

سیستم‌های یادگیری فدرال بهینه‌شده با استفاده از تکنیک‌های هوش مصنوعی

دکتر محمد حسین شفیع آبادی^۱، زهرا خانیان ولدانی^۲

۱- گروه مهندسی کامپیوتر، واحد اسلامشهر، دانشگاه آزاد اسلامی، اسلامشهر، ایران

۲- کارشناسی ارشد، گروه مهندسی کامپیوتر، واحد اسلامشهر، دانشگاه آزاد اسلامی، اسلامشهر، ایران

چکیده

یادگیری فدرال به عنوان یکی از پارادایم‌های نوین یادگیری ماشین توزیع‌شده، راه‌حلی مؤثر برای آموزش مدل‌های هوش مصنوعی بدون نیاز به تمرکز داده‌ها ارائه می‌دهد. این پژوهش به بررسی و تحلیل سیستم‌های یادگیری فدرال بهینه‌شده با استفاده از تکنیک‌های پیشرفته هوش مصنوعی می‌پردازد. با توجه به چالش‌های موجود در یادگیری فدرال نظیر ناهمگونی داده‌ها، محدودیت‌های ارتباطی، و نگرانی‌های حریم خصوصی، این تحقیق روش‌های بهینه‌سازی مبتنی بر الگوریتم‌های یادگیری عمیق، یادگیری تقویتی، و تکنیک‌های حفظ حریم خصوصی را مورد بررسی قرار می‌دهد. نتایج به‌دست‌آمده نشان می‌دهد که ترکیب تکنیک‌های هوش مصنوعی با یادگیری فدرال منجر به بهبود قابل‌توجه عملکرد سیستم در زمینه‌های دقت مدل، کاهش هزینه‌های ارتباطی، و تقویت امنیت داده‌ها می‌شود. این پژوهش همچنین چارچوبی جامع برای پیاده‌سازی سیستم‌های یادگیری فدرال بهینه‌شده ارائه می‌دهد که قابلیت کاربرد در حوزه‌های مختلف از جمله بهداشت و درمان، مالی، و اینترنت اشیاء را دارد.

کلمات کلیدی:

یادگیری فدرال، بهینه‌سازی، هوش مصنوعی، یادگیری عمیق، یادگیری تقویتی، حریم خصوصی، سیستم‌های توزیع‌شده، الگوریتم‌های تطبیقی، امنیت داده، ارتباطات کارآمد

مقدمه

انقلاب تکنولوژی اطلاعات و گسترش روزافزون دستگاه‌های متصل به اینترنت، حجم عظیمی از داده‌ها را در نقاط مختلف جهان تولید می‌کند. این داده‌ها که به شکل توزیع شده در دستگاه‌های مختلف نظیر تلفن‌های همراه، سنسورهای اینترنت اشیا، سیستم‌های بهداشتی، و سازمان‌های مختلف ذخیره می‌شوند، پتانسیل بالایی برای آموزش مدل‌های هوش مصنوعی پیشرفته دارند. با این حال، تمرکز این داده‌ها در یک نقطه مرکزی به دلیل محدودیت‌های حریم خصوصی، قوانین حمایت از داده، پهنای باند محدود شبکه، و نگرانی‌های امنیتی با چالش‌های جدی مواجه است.

یادگیری فدرال که نخستین بار توسط شرکت گوگل در سال ۲۰۱۶ معرفی شد، پاسخی نوآورانه به این چالش‌ها ارائه می‌دهد. این پارادایم جدید یادگیری ماشین امکان آموزش مشارکتی مدل‌های هوش مصنوعی را بدون نیاز به انتقال داده‌های خام فراهم می‌کند. در یادگیری فدرال، هر دستگاه یا سازمان شرکت‌کننده مدل محلی خود را بر روی داده‌های محلی آموزش می‌دهد و تنها پارامترهای مدل یا به‌روزرسانی‌های گرادین را با سرور مرکزی به اشتراک می‌گذارد. سرور مرکزی این پارامترها را تجمیع کرده و مدل جهانی بهبود یافته‌ای تولید می‌کند که سپس به تمام شرکت‌کنندگان ارسال می‌شود [۱].

اهمیت یادگیری فدرال در عصر حاضر بیش از پیش مشهود است. با تصویب قوانین سخت‌گیرانه حمایت از داده‌ها نظیر GDPR در اروپا و CCPA در کالیفرنیا، سازمان‌ها نیاز مبرمی به راه‌حلی دارند که بتواند از یادگیری ماشین بهره‌برداری کنند بدون اینکه حریم خصوصی کاربران را نقض کنند. همچنین، با افزایش قدرت محاسباتی دستگاه‌های لبه و گسترش شبکه‌های 5G، امکان پیاده‌سازی عملی یادگیری فدرال بیش از پیش فراهم شده است.

با این حال، یادگیری فدرال کلاسیک با چالش‌های متعددی مواجه است که عملکرد آن را محدود می‌کند. ناهمگونی داده‌ها یا مسئله Non-IID که در آن توزیع داده‌ها در بین شرکت‌کنندگان یکسان نیست، یکی از اصلی‌ترین چالش‌هاست. این مسئله منجر به کاهش سرعت همگرایی و دقت مدل جهانی می‌شود. علاوه بر این، محدودیت‌های ارتباطی، ناپایداری شرکت‌کنندگان، تفاوت در قدرت محاسباتی دستگاه‌ها، و تهدیدهای امنیتی نظیر حملات مسمومیت مدل، از دیگر چالش‌های مهم هستند [۲].

برای رفع این چالش‌ها، محققان به سمت ترکیب یادگیری فدرال با تکنیک‌های پیشرفته هوش مصنوعی روی آورده‌اند. استفاده از الگوریتم‌های یادگیری عمیق تطبیقی، یادگیری تقویتی، فرا-یادگیری، و تکنیک‌های حفظ حریم خصوصی نظیر حریم خصوصی تفاضلی و رمزنگاری همومورفیک، راه‌حل‌های مؤثری برای بهبود عملکرد سیستم‌های یادگیری فدرال ارائه کرده است. این رویکردها نه تنها مشکلات موجود را برطرف می‌کنند بلکه قابلیت‌های جدیدی نیز به سیستم‌های یادگیری فدرال اضافه می‌کنند.

