

طرح پژوهش برای دوره پسادکتری پژوهشگاه دانش‌های بنیادی

مرضیه اسماعیلی

مقدمه: تشخیص ناهنجاری یک مسئله چالش‌برانگیز در هوش مصنوعی می‌باشد. اکثر مدل‌های تشخیص ناهنجاری (AD) تنها با استفاده از نمونه‌های نرمال به روشی بدون نظارت آموزش داده می‌شوند، که ممکن است مرز بین داده نرمال و ناهنجار مبهم باشد و تمایز به درستی انجام نشود. در دنیای واقعی اغلب تعداد کمی از نمونه‌های ناهنجار در دسترس هستند. این مسئله در تصویربرداری پزشکی به دلیل اهمیت تشخیص ناهنجاری به عنوان یک نشانگر زیستی و یک داده حیاتی می‌تواند از اهمیت بالایی برخوردار باشد. علاوه بر آن، در اکثر موارد داده نرمال و ناهنجار با هم همپوشانی دارند و تمایز آنها و تشخیص مرز دقیق چالش‌برانگیز می‌باشد. همچنین، با موانعی چون کمیاب بودن تصاویر پزشکی برچسب‌گذاری شده، نامتوازن بودن مجموعه تصاویر پزشکی و حفظ محرمانگی داده‌های پزشکی نیز همراه است.

پژوهش‌های انجام شده: در پایان‌نامه دوره دکترای تخصصی انفورماتیک پزشکی در دانشگاه علوم پزشکی تهران و پژوهشکده علوم کامپیوتر پژوهشگاه دانش‌های بنیادی زیر نظر جناب آقای دکتر سبک‌رو، به موضوع تشخیص ناهنجاری در تصاویر پزشکی پرداخته شده است. این پژوهش با دو هدف اصلی طراحی و ایجاد شده است.

در هدف اول به بررسی عملکرد و قابلیت اطمینان روش‌های تشخیص ناهنجاری عمیق مانند روش‌های مبتنی بر شبکه‌های مولد رقابتی (GAN) و مبتنی بر بازنمایی ویژگی‌های عمیق بر روی تصاویر پزشکی پرداخته شد. این روش‌ها با استفاده از تصاویر نرمال و بدون نظر متخصص سعی می‌کنند ویژگی‌هایی را که سبب تمایز بیشتر تصاویر نرمال می‌شود را در قالب یک بردار پنهان (Latent vector) یاد بگیرند. این ویژگی سبب می‌شود این روش‌ها برای تشخیص ناهنجاری در تصاویر پزشکی مناسب باشند زیرا استخراج نشانگرهای زیستی در تصاویر پزشکی به صورت دستی پیچیده و چالش‌برانگیز است. با این وجود، همان‌طور که یافته‌های حاصل از آزمایش‌های انجام شده بر روی مجموعه تصاویر پزشکی متنوع نشان داد، آنچه سبب می‌شود این روش‌ها عملکرد قابل اطمینانی بر روی برخی مجموعه تصاویر نداشته باشند می‌توان ناشی از این مسئله دانست که ویژگی‌های متمایزکننده تصاویر نرمال و ناهنجار با فرضیه اولیه این روش‌ها که - ویژگی‌های داده ناهنجار فاصله قابل توجهی با ویژگی‌های داده نرمال دارد - مطابقت ندارد. همچنین، عواملی مانند تعداد کم نمونه‌های آموزشی، میزان ظرافت ناهنجاری، میزان پراکندگی ناهنجاری و روش تصویربرداری تأثیرگذار می‌باشند.

