاهش بعد خطی $\mathbf{P}$ در فضای ورودی معادل با محدودکردن مدل جستجو به یک زیرفضای پارامتر M_p است، که در رابطه (۲۹) تعریف شده است. از

¹ Slack Variables

² Quadratic Programming

تا این جا کاهش بعد با در نظر گرفتن جریمه نگاشت برای روش های کاهش بعد خطی و غیر خطی مبتنی بر هسته بیان شد. آنچه در این مقاله مورد بررسی قرار می گیرد، کاهش بعد محلی خطی با در نظر گرفتن جریمه نگاشت است. در روابط (۳۰) و (۳۳) عمل گر کاهش بعد P تنها از طریق جمله Px_i (در رابطه (۳۰)) و $P\Phi(x_i)$ (در رابطه (۳۳)) ظاهر شده است. می توان با یک ترفند ساده این دو جمله را با یک تابع کاهش بعد عمومی مانند $\Psi(x_i)$ جایگزین کرد [۱۶]. در این مقاله یک خمینه^۱ زیربنایی محلی خطی با استفاده از روش ارائه شده برای مدل ترکیبی تحلیل مؤلفه اصلی احتمالاتی بانظارت (SPPCMM) از نمونه داده ها به دست می آید، که همان تابع کاهش بعد $\Psi(x_i)$ خواهد بود. با جای گذاری این تابع به جای P در رابطه (۳۳) مسأله بهینه سازی زیر را خواهیم داشت:

$$\begin{aligned} \argmin_{\tilde{w} \in R^P, v \in R^d, b, \xi_i} & \frac{1}{2} \|\tilde{w}\|_2^2 + C \sum_{i=1}^n \xi_i \\ s.t. & y_i (\tilde{w}^T \Phi(x_i) + v^T \Psi(x_i) + b) \geq 1 - \xi_i, \forall i \\ & \xi_i \geq 0 \quad \forall i \end{aligned} \quad (34)$$

درواقع در این روش به جای نگاشت داده بر خمینه زیربنایی و جستجوی پارامترهای مدل پیش گو در آن که موجب از دست رفتن اطلاعات می شود، همان طور که در روابط بالا بیان شد، جستجو را در فضای اصلی انجام می دهیم و فاصله پارامترهای به دست آمده در فضای اصلی از پارامترهای به دست آمده در خمینه محلی خطی حاصل از SPPCMM را به عنوان یک جریمه نگاشت در یادگیری مدل پیش گو مورد استفاده قرار می دهیم.

۴- نتایج پیاده سازی

در این بخش به بررسی کارایی روش پیشنهادی پرداخته و آن را با چند روش کاهش بعد دیگر مورد مقایسه قرار می دهیم. برای این منظور مسئله شناسایی چهره را در نظر گرفته ایم، و آزمایش ها را بر روی پایگاه داده های چهره ORL و Yale انجام داده ایم. پایگاه داده Yale شامل ۱۶۵ تصویر سیاه سفید از پانزده نفر است. برای هر شخص یازده تصویر با حالت های مختلف چهره و شرایط نوری متفاوت وجود دارد. در این مقاله از تصاویر این پایگاه داده که بر روی سایت دانشگاه MIT قرار دارد^۲ و شامل تصاویر هنجار سازی شده

است، استفاده گردید. در این هنجار سازی، تصاویر بر اساس فاصله بین دو چشم تنظیم شده و سپس تصویر چهره از عکس استخراج می شود. پایگاه داده ORL شامل تصاویری از چهل نفر است. از هر نفر ده تصویر متفاوت وجود دارد. تصاویر این پایگاه داده، تغییراتی از قبیل شدت روشنایی، حالت های چهره (چشمان باز و بسته، و لبخند)، تصاویر با عینک و بدون عینک، مقیاس و فاصله و چرخش زاویه سر را شامل می شود. در آزمایش بر روی این پایگاه داده ها ابتدا هر تصویر به 32×32 پیکسل تغییر اندازه داده می شود.