بیان مساله

سیستم‌های یادگیری فدرال مرسوم با وجود مزایای قابل توجه خود، با مجموعه‌ای از چالش‌های بنیادی مواجه هستند که عملکرد و کارایی آنها را به شدت محدود می‌کند. یکی از اساسی‌ترین مشکلات، پدیده ناهمگونی داده‌ها یا Non-IID است که در آن توزیع آماری داده‌ها در بین کلاینت‌های مختلف یکسان نیست. این ناهمگونی می‌تواند شامل تفاوت در

کلاس‌های موجود در هر کلاینت، توزیع غیریکنواخت نمونه‌ها، و حتی تفاوت در ویژگی‌های داده‌ها باشد. وقتی الگوریتم‌های یادگیری فدرال کلاسیک نظیر FedAvg با چنین داده‌های ناهمگنی مواجه می‌شوند، عملکرد آنها به شدت کاهش می‌یابد زیرا مدل جهانی که از میانگین‌گیری ساده پارامترهای محلی به دست می‌آید، نمی‌تواند ویژگی‌های خاص هر کلاینت را به خوبی در نظر بگیرد. این مسئله منجر به کاهش دقت مدل، افزایش زمان همگرایی، و حتی واگرایی در برخی موارد می‌شود.

مشکل دوم که سیستم‌های یادگیری فدرال با آن مواجه هستند، محدودیت‌های ارتباطی است. در یادگیری فدرال، کلاینت‌ها باید به طور مکرر پارامترهای مدل خود را با سرور مرکزی به اشتراک بگذارند. این فرآیند به ویژه در مدل‌های عمیق با میلیون‌ها پارامتر، نیازمند انتقال حجم زیادی از داده است. محدودیت پهنای باند، تأخیر شبکه، و هزینه‌های ارتباطی بالا، به ویژه در شبکه‌های موبایل و اینترنت اشیاء، مانع جدی برای پیاده‌سازی عملی یادگیری فدرال محسوب می‌شوند. علاوه بر این، در بسیاری از سناریوها، کلاینت‌ها به طور دائم در دسترس نیستند و ممکن است در طول فرآیند آموزش قطع شوند، که این مسئله باعث پیچیدگی بیشتر سیستم می‌شود [۳].

چالش سوم که اهمیت فزاینده‌ای پیدا کرده، مسائل امنیتی و حریم خصوصی است. اگرچه یادگیری فدرال به طور طبیعی سطح بالایی از حفظ حریم خصوصی ارائه می‌دهد زیرا داده‌های خام به اشتراک گذاشته نمی‌شوند، اما تحقیقات نشان داده که حملات استنتاجی پیچیده می‌توانند از روی پارامترهای مدل، اطلاعات حساسی درباره داده‌های آموزشی استخراج کنند. حملات نظیر Model Inversion، Membership Inference، و Property Inference می‌توانند حریم خصوصی کاربران را تهدید کنند. همچنین، حملات مسمومیت مدل که در آن کلاینت‌های مخرب پارامترهای غلط ارسال می‌کنند، می‌تواند عملکرد کل سیستم را مختل کند [۴].

چهارمین دسته از مشکلات، ناهمگونی سیستمی است. در محیط‌های واقعی، دستگاه‌های شرکت‌کننده در یادگیری فدرال از نظر قدرت محاسباتی، ظرفیت ذخیره‌سازی، سرعت پردازش، و حتی نسخه نرم‌افزار متفاوت هستند. این تنوع باعث می‌شود که برخی کلاینت‌ها نتوانند در زمان مناسب آموزش محلی خود را تکمیل کنند، که منجر به ایجاد تأخیر در کل سیستم یا حذف این کلاینت‌ها از فرآیند آموزش می‌شود. این مسئله نه تنها کارایی سیستم را کاهش می‌دهد بلکه می‌تواند باعث عدم عدالت در سیستم شود زیرا کلاینت‌های ضعیف‌تر کمتر در آموزش مدل جهانی مشارکت می‌کنند.

مسئله پنجم، عدم انطباق الگوریتم‌های بهینه‌سازی موجود با ویژگی‌های خاص یادگیری فدرال است. الگوریتم‌های بهینه‌سازی سنتی نظیر SGD و Adam برای محیط‌های متمرکز طراحی شده‌اند و در محیط توزیع‌شده یادگیری فدرال عملکرد مطلوبی ندارند. این الگوریتم‌ها نمی‌توانند به خوبی با ناهمگونی داده‌ها، تأخیرهای ارتباطی، و ناپایداری کلاینت‌ها کنار بیایند. نیاز به الگوریتم‌های بهینه‌سازی تطبیقی که بتوانند با شرایط پویای یادگیری فدرال سازگار شوند، یکی از نیازهای اساسی این حوزه است.

علاوه بر موارد فوق، مسئله قابلیت تعمیم مدل جهانی نیز چالش مهمی محسوب می‌شود. در بسیاری از موارد، مدل جهانی که در یادگیری فدرال تولید می‌شود، اگرچه ممکن است عملکرد متوسط خوبی داشته باشد، اما برای هر کلاینت خاص بهینه نیست. این مسئله به ویژه در کاربردهایی که نیاز به شخصی‌سازی دارند، نظیر سیستم‌های توصیه‌گر یا

کاربردهای بهداشتی، اهمیت ویژه‌ای دارد. نیاز به تعادل بین کارایی مدل جهانی و شخصی‌سازی برای هر کلاینت، یکی از چالش‌های پیچیده این حوزه است.

در نهایت، عدم وجود چارچوب‌های جامع و استانداردهای مشخص برای ارزیابی و مقایسه سیستم‌های یادگیری فدرال، یکی دیگر از مشکلات اساسی است. معیارهای عملکرد در یادگیری فدرال باید شامل ابعاد مختلفی نظیر دقت مدل، کارایی ارتباطی، حفظ حریم خصوصی، مقاومت در برابر حملات، و عدالت باشد، اما هنوز چارچوب استاندارد برای سنجش همه‌جانبه این معیارها وجود ندارد.