در هدف دوم بر اساس یافته‌های هدف اول، روشی برای تشخیص ناهنجاری در تصاویر پزشکی طراحی و توسعه داده شد که بتوان ویژگی‌های متمایزکننده از تصاویر نرمال را با دقت بالاتری استخراج کرد. برای این منظور، ابتدا با به کارگیری شبکه‌های مولد مبتنی بر انتشار (SDE) - که نسبت به GANها پایدارتر و تصاویر باکیفیت‌تری را تولید می‌کنند - نمونه‌های نزدیک به ناهنجار را تولید کردیم. به دلیل اینکه این مدل‌ها قابلیت توقف آموزش در مرحله‌ای قبل از تکمیل آموزش مدل برای تولید تصاویر نرمال را دارد بنابراین می‌توان مدلی را آموزش داد که تصاویری تولید کند که ناهنجاری‌هایی دارد. بدین ترتیب، برای هر مجموعه تصاویر پزشکی، مجموعه‌ای از تصاویر نزدیک به ناهنجار (جعلی) تولید کردیم. سپس، با داشتن مجموعه تصاویر نرمال و مجموعه تصاویر نزدیک به ناهنجار تولید شده (دو کلاس متوازن) با به کارگیری استخراج‌کننده ویژگی از پیش‌آموزش‌دیده بر روی تصاویر پزشکی (MedCLIP) می‌توان بردارهایی از ویژگی‌های متمایزکننده تصاویر نرمال استخراج و ذخیره کرد. در مرحله نهایی، امتیاز ناهنجاری برای هر تصویر با استفاده از روش k -نزدیک‌ترین همسایه (k -nn) محاسبه می‌شود و تصاویر ناهنجار بر اساس میزان فاصله تشخیص داده می‌شود.

نتایج ارزیابی روش ارائه شده با روش‌های پیشین نشان داد که به میزان قابل‌توجهی (۳ تا ۱۴ درصد افزایش سطح زیر نمودار ROC) تشخیص ناهنجاری در مجموعه تصاویر با تعداد نمونه کم و با ناهنجاری ظریف بهبود حاصل شد.

پژوهش‌های آتی: همان‌طور که نتایج پژوهش نشان داد روش‌های ارائه شده با وجود اینکه بر روی برخی مجموعه تصاویر با دقت بسیار بالا ناهنجاری را تشخیص داده‌اند اما بر روی برخی تصاویر نتایج قابل قبولی نداشته‌اند. این نشان از عدم تعمیم‌پذیری مدل‌ها دارد که یکی از موانع اصلی در به کارگیری آنها به خصوص در تصاویر پزشکی که تنوع و واریانس بالا می‌باشد.

همچنین، با توجه به همپوشانی تصاویر نرمال و ناهنجار در تصاویر پزشکی نیازمند روش‌هایی برای تعیین مرزهای دقیق و استخراج ویژگی‌های متمایزکننده می‌باشیم تا بتوان با دقت بالاتری موارد ناهنجار را تشخیص داد زیرا در پزشکی عدم تشخیص ناهنجاری می‌تواند پیامدهای حیاتی داشته باشد و از طرفی تشخیص ناهنجاری جدید می‌تواند منجر به یافته‌های جدید که ارزشمند است.

استفاده از گزارش‌های متنی پزشکی همراه با تصاویر پزشکی برای استخراج ویژگی‌ها با استفاده از استخراج‌کننده‌های تصویر-متن می‌تواند منجر به افزایش بهبود دقت تشخیص ناهنجاری شود.

تفسیر نتایج روش‌های تشخیص ناهنجاری مبتنی بر یادگیری عمیق می‌تواند کمک شایانی به طراحی و توسعه مدل‌های با قابلیت اعتماد بالاتر کند. همچنین، قابلیت پذیرش و قابلیت اعتماد این سیستم‌ها را به ویژه در

سیستم مراقبت بهداشتی بالاتر می‌برد. بنابراین، یکی از پژوهش‌های آتی به کارگیری روش‌های هوش مصنوعی تفسیرپذیر جهت توضیح و تفسیر چگونگی عملکرد و تصمیمات نهایی روش‌های تشخیص ناهنجاری می‌باشد. ارائه مدل تشخیص ناهنجاری تفسیرپذیر می‌تواند به افزایش قابلیت اطمینان، قابلیت پذیرش و اعتماد آن در محیط‌های حساس مانند سیستم مراقبت سلامت کمک کند.