در آزمایش ها از روش استخراج ویژگی گابور^۳ برای به دست آوردن بردار ویژگی مربوط به هر تصویر استفاده شده است. فیلترهای گابور ابزاری برای استخراج ویژگی هستند. این توابع دوبعدی، لبه اشکال و همچنین گودی ها و برآمدگی های تصویر را تقویت می کنند. برای افزایش تمایز چشم ها، دهان و بینی که به عنوان اجزای مهم و اساسی چهره مطرح هستند، فیلترهای گابور می توانند مفید باشند (حقیقت و همکاران ۲۰۱۳). طول بردار ویژگی استخراج شده با این روش ۲۵۶۰ است.

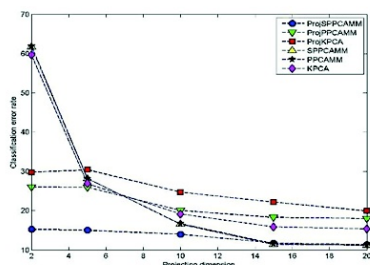
در آزمایش ها روش های زیر با هم مقایسه شده اند: (۱) روش ماشین بردار پشتیبان (SVM) در فضای ویژگی کاهش بعد یافته با روش تحلیل مؤلفه اصلی هسته (KPCA) برای کاهش بعد، که به اختصار آن را SVM-KPCA می نامیم. این روش ابتدا KPCA را اجرا کرده، و سپس SVM خطی را روی فضای کاهش بعد یافته اجرا می کند. (۲) روش ماشین بردار پشتیبان در فضای ویژگی کاهش بعد یافته با روش ترکیبی تحلیل مؤلفه اصلی احتمالاتی که آن را به اختصار SVM-PPCMM می نامیم. این روش ابتدا PPCMM را اجرا کرده، و سپس SVM خطی را روی فضای کاهش بعد یافته اجرا می کند. (۳) روش ماشین بردار پشتیبان در فضای ویژگی کاهش بعد یافته با روش پیشنهادی مدل ترکیبی تحلیل مؤلفه اصلی احتمالاتی بانظارت که آن را به اختصار SVM-SPPCMM می نامیم. در این روش ابتدا SPPCMM را اجرا کرده، و سپس SVM خطی را روی فضای کاهش بعد یافته اجرا می کنیم. (۴) روش KPCA با در نظر گرفتن جریمه نگاشت که به اختصار آن را Proj-KPCA می نامیم. (۵) مدل ترکیبی تحلیل مؤلفه اصلی احتمالاتی با در نظر گرفتن جریمه نگاشت، که آن را به اختصار Proj-PPCMM می نامیم. (۶) به کارگیری روش پیشنهادی

³ Gabor

¹ Manifold

² <http://vismod.media.mit.edu/vismod/classes/mas622-00/datasets/>

به تعداد کل چهره‌ها برای ارزیابی گزارش شده است. آزمایش‌ها برای سی بار اجرای تصادفی بر روی پایگاه داده Yale با مجموعه آموزشی شامل هفت تصویر از هر شخص و کاهش بعد به ده انجام شده است.



(شکل-۱): نمودار نرخ خطای دسته بندی (به درصد) در مقابل تعداد ویژگی‌های استخراج شده

در این جدول همچنین دو معیار ارزیابی دقت^۲ و فراخوانی^۳ برای یکبار اجرای تصادفی با همان تنظیمات قبلی گزارش شده است. این مقادیر از میانگین دقت و فراخوانی دسته‌های مختلف مسأله روی نمونه‌های آزمایشی، که در ادامه تعریف شده‌اند، محاسبه می‌شوند. دقت دسته نام نسبت تعداد نمونه‌هایی است که الگوریتم به درستی آن‌ها را به دسته نام منسوب کرده است، به تعداد کل نمونه‌هایی که الگوریتم آن‌ها را به دسته نام منسوب کرده است. فراخوانی برای دسته نام نیز نسبت تعداد نمونه‌هایی است که الگوریتم به درستی آن‌ها را به دسته نام منسوب کرده است، به تعداد کل نمونه‌هایی که به درستی متعلق به دسته نام هستند.

همچنین در جدول (۶) روش‌های مورد مقایسه از نظر زمان اجرا مورد بررسی قرار گرفته‌اند. زمان‌های اجرای به دست آمده برای کاهش بعد به ده بر روی پایگاه داده Yale و مجموعه آموزشی شامل سه تصویر از هر شخص بوده و بر حسب ثانیه گزارش شده‌اند.