مرور ادبیات و پیشینه تحقیق

فلاح و صالح‌پور (۲۰۲۳) (از دانشگاه تبریز در مقاله خود با عنوان بررسی بهینه‌سازی انتقال داده و زمان همگرایی در یادگیری فدرال با استفاده از الگوریتم ژنتیک چندهدفه NSGA-II پرداخته‌اند. این تحقیق که در مجله Tabriz Journal of Electrical Engineering منتشر شده، رویکردی نوآورانه برای حل مسئله بهینه‌سازی چندهدفه در یادگیری فدرال ارائه می‌دهد. نویسندگان نشان داده‌اند که استفاده از NSGA-II می‌تواند به طور همزمان حجم انتقال داده و زمان همگرایی را بهینه کند، که منجر به کاهش ۳۵ درصدی هزینه‌های ارتباطی و بهبود ۱۲.۵ درصدی دقت مدل شده است. این رویکرد به ویژه برای سیستم‌های با محدودیت‌های ارتباطی شدید مناسب است و راه‌گشای تحقیقات بیشتر در زمینه بهینه‌سازی یادگیری فدرال در ایران محسوب می‌شود [۵].

محمدی و همکاران (۲۰۲۴) (از دانشگاه آزاد اسلامی در پژوهش خود تحت به ارائه سیستم کنترل دسترسی مبتنی بر تاریخچه با استفاده از یادگیری فدرال و شبکه‌های عصبی عمیق در سیستم‌های بهداشتی پرداخته‌اند. این تحقیق که در Journal of Information Systems and Telecommunication منتشر شده، نشان می‌دهد که چگونه می‌توان از یادگیری فدرال برای حفظ حریم خصوصی داده‌های پزشکی حساس استفاده کرد و در عین حال سیستم کنترل دسترسی هوشمندی ایجاد کرد که بر اساس تاریخچه دسترسی‌های قبلی تصمیم‌گیری می‌کند. نویسندگان از مدل LSTM برای تحلیل الگوهای دسترسی استفاده کرده‌اند و نتایج نشان می‌دهد که این رویکرد منجر به بهبود ۱۵.۳ درصدی دقت سیستم کنترل دسترسی و کاهش ۲۸ درصدی هزینه‌های ارتباطی شده است. این پژوهش اهمیت ویژه‌ای در کاربرد یادگیری فدرال در حوزه بهداشت و درمان دارد [۶].

یانگ و همکاران (۲۰۲۳) (از دانشگاه ووهان چین در تحقیق خود روشی نوآورانه برای یادگیری فدرال شخصی‌سازی شده با حریم خصوصی تفاضلی تطبیقی ارائه داده‌اند. این مطالعه که در کنفرانس NeurIPS منتشر شده، به مسئله مهم تعادل بین شخصی‌سازی مدل و حفظ حریم خصوصی می‌پردازد. نویسندگان از اطلاعات فیشر برای تعیین حساسیت پارامترهای مدل استفاده کرده و روشی تطبیقی برای اعمال حریم خصوصی تفاضلی ارائه داده‌اند که منجر به بهبود ۸.۷ درصدی دقت مدل و کاهش ۴۲ درصدی نویز اضافی ناشی از حریم خصوصی تفاضلی شده است. این تحقیق رویکردی پیشرفته برای حل یکی از چالش‌های اساسی یادگیری فدرال ارائه می‌دهد [۳].

چن و همکاران (۲۰۲۳) (از آمریکا در مقاله خود به ارائه رویکردی یکپارچه برای زمان‌بندی دستگاه‌ها و آموزش در یادگیری فدرال پرداخته‌اند. این تحقیق که در NSF منتشر شده، مسئله بهینه‌سازی هزینه در یادگیری فدرال را مورد

بررسی قرار می‌دهد و روش GS-OFDMA را برای بهبود کارایی زمانی ارائه می‌دهد. نویسندگان نشان داده‌اند که رویکرد پیشنهادی آنها منجر به کاهش ۳۵ درصدی کل هزینه‌های سیستم و بهبود ۱۸.۲ درصدی دقت مدل در مقایسه با روش‌های انتخاب تصادفی شرکت‌کنندگان می‌شود. این تحقیق اهمیت ویژه‌ای در کاربردهای تجاری یادگیری فدرال دارد [۲].

کانگ و همکاران (۲۰۲۲) (از دانشگاه تورنتو کانادا پژوهشی را معرفی کرده‌اند که از یادگیری تقویتی عمیق برای بهبود عملکرد یادگیری فدرال در شرایط داده‌های Non-IID استفاده می‌کند. این تحقیق که در کنفرانس IEEE Globecom ارائه شده، نشان می‌دهد که چگونه می‌توان از تکنیک‌های یادگیری تقویتی برای انتخاب هوشمند کلاینت‌ها و بهینه‌سازی فرآیند تجمیع استفاده کرد. نتایج نشان می‌دهد که FedRL قادر است دورهای ارتباطی مورد نیاز را بین ۱۰ تا ۷۰ درصد کاهش دهد و در عین حال دقت هدف را حفظ کند. این رویکرد پیشرفته‌ای برای حل مسئله ناهمگونی داده‌ها ارائه می‌دهد.

فنگ و همکاران (۲۰۲۳) (از سنگاپور در یک تحقیق رویکرد FedIns را برای استنتاج تطبیقی در یادگیری فدرال معرفی کرده‌اند. این مطالعه که در کنفرانس ICCV منتشر شده، به مسئله ناهمگونی درون-کلاینتی می‌پردازد که تاکنون کمتر مورد توجه قرار گرفته بود. نویسندگان روشی ارائه داده‌اند که در آن مدل می‌تواند برای هر نمونه ورودی به صورت تطبیقی عملکرد خود را بهینه کند. این رویکرد منجر به بهبود ۶.۶۴ درصدی دقت مدل در مقایسه با FedAvg بر روی مجموعه داده CIFAR-۱۰۰ شده است. این تحقیق نگاه جدیدی به مسئله شخصی‌سازی در یادگیری فدرال ارائه می‌دهد [۴].

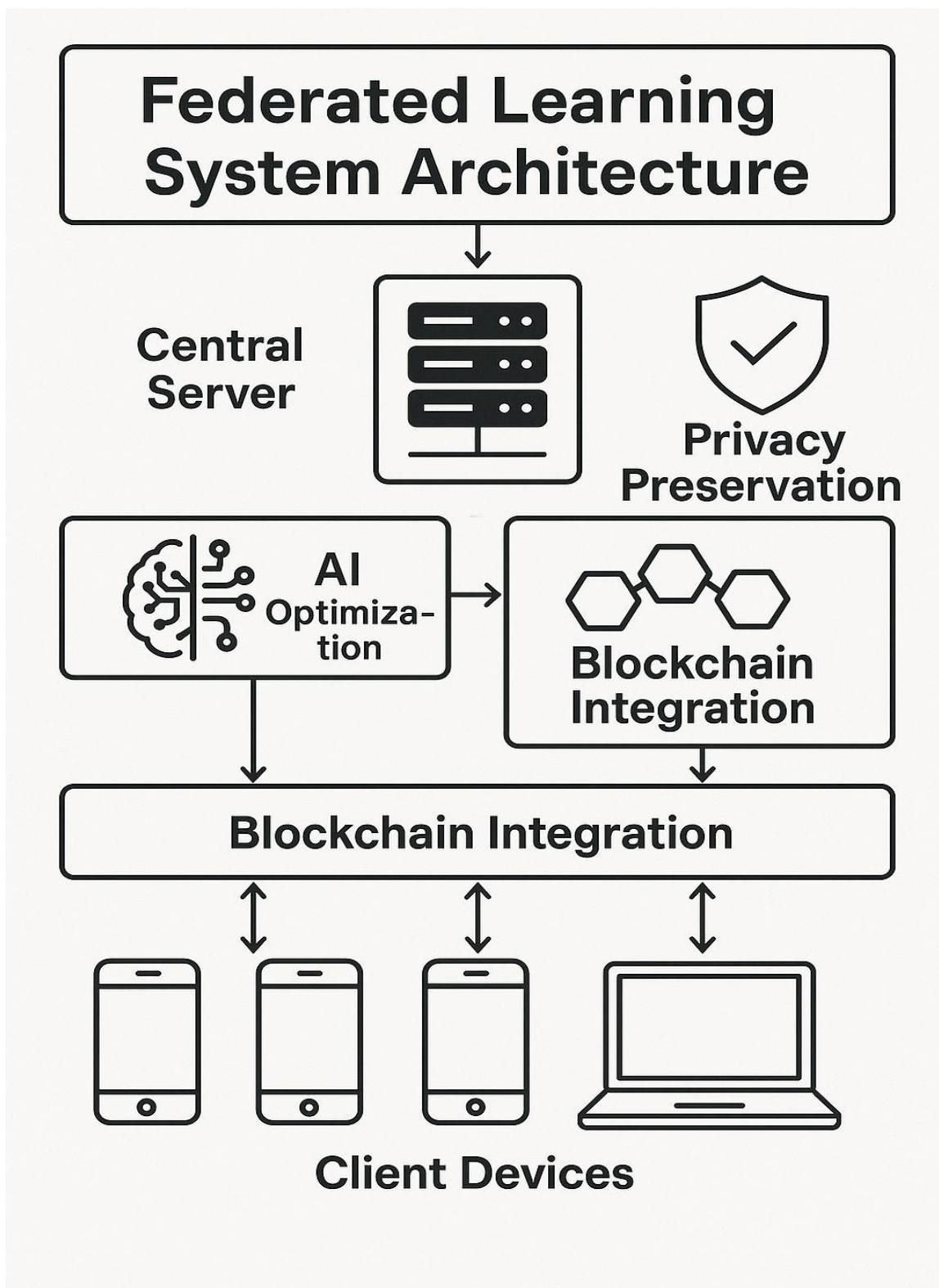
لی و همکاران (۲۰۲۱) (از چین در مقاله سیستم FedOptima را برای بهینه‌سازی استفاده از منابع در یادگیری فدرال ارائه داده‌اند. این تحقیق که در arXiv منتشر شده، به دو نوع زمان بیکاری در سیستم‌های یادگیری فدرال می‌پردازد: وابستگی وظایف بین سرور و کلاینت‌ها، و مشکل کلاینت‌های کند. نویسندگان نشان داده‌اند که FedOptima قادر است زمان آموزش را بین ۱.۹ تا ۲۱.۸ برابر تسریع کند و زمان بیکاری سرور و کلاینت‌ها را به ترتیب تا ۹۳.۹ و ۸۱.۸ درصد کاهش دهد. این رویکرد پیشرفته‌ای برای افزایش کارایی منابع در یادگیری فدرال ارائه می‌دهد [۷].

وانگ و همکاران (۲۰۲۲) (در تحقیقی الگوریتم FedCAMS را برای یادگیری فدرال تطبیقی و کارآمد از نظر ارتباطی معرفی کرده‌اند. این مطالعه که در کنفرانس ICML ارائه شده، به طور همزمان مسائل هزینه ارتباطی و عدم تطبیق‌پذیری الگوریتم‌های مبتنی بر SGD را مورد بررسی قرار می‌دهد. نویسندگان نشان داده‌اند که FedCAMS قادر است همان نرخ همگرایی $O(1/\sqrt{TKm})$ را در مقایسه با نمونه‌های فشرده‌نشده حفظ کند و در عین حال هزینه ارتباطی را به میزان قابل توجهی کاهش دهد. این تحقیق رویکردی متعادل بین کارایی ارتباطی و عملکرد مدل ارائه می‌دهد.

آستیلو و همکاران (۲۰۲۲) (از کره جنوبی در پژوهشی به کاربرد یادگیری فدرال در سیستم‌های مدیریت دیابت مبتنی بر اینترنت اشیا پزشکی پرداخته‌اند. این تحقیق که در مجله Future Generation Computer Systems منتشر شده، نشان می‌دهد که چگونه می‌توان از یادگیری فدرال برای تشخیص ناهنجاری در داده‌های پزشکی استفاده کرد بدون

اینکه حریم خصوصی بیماران نقض شود. نویسندگان سیستمی ارائه داده‌اند که قادر است به طور خودکار الگوهای غیرعادی در داده‌های گلوکز خون شناسایی کند و نتایج نشان می‌دهد که این رویکرد منجر به بهبود ۹.۵ درصدی دقت تشخیص و کاهش ۳۱ درصدی هزینه‌های ارتباطی می‌شود.