از آن‌جا که روش پیشنهادی یک روش بانظارت است، با روش‌های کاهش بعد بانظارت LDA و SPPCA نیز مقایسه شده است. برای این منظور از پایگاه داده Yale استفاده شده، و کاهش بعد به ده بر روی نمونه داده‌ها انجام گرفته است. برای روش‌های مورد بررسی از مجموعه داده‌های آموزشی و آزمایشی یکسان استفاده شده است. نتایج این آزمایش‌ها در جدول شماره (۷) آمده است.

SPPCamm با جریمه نگاشت، که آن را به اختصار Proj-SPPCamm می‌نامیم.

در جدول (۱) تا (۴) در همه روش‌ها نمونه داده‌ها به فضایی با بعد ده کاهش داده شده‌اند. در آزمایش‌های انجام شده هر بار به صورت تصادفی تعدادی از نمونه‌ها برای آموزش در نظر گرفته شده و مابقی برای آزمایش استفاده می‌شوند. ما از مجموعه‌های آموزشی با ۳، ۵، و ۷ تصویر از هر شخص برای مقایسه روش‌ها استفاده کرده‌ایم. مقدار نخستین پارامترهای δ_x^2 و δ_y^2 در روش‌هایی که این پارامتر را دارند، برابر 10^{-5} قرار داده شده است. هر آزمایش سی بار و به صورت مستقل تکرار شده است. در آزمایش‌هایی که در آن‌ها از روش استخراج ویژگی استفاده شده است، چون ویژگی‌های استخراج شده دارای مقادیر عددی کوچک هستند، برای عملکرد بهتر روش‌ها به خصوص روش‌های PPCamm و SPPCamm مقادیر با ضرب در یک مقدار عددی مناسب تغییر مقیاس داده شده‌اند.

جزئیات مربوط به پیاده‌سازی جریمه نگاشت به صورتی که در (ژانگ و اشاید، ۲۰۱۰) بیان شده، در نظر گرفته شده است. در استفاده از SVM با به کارگیری جریمه نگاشت از هسته چندجمله‌ای درجه دو استفاده شده است. همچنین پارامتر C در آزمایش‌ها برابر 10^3 قرار داده شده است. در جدول (۱) و (۲) میانگین خطای دسته بندی برای هر روش بر روی پایگاه داده ORL گزارش شده است، و جدول (۳) و (۴) نتایج مربوط به پیاده‌سازی‌ها بر روی پایگاه داده Yale را نشان می‌دهند.

برای بررسی بیش تر روش‌های مورد مقایسه، به ازای تعداد ویژگی‌های استخراج شده مختلف، نه فقط ده ویژگی، آزمایش‌ها را تکرار کرده‌ایم. در آزمایش‌ها بر روی پایگاه داده Yale با مجموعه آموزشی شامل هفت تصویر از هر شخص، بعد داده به ۲، ۵، ۱۰، ۱۵، و ۲۰ کاهش داده شده، و متوسط نرخ خطای دسته بندی در مقابل تعداد ویژگی‌های استخراج شده برای سی بار اجرای تصادفی در قالب نمودار شکل (۱) رسم شده است. در این آزمایش‌ها خطای به دست آمده برای روش SPPCamm در بعدهای مختلف به خطاهای روش PPCamm بسیار نزدیک است. از این رو نمودار خطای این روش‌ها در این شکل بر روی هم قرار گرفته است.

در جدول (۵) معیار نرخ شناسایی درست^۱ به معنای نسبت تعداد چهره‌هایی که درست تشخیص داده می‌شوند،