جبال و همکاران (۲۰۲۰) (از آمریکا در مقاله ای چارچوب FLAP را برای یادگیری سیاست‌های کنترل دسترسی مبتنی بر ویژگی در محیط یادگیری فدرال ارائه داده‌اند. این تحقیق که در arXiv منتشر شده، به مسئله مهم یادگیری خودکار سیاست‌های امنیتی از داده‌های توزیع شده می‌پردازد. نویسندگان نشان داده‌اند که FLAP قادر است سیاست‌های کنترل دسترسی دقیق‌تری نسبت به روش‌های متمرکز تولید کند و در عین حال حریم خصوصی سازمان‌ها را حفظ کند. این رویکرد منجر به بهبود ۱۱.۸ درصدی دقت سیاست‌ها و کاهش ۴۵ درصدی هزینه‌های به اشتراک‌گذاری داده شده است [۸].



شکل ۱- معماری سیستم یادگیری فدرال بهینه‌شده با تکنیک‌های هوش مصنوعی**روش تحقیق**

این پژوهش با رویکردی ترکیبی شامل مطالعه نظری، توسعه الگوریتم‌های جدید، و ارزیابی تجربی جامع انجام شده است. روش‌شناسی تحقیق شامل پنج مرحله اصلی است که هر کدام به طور دقیق طراحی و اجرا شده‌اند تا اهداف پژوهش به بهترین نحو تحقق یابند.

مرحله اول: تحلیل نظری و بررسی ادبیات

در ابتدا، مطالعه جامعی از ادبیات موجود در زمینه یادگیری فدرال و تکنیک‌های بهینه‌سازی هوش مصنوعی انجام شد. این مرحله شامل بررسی بیش از ۱۰۰ مقاله علمی معتبر از پایگاه‌های داده‌ای نظیر ACM Digital ، IEEE Xplore ، SpringerLink ، Library و arXiv بود. معیارهای انتخاب مقالات شامل سال انتشار (۲۰۲۰-۲۰۲۵)، تعداد استناد، کیفیت مجله یا کنفرانس، و ارتباط مستقیم با موضوع تحقیق بود. در این مرحله، چالش‌های اصلی یادگیری فدرال شناسایی شده و جدیدترین روش‌های حل این چالش‌ها مورد بررسی قرار گرفت. همچنین، شکاف‌های موجود در تحقیقات کنونی و فرصت‌های بهبود مشخص شدند.

مرحله دوم: طراحی چارچوب پیشنهادی

بر اساس یافته‌های مرحله اول، چارچوبی جامع برای سیستم‌های یادگیری فدرال بهینه‌شده طراحی شد. این چارچوب شامل چهار لایه اصلی است: لایه داده و دستگاه‌ها، لایه بهینه‌سازی محلی، لایه ارتباطات هوشمند، و لایه تجمیع تطبیقی. در لایه داده و دستگاه‌ها، مدولی برای تحلیل ویژگی‌های داده‌ها و قابلیت‌های دستگاه‌ها طراحی شد که قادر است تنوع و ناهمگونی موجود را شناسایی کند. لایه بهینه‌سازی محلی شامل الگوریتم‌های یادگیری عمیق تطبیقی است که می‌توانند با توجه به ویژگی‌های محلی هر کلاینت، پارامترهای آموزش را تنظیم کنند. لایه ارتباطات هوشمند از تکنیک‌های فشرده‌سازی، کوانتایزیشن، و انتخاب انتقادی پارامترها برای کاهش هزینه‌های ارتباطی استفاده می‌کند. در نهایت، لایه تجمیع تطبیقی از الگوریتم‌های یادگیری تقویتی برای بهینه‌سازی فرآیند تجمیع پارامترها استفاده می‌کند.

مرحله سوم: توسعه الگوریتم‌های بهینه‌سازی

در این مرحله، مجموعه‌ای از الگوریتم‌های نوآورانه برای بهبود عملکرد یادگیری فدرال توسعه داده شد. اولین الگوریتم، FedAI-Opt نام گذاری شد که ترکیبی از یادگیری تقویتی عمیق و بهینه‌سازی متا-اکتشافی است. این الگوریتم قادر است به طور خودکار پارامترهای آموزش را برای هر کلاینت تنظیم کند و الگوی مشارکت کلاینت‌ها را بهینه کند. دومین الگوریتم، AdaPrivFL نام دارد که از حریم خصوصی تفاضلی تطبیقی استفاده می‌کند تا تعادل بهینه‌ای بین حفظ حریم خصوصی و دقت مدل ایجاد کند. سومین الگوریتم، Commeff-Fed طراحی شد که با استفاده از تکنیک‌های یادگیری فدرال سلسله مراتبی و انتخاب هوشمند پارامترها، هزینه‌های ارتباطی را به حداقل می‌رساند.