² Precision

³ Recall

¹ Recognition rate

۵- تحلیل نتایج

همان‌طور که از نتایج به‌دست‌آمده از آزمایش‌ها مشاهده می‌شود، روش پیشنهادی با در نظر گرفتن جریمه نداشت Proj-SPPCamm (در مقایسه با سایر روش‌های کاهش بعد آزمایش‌شده، در تمام آزمایش‌ها خطای دسته‌بندی کمتری دارد. با مقایسه روش PPCamm با در نظر گرفتن جریمه نداشت و مدل بانظارت آن یعنی Proj-SPPCamm برتری روش Proj-SPPCamm به‌طور کامل مشهود است. این به بدان معناست که با استفاده از اطلاعات خروجی و اعمال جریمه نداشت، قادر به یادگیری بهتر مدل در مسائل بانظارت هستیم. به‌علاوه همان‌طور که از نتایج آزمایش‌ها مشاهده می‌شود، روش PPCamm در مقایسه با KPCA در بیش‌تر آزمایش‌ها خطای کمتری دارد. این امر می‌تواند نشان‌دهنده مناسب‌تر بودن روش‌های کاهش بعد محلی خطی نسبت به روش‌های کاهش بعد غیر خطی سراسری در ارائه مجموعه داده‌ها (به‌دلیل انعطاف‌پذیری بیش‌تر و ساده‌تر بودن این مدل‌ها) باشد. از آن‌جا که مدل بانظارت این روش، یعنی SPPCamm که در این مقاله ارائه شده است، هم دارای خاصیت محلی خطی است، و هم می‌تواند از اطلاعات مربوط به خروجی در مسائل پیش‌بینی و دسته‌بندی بهره‌گیر، برتری این مدل نسبت به دیگر روش‌های کاهش بعد خطی و غیرخطی موجه بوده، و این ویژگی‌ها امکان استفاده از آن در مسائل پیچیده و بانظارت را فراهم می‌کند.

مقایسه روش پیشنهادی با روش LDA در جدول (۷) نیز برتری روش پیشنهادی به‌خصوص برای تعداد داده‌های آموزشی کم را نشان می‌دهد؛ مقایسه نتایج روش SPPCamm و SPPCA برتری نسبی روش SPPCamm را نشان می‌دهد؛ اما همان‌طور که از جدول (۷) مشاهده می‌شود، نتایج بسیار به هم نزدیک هستند. این امر را می‌توان به دلیل کوچک بودن مجموعه آموزشی و یا در حالت کلی کوچک بودن مجموعه نمونه داده‌ها دانست. چون برای یک مجموعه داده کوچک چند مؤلفه در نظر گرفته می‌شود؛ درحالی که یک مؤلفه هم می‌تواند بازنمایی داده را به خوبی انجام دهد. اگر مجموعه داده بزرگ باشد استفاده از روش SPPCamm و در نظر گرفتن چند مؤلفه برای بازنمایی داده، کارایی را به نحو قابل توجهی می‌تواند بهبود بخشد.

همان‌طور که از نتایج به‌دست‌آمده از آزمایش‌ها مشاهده می‌شود، با افزایش تعداد داده‌های آموزشی متوسط خطای دسته‌بندی به‌طور تقریبی در همه روش‌ها کاهش

می‌یابد، که این ویژگی برای روش پیشنهادی نیز برقرار است. همچنین با توجه به نمودارهای شکل (۱) که خطای دسته‌بندی در مقابل تغییرات بعد را نشان می‌دهد، مشاهده می‌شود که با افزایش بعد، خطای دسته‌بندی کاهش می‌یابد. بررسی زمان اجرای روش‌های آزمایش‌شده در جدول (۶) نشان می‌دهد که زمان اجرا در روش پیشنهادی Proj-SPPCamm نسبت به سایر روش‌ها بیشتر است. بعد از این روش و روش SPPCamm زمان اجرای روش PPCamm و PPCamm زیاد است. در این روش‌ها قسمت بیش‌تر زمان اجرا مربوط به آموزش است؛ که دلیل آن استفاده از روش تکراری الگوریتم EM در آن‌ها است. زمان بالای آموزش را می‌توان ضعف روش پیشنهادی دانست؛ اما دقت به‌دست آمده از این روش به‌خصوص در بعدهای پایین قابل توجه است.

در روش پیشنهادی در بخش ۳-۱ روابط با فرض گاسی بودن توزیع داده به دست آمده است؛ لذا ممکن است این سؤال به ذهن آید که اگر توزیع داده غیر گاسی باشد آیا این روش همچنان عملکرد قابل قبولی خواهد داشت؟ همان‌طور که می‌دانیم در مسائل دنیای واقعی فرض گاسی بودن داده همیشه برقرار نیست. در مسایل بررسی شده در این مقاله، یعنی شناسایی چهره بر روی پایگاه داده Yale و ORL، نیز به الزام فرض گاسی بودن داده برقرار نیست. این در حالی است که نتایج آزمایش‌ها نشان می‌دهند روش پیشنهادی برای این مجموعه داده‌ها عملکرد خوبی داشته است؛ لذا مشاهده می‌شود که روش پیشنهادی برای داده غیر گاسی هم عملکرد قابل قبولی دارد.

به‌منظور کسب بینش شهودی در این زمینه سه توزیع یکنواخت در فضای سه‌بعدی را که دامنه آنها سه مکعب نزدیک به یکدیگر هستند، در نظر گرفته و از هر یک‌صد نمونه به تصادف برداشته‌ایم. بدین ترتیب چنان که در شکل (۲) نشان داده شده است، سیصد نمونه داده با بعد سه به‌دست آمده است. برای نمونه‌های تولیدشده از هر توزیع یک برچسب دسته در نظر گرفته‌ایم. بدین ترتیب نمونه‌ها از سه دسته متفاوت در نظر گرفته می‌شوند؛ سپس با استفاده از روش پیشنهادی SPPCamm با سه مؤلفه، نمونه داده‌ها را به فضای با بعد دو نگاشت کرده‌ایم. شکل (۳) نمونه‌های کاهش‌یافته بدین ترتیب را نشان می‌دهد. همان‌طور که در شکل (۳) مشاهده می‌شود، نمونه داده‌ها به‌صورت سه دسته به‌طور تقریبی جدا از هم در فضای

۶- جمع‌بندی

پیشنهادهایی به شرح زیر برای ارتقای روش پیشنهادی در این مقاله مطرح هستند. همان‌طور که بیان شد، در این مقاله از جریمه نگاشت با دسته‌بند SVM استفاده شده است. به کارگیری ایده جریمه نگاشت با استفاده از سایر روش‌های دسته‌بندی، به عنوان یک موضوع پژوهشی قابل بررسی است. همچنین در روش کاهش بعد بدون اتلاف به کاررفته در این مقاله از هسته چندجمله‌ای استفاده شده است. استفاده از هسته‌های دیگر و نیز یادگیری چندحسته‌ای در این مورد می‌تواند مفید باشد.

(جدول-۱): متوسط خطای دسته‌بندی روش‌های مختلف بر روی پایگاه داده ORL برای ۳۰ بار اجرای تصادفی

ORL Dataset - Classification error rate (in %) $\pm$ standard deviation (p -value)						
Training face s(#)	KPC A	PPC AM M	SPPC AMM	Proj-KPC A	Proj-PPC AM M	Proj-SPPC AMM
3	18.39 $\pm$ 3.27	18.01 $\pm$ 3.40	16.83 $\pm$ 2.88	15.50 $\pm$ 2.54	16.63 $\pm$ 3.47	10.37 $\pm$ 2.37
5	9.08 $\pm$ 2.64	8.85 $\pm$ 2.27	8.37 $\pm$ 2.69	7.77 $\pm$ 2.01	9.71 $\pm$ 2.40	4.25 $\pm$ 1.92
7	5.61 $\pm$ 1.61	5.31 $\pm$ 2.42	5.58 $\pm$ 3.85	4.22 $\pm$ 1.69	4.86 $\pm$ 2.10	2.14 $\pm$ 1.38

(جدول-۲): متوسط خطای دسته‌بندی روش‌های مختلف بر روی پایگاه داده ORL، استخراج ویژگی شده به روش گابور برای ۳۰ بار اجرای تصادفی

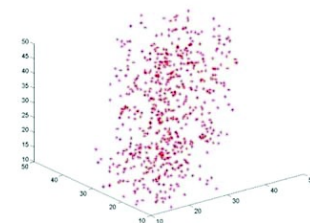
ORL Dataset - Gabor Feature Extraction- Classification error rate (in %) $\pm$ standard deviation (p -value)						
Training face s(#)	KPC A	PPC AM M	SPPC AMM	Proj-KPC A	Proj-PPC AM M	Proj-SPPC AMM
3	21.61 $\pm$ 2.56	17.99 $\pm$ 3.91	17.99 $\pm$ 3.91	18.07 $\pm$ 3.11	14.90 $\pm$ 3.16	7.23 $\pm$ 1.98
5	11.00 $\pm$ 2.23	9.91 $\pm$ 2.21	9.91 $\pm$ 2.21	9.32 $\pm$ 1.99	8.07 $\pm$ 1.76	1.96 $\pm$ 0.84
7	7.22 $\pm$ 1.54	5.95 $\pm$ 1.76	5.95 $\pm$ 1.76	6.31 $\pm$ 1.77	5.92 $\pm$ 1.70	1.06 $\pm$ 0.90