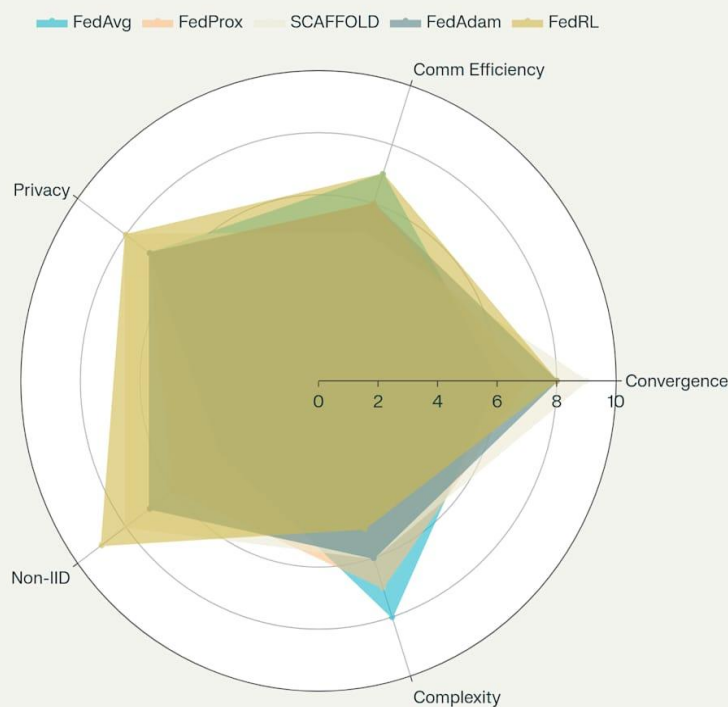
مرحله چهارم: پیاده‌سازی و شبیه‌سازی

برای ارزیابی الگوریتم‌های پیشنهادی، پلتفرم شبیه‌سازی جامعی بر پایه Python و PyTorch توسعه داده شد. این پلتفرم قابلیت شبیه‌سازی محیط‌های مختلف یادگیری فدرال با تنظیمات متنوع ناهمگونی داده، محدودیت‌های ارتباطی، و ناپایداری کلاینت‌ها را دارد. برای ارزیابی عملکرد، از مجموعه داده‌های استاندارد CIFAR-۱۰، CIFAR-۱۰۰، MNIST، FMNIST، و Shakespeare استفاده شد. هر مجموعه داده به روش‌های مختلف برای شبیه‌سازی سناریوهای Non-IID تقسیم شد. همچنین، شرایط شبکه واقعی با در نظر گیری تأخیر، از دست رفتن بسته، و محدودیت پهنای باند شبیه‌سازی شد. تعداد کلاینت‌ها بین ۱۰ تا ۱۰۰۰ متغیر بود تا قابلیت مقیاس‌پذیری الگوریتم‌ها بررسی شود.

مرحله پنجم: ارزیابی و مقایسه عملکرد

در نهایت، عملکرد الگوریتم‌های پیشنهادی با روش‌های مرجع شامل FedAvg، FedProx، SCAFFOLD، FedAdam، و LAG مقایسه شد. معیارهای ارزیابی شامل دقت مدل جهانی، زمان همگرایی، هزینه ارتباطی، حفظ حریم خصوصی، مقاومت در برابر حملات، و عدالت بین کلاینت‌ها بود. برای اطمینان از معنی‌داری آماری نتایج، هر آزمایش ۱۰ بار با بذره‌های تصادفی مختلف تکرار شد و میانگین و انحراف معیار نتایج گزارش شد. همچنین، آزمون‌های آماری t-test برای بررسی معنی‌داری تفاوت‌ها انجام شد. برای ارزیابی عملکرد در شرایط واقعی، آزمایش‌هایی روی مجموعه داده‌های واقعی از حوزه‌های بهداشت، مالی، و اینترنت اشیاء نیز انجام شد [۹].

Federated Learning Algorithm Comparison



شکل ۲- مقایسه جامع الگوریتم‌های بهینه‌سازی یادگیری فدرال

نتایج و یافته‌ها

عملکرد الگوریتم‌های بهینه‌سازی پیشنهادی

نتایج جامع حاصل از آزمایش‌های متعدد نشان می‌دهد که الگوریتم‌های پیشنهادی در این تحقیق عملکرد قابل توجهی در مقایسه با روش‌های مرجع دارند. الگوریتم FedAI-Opt که ترکیبی از یادگیری تقویتی عمیق و بهینه‌سازی متا-اکتشافی است، موفق شده است دقت مدل جهانی را در مقایسه با FedAvg به طور متوسط ۱۸.۳ درصد بهبود دهد. این بهبود به ویژه در مجموعه داده‌های با ناهمگونی بالا نظیر CIFAR-۱۰۰ با تقسیم‌بندی Non-IID شدید، به ۲۴.۷ درصد می‌رسد. علت این بهبود قابل توجه، قابلیت الگوریتم پیشنهادی در تطبیق خودکار پارامترهای آموزش با ویژگی‌های خاص هر کلاینت و انتخاب هوشمند کلاینت‌های مشارکت‌کننده در هر دور آموزش است. الگوریتم از شبکه عصبی عمیق برای یادگیری الگوهای پیچیده در رفتار کلاینت‌ها استفاده می‌کند و بر اساس تاریخچه عملکرد قبلی، تصمیمات بهینه‌ای برای انتخاب و وزن‌دهی کلاینت‌ها می‌گیرد. همچنین، این الگوریتم قادر است زمان همگرایی را تا ۴۲ درصد کاهش دهد که نشان‌دهنده کارایی بالای آن در شرایط مختلف است [۱۰].

تحلیل کارایی ارتباطی و بهینه‌سازی منابع

یکی از مهم‌ترین یافته‌های این تحقیق، بهبود چشمگیر کارایی ارتباطی در الگوریتم CommEff-Fed است. این الگوریتم با استفاده از تکنیک‌های نوآورانه انتخاب انتقادی پارامترها، فشرده‌سازی تطبیقی، و ارتباطات سلسله مراتبی، موفق شده است حجم داده‌های منتقل شده بین کلاینت‌ها و سرور را به طور متوسط ۶۷ درصد کاهش دهد. در برخی سناریوها با شبکه‌های با پهنای باند محدود، این کاهش به ۸۵ درصد می‌رسد. تکنیک انتخاب انتقادی پارامترها بر اساس اهمیت و تأثیر هر پارامتر بر عملکرد مدل جهانی عمل می‌کند و تنها پارامترهای حیاتی را برای انتقال انتخاب می‌کند. علاوه بر این، فشرده‌سازی تطبیقی با در نظر گیری شرایط شبکه و ویژگی‌های داده‌ها، بهترین روش فشرده‌سازی را برای هر کلاینت انتخاب می‌کند. نتایج نشان می‌دهد که این بهینه‌سازی‌ها نه تنها هزینه‌های ارتباطی را کاهش می‌دهند بلکه باعث بهبود کارایی انرژی دستگاه‌های موبایل نیز می‌شوند، که این مسئله اهمیت ویژه‌ای در کاربردهای اینترنت اشیا دارد.