(جدول-۳): متوسط خطای دسته‌بندی روش‌های مختلف بر روی پایگاه داده Yale برای ۳۰ بار اجرای تصادفی

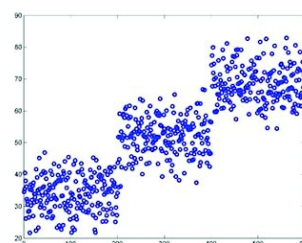
Yale Dataset - Classification error rate (in %) $\pm$ standard deviation (p -value)						
Training faces(#)	KPCA	PPCamm	SPPCamm	Proj-KPCA	Proj-PPCamm	Proj-SPPCamm
3	25.91 $\pm$ 3.85	26.52 $\pm$ 3.11	26.52 $\pm$ 3.11	25.96 $\pm$ 3.46	26.17 $\pm$ 2.63	19.50 $\pm$ 2.67
5	22.06 $\pm$ 4.10	21.62 $\pm$ 3.38	21.62 $\pm$ 3.38	21.97 $\pm$ 4.16	21.70 $\pm$ 2.85	16.27 $\pm$ 3.55
7	20.11 $\pm$ 3.63	18.62 $\pm$ 3.97	18.01 $\pm$ 4.27	21.39 $\pm$ 4.91	20.84 $\pm$ 3.55	15.67 $\pm$ 3.55

کاهش بعد یافته ظاهر شده‌اند، که بیان‌گر کارایی مناسب روش پیشنهادی برای داده با توزیع غیر گاسی است.

در این مقاله ابتدا مدل بانظارت روش ترکیبی تحلیل مؤلفه اصلی احتمالاتی (SPPCamm) ارائه شده است. این مدل توسعه‌ای از مدل ترکیبی تحلیل مؤلفه اصلی احتمالاتی است، که اطلاعات مربوط به برچسب داده را در نگاشت یادگیری شده منظور می‌کند. شکل (۴) روندنمای این مدل پیشنهادی را نشان می‌دهد؛ سپس این رویکرد بانظارت کاهش بعد محلی خطی SPPCamm در چارچوبی به نام جریمه نگاشت با استفاده از دسته‌بند SVM به کار گرفته شده است. جریمه نگاشت، این امکان را فراهم می‌آورد که از کاهش بعد برای بهبود دقت مدل پیش‌گو بدون خطر از دست دادن اطلاعات مرتبط استفاده کنیم. نتایج آزمایش‌ها نشان‌دهنده برتری روش پیشنهادی (SPPCamm) با در نظر گرفتن جریمه نگاشت (Proj-SPPCamm) در مقایسه با سایر روش‌های کاهش بعد بررسی شده در زمینه دقت دسته‌بندی است.



(شکل-۲): نمونه داده‌های تصادفی به دست آمده از سه توزیع یکنواخت نزدیک به هم



(شکل-۳): نسخه کاهش بعد یافته نمونه داده‌های تصادفی SPPCamm با روش (۲)

(جدول-۴): متوسط خطای دسته‌بندی روش‌های مختلف بر روی پایگاه داده Yale، استخراج ویژگی شده به روش گابور برای

۳۰ بار اجرای تصادفی

Yale Dataset – Gabor Feature Extraction- Classification error rate (in %) $\pm$ standard deviation (p -value)						
Training faces(#)	KPCA	PPCMM	SPPCMM	Proj-KPCA	Proj-PPCMM	Proj-SPPCMM
3	16.78 $\pm$ 2.89	14.17 $\pm$ 3.52	13.63 $\pm$ 3.11	14.97 $\pm$ 2.49	14.03 $\pm$ 2.98	10.46$\pm$2.65
5	13.11 $\pm$ 2.80	11.87 $\pm$ 3.13	11.85 $\pm$ 3.07	12.78 $\pm$ 2.69	12.07 $\pm$ 1.92	8.11$\pm$2.80
7	11.83 $\pm$ 3.51	10.97 $\pm$ 4.18	10.97 $\pm$ 4.18	11.17 $\pm$ 4.38	11.56 $\pm$ 3.06	7.74$\pm$3.41

(جدول-۵): بررسی معیارهای ارزیابی نرخ شناسایی درست، دقت، و فراخوانی بر روی پایگاه داده Yale و برای مجموعه آموزشی با ۷ تصویر

از هر شخص

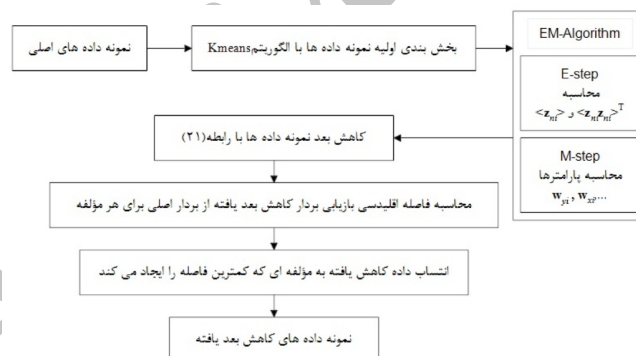
	KPCA	PPCMM	SPPCMM	Proj-KPCA	Proj-PPCMM	Proj-SPPCMM
Recognition rate	0.81	0.83	0.84	0.75	0.82	0.85
Precision	0.90	0.93	0.94	0.93	0.94	0.93
Recall	0.88	0.90	0.85	0.92	0.92	0.90

(جدول-۶): بررسی زمان آموزش و آزمایش روش‌های مورد مقایسه بر حسب ثانیه

	KPCA	PPCMM	SPPCMM	Proj-KPCA	Proj-PPCMM	Proj-SPPCMM
زمان آزمایش	0.0960	0.1968	0.4242	3.8218	3.6171	4.8276
زمان آموزش	0.8156	43.5106	56.1545	3.0556	46.4498	72.4550

(جدول-۷): متوسط خطای دسته‌بندی روش‌های با نظارت مختلف بر روی پایگاه داده Yale برای ۳۰ بار اجرای تصادفی

Yale Dataset - Classification error rate (in %) $\pm$ standard deviation (p -value)						
Training faces(#)	LDA	SPPCA	SPPCMM	Proj-LDA	Proj-SPPCA	Proj-SPPCMM
3	32.32 $\pm$ 3.21	23.58 $\pm$ 3.91	23.61 $\pm$ 3.46	19.38 $\pm$ 3.01	18.83 $\pm$ 2.78	18.83 $\pm$ 2.78
5	31.89 $\pm$ 4.90	17.81 $\pm$ 3.45	17.78 $\pm$ 3.48	15.63 $\pm$ 2.87	15.22 $\pm$ 2.66	15.22 $\pm$ 2.66
7	27.44 $\pm$ 3.48	17.56 $\pm$ 4.27	17.56 $\pm$ 4.27	14.78 $\pm$ 3.62	15.17 $\pm$ 3.73	15.07 $\pm$ 3.57



(شکل-۴): کارنمای روش پیشنهادی SPPCMM

Control, 2006. ICICIC'06. First International Conference on, 2006, pp. 345-348.

C.M. Bishop, Pattern Recognition and Machine Learning, Springer, 2006.

F. Han, T. Zhao, and H. Liu, "Coda: High dimensional copula discriminant analysis," Journal of Machine Learning Research, vol. 14, pp. 629-671, 2013.

۷- مراجع

C. Chen, J. Zhang, and R. Fleischer, "Distance approximating dimension reduction of Riemannian manifolds," Systems, Man, and Cybernetics, Part B: Cybernetics, IEEE Transactions on, vol. 40, pp. 208-217, 2010.

C.-G. Li and J. Guo, "Supervised isomap with explicit mapping," in Innovative Computing, Information and

Marginal Discriminant Embedding," *Procedia Engineering*, vol. 29, pp. 494-498, 2012.



سمیه احمدخانی کارشناسی خود را در رشته مهندسی کامپیوتر گرایش سخت‌افزار در سال ۱۳۸۹ از دانشگاه صنعتی همدان اخذ کرد. ایشان هم‌اکنون دانشجوی کارشناسی ارشد رشته مهندسی کامپیوتر گرایش هوش مصنوعی در دانشگاه اصفهان است. زمینه‌های مورد علاقه ایشان یادگیری ماشین، شناسایی الگو و پردازش تصویر می‌باشد.