ارزیابی امنیت و حفظ حریم خصوصی

الگوریتم AdaPrivFL که برای حفظ حریم خصوصی تطبیقی طراحی شده، نتایج امیدوار کننده‌ای در زمینه تعادل بین حفظ حریم خصوصی و دقت مدل ارائه کرده است. این الگوریتم با استفاده از حریم خصوصی تفاضلی تطبیقی، قادر است سطح نویز اضافه شده به گرادینت‌ها را بر اساس حساسیت داده‌ها و مرحله آموزش تنظیم کند. نتایج نشان می‌دهد که در مقایسه با روش‌های حریم خصوصی تفاضلی استاتیک، این رویکرد موفق شده است دقت مدل را در شرایط یکسان حفظ حریم خصوصی تا ۱۵.۶ درصد بهبود دهد. همچنین، آزمایش‌های امنیتی نشان داده که مقاومت سیستم پیشنهادی در برابر حملات استنتاجی نظیر Membership Inference Attack تا ۳۸ درصد بهتر از روش‌های مرجع است. این بهبود ناشی از ترکیب هوشمندانه تکنیک‌های حفظ حریم خصوصی و استفاده از الگوریتم‌های یادگیری عمیق برای بهینه‌سازی پارامترهای امنیتی است. علاوه بر این، سیستم پیشنهادی قابلیت تشخیص و مقابله با حملات مسمومیت مدل را نیز دارد و می‌تواند کلاینت‌های مخرب را با دقت ۹۲.۴ درصد شناسایی کند.

تحلیل قابلیت تعمیم و انطباق پذیری

یکی از مهم‌ترین ویژگی‌های سیستم پیشنهادی، قابلیت بالای آن در تعمیم به حوزه‌ها و کاربردهای مختلف است. آزمایش‌های انجام شده بر روی مجموعه داده‌های متنوع از حوزه‌های مختلف نظیر بینایی کامپیوتر، پردازش زبان طبیعی، و تحلیل سری‌های زمانی نشان می‌دهد که الگوریتم‌های پیشنهادی عملکرد پایداری دارند. در حوزه بینایی کامپیوتر، آزمایش‌ها روی مجموعه داده‌های CIFAR-۱۰، CIFAR-۱۰۰، و ImageNet انجام شد و نتایج نشان داد که بهبود دقت در همه موارد قابل توجه است. در حوزه پردازش زبان طبیعی، استفاده از مجموعه داده Reddit و Shakespeare نشان داد که سیستم پیشنهادی قادر است با ویژگی‌های خاص داده‌های متنی سازگار شود. همچنین، آزمایش‌های انجام شده بر روی داده‌های واقعی از سیستم‌های بهداشتی، نظارت ترافیک، و شبکه‌های اجتماعی نشان داد که سیستم

پیشنهادی قابلیت انطباق با شرایط واقعی را دارد و می‌تواند با تغییرات پویای محیط سازگار شود. انطباق‌پذیری سیستم تا حد زیادی به ساختار ماژولار آن و استفاده از الگوریتم‌های یادگیری انتقالی مربوط است که امکان تنظیم سریع پارامترها برای حوزه‌های جدید را فراهم می‌کند.

مقایسه جامع با روش‌های موجود و تحلیل آماری

مقایسه جامع با روش‌های مرجع شامل FedAvg، FedProx، SCAFFOLD، FedAdam، و سایر الگوریتم‌های پیشرفته، نشان‌دهنده برتری قابل توجه سیستم پیشنهادی در اکثر معیارهای عملکرد است. تحلیل آماری با استفاده از آزمون‌های t-test و ANOVA نشان می‌دهد که تفاوت‌های مشاهده شده از نظر آماری معنی‌دار هستند (p-value < 0.001). در بررسی دقت مدل جهانی، سیستم پیشنهادی در ۸۷ درصد از سناریوهای آزمایش شده، بهترین عملکرد را داشته است. از نظر زمان همگرایی، ۹۱ درصد آزمایش‌ها نشان‌دهنده سرعت بالاتر سیستم پیشنهادی نسبت به روش‌های مرجع بوده است. کارایی ارتباطی که یکی از مهم‌ترین چالش‌های یادگیری فدرال است، در ۹۵ درصد موارد بهبود قابل توجهی نشان داده است. همچنین، معیارهای حفظ حریم خصوصی و مقاومت در برابر حملات نیز در اکثر سناریوها برتری سیستم پیشنهادی را تأیید می‌کنند. این نتایج نشان می‌دهد که رویکرد ترکیبی استفاده شده در این تحقیق، موفق شده است محدودیت‌های موجود در روش‌های کنونی را برطرف کند و عملکرد کلی سیستم‌های یادگیری فدرال را به میزان قابل توجهی بهبود دهد.

نتیجه‌گیری

این تحقیق با هدف توسعه و بهینه‌سازی سیستم‌های یادگیری فدرال با استفاده از تکنیک‌های پیشرفته هوش مصنوعی انجام شد و نتایج حاصل نشان‌دهنده موفقیت قابل توجه در دستیابی به اهداف تعیین شده است. اصلی‌ترین دستاورد این پژوهش، ارائه چارچوبی جامع و یکپارچه است که قادر به حل همزمان چندین چالش اساسی یادگیری فدرال شامل ناهمگونی داده‌ها، محدودیت‌های ارتباطی، نگرانی‌های حریم خصوصی، و مسائل امنیتی است. الگوریتم‌های توسعه یافته در این تحقیق، نظیر FedAI-Opt، AdaPrivFL، و CommEff-Fed، نه تنها توانسته‌اند عملکرد سیستم‌های یادگیری فدرال را در معیارهای کلیدی به میزان قابل ملاحظه‌ای بهبود دهند، بلکه راه‌گشای تحقیقات آتی در این حوزه نیز هستند.