نشانی رایانامه ایشان عبارت است از:
s.ahmadkhani@eng.ui.ac.ir



پیمان ادیبی مدرک کارشناسی خود را در رشته مهندسی کامپیوتر گرایش سخت‌افزار در سال ۱۳۷۶ از دانشگاه صنعتی اصفهان، و مدرک کارشناسی ارشد و دکترای خود را در رشته مهندسی کامپیوتر گرایش هوش مصنوعی از دانشگاه صنعتی امیرکبیر تهران به ترتیب در سال‌های ۱۳۸۰ و ۱۳۸۸ دریافت کرد. ایشان از سال ۱۳۸۹ تاکنون استادیار دانشکده مهندسی کامپیوتر دانشگاه اصفهان هستند. زمینه‌های پژوهشی مورد علاقه ایشان در حال حاضر شامل یادگیری ماشین و شناسایی الگو، هوش محاسباتی و محاسبات نرم و بینایی کامپیوتر است.

نشانی رایانامه ایشان عبارت است از:

adibi@eng.ui.ac.ir

J. Chen and Y. Liu, "Locally linear embedding: a survey," *Artificial Intelligence Review*, vol. 36, pp. 29-48, 2011.

K. Pearson, "LIII. On lines and planes of closest fit to systems of points in space," *The London, Edinburgh, and Dublin Philosophical Magazine and Journal of Science*, vol. 2, pp. 559-572, 1901.

L. Van der Maaten, E. Postma, and H. Van Den Herik, "Dimensionality reduction: A comparative review," *Journal of Machine Learning Research*, vol. 10, pp. 1-41, 2009.

M. E. Tipping and C. M. Bishop, "Mixtures of probabilistic principal component analyzers," *Neural computation*, vol. 11, pp. 443-482, 1999.

M. E. Tipping and C. M. Bishop, "Probabilistic principal component analysis," *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, vol. 61, pp. 611-622, 1999.
<http://vismod.media.mit.edu/vismod/classes/mas622-00/datasets/>

M. Haghighat, S. Zonouz, and M. Abdel-Mottaleb, "Identification Using Encrypted Biometrics," in *Computer Analysis of Images and Patterns*, 2013, pp. 440-448.

M. Zhu and A. M. Martinez, "Subclass discriminant analysis," *Pattern Analysis and Machine Intelligence, IEEE Transactions on*, vol. 28, pp. 1274-1286, 2006.

N. Gkalelis, V. Mezaris, I. Kompatsiaris, and T. Stathaki, "Mixture subclass discriminant analysis link to restricted Gaussian model and other generalizations," *Neural Networks and Learning Systems, IEEE Transactions on*, vol. 24, pp. 8-21, 2013.

R.A. Fisher, "The Statistical Utilization of Multiple Measurements," *Annals of Eugenics*, vol. 8, pp. 376-386, 1938.

S. Mika, G. Ratsch, J. Weston, B. Scholkopf, and K. Muller, "Fisher Discriminant Analysis with Kernels," *Proc. IEEE Signal Processing Soc. Workshop Neural Networks for Signal Processing IX*, pp. 41-48, 1999.

S. Yan, D. Xu, B. Zhang, H.-J. Zhang, Q. Yang, and S. Lin, "Graph embedding and extensions: a general framework for dimensionality reduction," *Pattern Analysis and Machine Intelligence, IEEE Transactions on*, vol. 29, pp. 40-51, 2007.

S. Yu, K. Yu, V. Tresp, H.-P. Kriegel, and M. Wu, "Supervised probabilistic principal component analysis," in *Proceedings of the 12th ACM SIGKDD international conference on Knowledge discovery and data mining*, 2006, pp. 464-473.

Y. Zhang and J. G. Schneider, "Projection penalties: Dimension reduction without loss," in *Proceedings of the 27th International Conference on Machine Learning (ICML-10)*, 2010, pp. 1223-1230.

Y.-D. Lan, H. Deng, and T. Chen, "Dimensionality Reduction Based on Neighborhood Preserving and