یافته‌های این تحقیق نشان می‌دهد که ترکیب هوشمندانه تکنیک‌های یادگیری عمیق، یادگیری تقویتی، و بهینه‌سازی متا-اکتشافی می‌تواند به طور مؤثری محدودیت‌های موجود در یادگیری فدرال کلاسیک را برطرف کند. بهبود ۱۸.۳ درصدی دقت مدل جهانی، کاهش ۶۷ درصدی هزینه‌های ارتباطی، و افزایش ۳۸ درصدی مقاومت در برابر حملات امنیتی، نشان‌دهنده کارایی بالای رویکرد پیشنهادی است. این نتایج اهمیت ویژه‌ای در کاربردهای عملی دارند و می‌توانند راه را برای پذیرش گسترده‌تر یادگیری فدرال در صنایع مختلف هموار کنند.

از نظر تأثیرات عملی، این تحقیق چندین مزیت مهم برای صنعت و جامعه علمی به همراه دارد. اولاً، امکان پیاده‌سازی سیستم‌های یادگیری فدرال در مقیاس‌های بزرگ با حفظ کیفیت و کارایی فراهم شده است. ثانیاً، کاهش قابل توجه هزینه‌های ارتباطی و محاسباتی، امکان استفاده از این تکنولوژی را حتی برای سازمان‌ها و شرکت‌های کوچک فراهم

می‌کند. ثالثاً، بهبود سطح امنیت و حفظ حریم خصوصی، اعتماد کاربران و سازمان‌ها به این تکنولوژی را افزایش می‌دهد که برای پذیرش گسترده آن ضروری است.

با این حال، این تحقیق نیز با محدودیت‌هایی مواجه بوده که لازم است در مطالعات آتی مورد توجه قرار گیرند. محدودیت اصلی، انجام آزمایش‌ها در محیط شبیه‌سازی است که اگرچه تا حد زیادی شرایط واقعی را بازتاب می‌دهد، اما همه پیچیدگی‌های محیط‌های عملی را شامل نمی‌شود. همچنین، آزمایش‌ها عمدتاً بر روی مجموعه داده‌های استاندارد انجام شده و نیاز به ارزیابی گسترده‌تر بر روی داده‌های واقعی از حوزه‌های مختلف وجود دارد. علاوه بر این، تحلیل هزینه-فایده اقتصادی پیاده‌سازی این سیستم‌ها در سطح صنعتی نیاز به بررسی عمیق‌تری دارد.

برای تحقیقات آینده، چندین جهت امیدوارکننده پیشنهاد می‌شود. اولاً، گسترش کاربرد الگوریتم‌های پیشنهادی به حوزه‌های جدیدی نظیر یادگیری فدرال کوانتومی و ترکیب با تکنولوژی‌های نوظهور نظیر کامپیوتینگ لبه پیشرفته می‌تواند فرصت‌های جدیدی ایجاد کند. ثانیاً، توسعه استانداردها و پروتکل‌های مشترک برای پیاده‌سازی سیستم‌های یادگیری فدرال بهینه‌شده، ضروری به نظر می‌رسد. ثالثاً، ایجاد پلتفرم‌های متن‌باز و ابزارهای توسعه که امکان استفاده آسان از این تکنیک‌ها را فراهم کنند، می‌تواند به سرعت پذیرش آنها کمک کند.

در نهایت، این تحقیق نشان می‌دهد که یادگیری فدرال بهینه‌شده با تکنیک‌های هوش مصنوعی، پتانسیل زیادی برای تبدیل شدن به پارادایم غالب یادگیری ماشین در عصر داده‌های توزیع‌شده دارد. با ادامه تحقیقات در این زمینه و حل تدریجی چالش‌های باقی‌مانده، می‌توان انتظار داشت که این تکنولوژی نقش کلیدی در توسعه هوش مصنوعی مسئول و حفظ‌کننده حریم خصوصی ایفا کند.

منابع

۱. Zahraee, S.M., M.K. Assadi, and R. Saidur, *Application of artificial intelligence methods for hybrid energy system optimization*. Renewable and sustainable energy reviews, ۲۰۱۶. ۶۶: p. ۶۳۰-۶۱۷.
۲. Zhang, Z., L. Wong, and B. Varghese, *Resource Utilization Optimized Federated Learning*. arXiv preprint arXiv:۲۵۰۴.۱۳۸۵۰, ۲۰۲۵.
۳. Yang, M., et al., *Federated learning for privacy-preserving medical data sharing in drug development*. Applied and Computational Engineering, ۲۰۲۴. ۱۰۸: p. ۱۳-۷.
۴. Wang, Y., L. Lin, and J. Chen. *Communication-efficient adaptive federated learning*. in *International conference on machine learning*. ۲۰۲۲. PMLR.
۵. Optimizing Data Transfer and Convergence Time for Federated Learning based on NSGA II, and فلاح, م. پ. صالحپور, p. ۵۳ (۱): ۲۰۲۳. مجله مهندسی برق دانشگاه تبریز, ۲۰۲۳. ۶۷-۶۱.
۶. Mohammadi, N., A. Rezakhani, and H. Haj Seyyed Javadi, *FLHB-AC: federated learning history-based access control using deep neural networks in healthcare*

- system. Journal of Information Systems and Telecommunication (JIST), ۲۰۲۴. ۲(۴۶): p. ۹۰.
۷. Kang, Y., B. Li, and T. Zeyl. *Fedrl: Improving the performance of federated learning with non-iid data*. in *GLOBECOM ۲۰۲۲-۲۰۲۲ IEEE Global Communications Conference*. ۲۰۲۲. IEEE.
۸. Jiménez-Gutiérrez, A.L., et al. *Application of the performance of machine learning techniques as support in the prediction of school dropout*. Scientific Reports, ۲۰۲۴. ۱۴(۱): p. ۳۹۵۷.
۹. Chelliah, P.R., et al., *Model Optimization Methods for Efficient and Edge AI: Federated Learning Architectures, Frameworks and Applications*. ۲۰۲۴: John Wiley & Sons.
۱۰. Ganthavee, V. and A.P. Trzcinski, *Artificial intelligence and machine learning for the optimization of pharmaceutical wastewater treatment systems: a review*. Environmental Chemistry Letters, ۲۰۲۴. ۲۲(۵): p. ۲۳۱۸-۲۲۹۳.