

بررسی تعارض میان توسعه مدل‌های هوش مصنوعی و الزامات حریم خصوصی داده‌ها با تأکید بر مقررات بین‌المللی حفاظت از داده

علیرضا براتی

کارشناس ارشد حقوق خصوصی، گروه حقوق خصوصی، واحد نورآباد ممسنی، دانشگاه آزاد نورآباد ممسنی، ایران

alireza.barati5396@iau.ir

چکیده

پیشرفت سریع هوش مصنوعی (AI) و به‌ویژه مدل‌های یادگیری عمیق، انقلابی در صنایع مختلف ایجاد کرده است. با این حال، توسعه این مدل‌ها مستلزم جمع‌آوری و پردازش حجم عظیمی از داده‌هاست که اغلب حاوی اطلاعات شخصی هستند. این امر منجر به تنش قابل توجهی میان ضرورت توسعه فناوری‌های هوش مصنوعی و الزامات حریم خصوصی داده‌ها شده است. این مقاله به بررسی نظام‌مند تعارضات میان توسعه مدل‌های هوش مصنوعی و حریم خصوصی داده‌ها می‌پردازد و نقش مقررات بین‌المللی حفاظت از داده را در تعدیل این تعارض تحلیل می‌کند. با استفاده از روش تحقیق توصیفی-تحلیلی و بررسی اسناد معتبر، یافته‌ها نشان می‌دهند که اگرچه مقرراتی مانند GDPR اتحادیه اروپا، CCPA کالیفرنیا و قوانین نوظهور در سایر کشورها چارچوب‌هایی برای حفاظت ارائه می‌دهند، اما چالش‌های عملی قابل توجهی در همسو کردن توسعه هوش مصنوعی با این مقررات وجود دارد. مقاله حاضر راهکارهای فنی (مانند تفاضل خصوصی، یادگیری فدرال) و حقوقی-سیاستی را برای کاهش این تعارض پیشنهاد می‌دهد و بر ضرورت همکاری بین‌المللی برای تدوین استانداردهای جهانی تأکید می‌کند.

کلمات کلیدی: هوش مصنوعی، حریم خصوصی داده‌ها، مقررات بین‌المللی حفاظت از داده، GDPR، یادگیری عمیق

۱. مقدمه

در عصر دیجیتال، داده‌ها به منبعی حیاتی برای نوآوری و رشد اقتصادی تبدیل شده‌اند. هوش مصنوعی، به‌عنوان یکی از پیشرفته‌ترین محصولات این عصر، برای آموزش و بهبود عملکرد خود به حجم انبوهی از داده‌ها وابسته است (al et Zhao, ۲۰۲۳). بسیاری از این داده‌ها، از تصاویر و متن‌های منتشرشده در شبکه‌های اجتماعی تا سوابق درمانی و تراکنش‌های مالی، ماهیتی شخصی دارند. این امر حریم خصوصی افراد را در معرض تهدیدات جدیدی قرار داده است. همزمان، نگرانی‌های عمومی و پاسخ مقنن به این نگرانی‌ها منجر به تصویب مقررات سخت‌گیرانه‌ای در زمینه حفاظت از داده‌ها در سطح جهان شده است. معروف‌ترین این مقررات، «آیین‌نامه عمومی حفاظت از داده» (GDPR) اتحادیه اروپا است که در سال ۲۰۱۸ به اجرا درآمد و تأثیر گسترده‌ای بر کسب‌وکارهای جهانی داشته است (Bussche dem Von & Voigt, ۲۰۱۷).

تعارض ذاتی میان توسعه مدل‌های هوش مصنوعی و حریم خصوصی از چند جهت قابل بررسی است: نخست، ماهیت داده‌محور هوش مصنوعی که نیازمند دسترسی به داده‌های گسترده و گاه حساس است. دوم، ویژگی‌های مدل‌های پیچیده مانند شبکه‌های عصبی عمیق که ممکن است به‌طور ناخواسته اطلاعات شخصی موجود در داده‌های آموزشی را به خاطر بسپارند و آن‌ها را افشا کنند (et Carlini, ۲۰۲۱). سوم، اهداف متفاوت ذی‌نفعان؛ توسعه‌دهندگان هوش مصنوعی به دنبال دسترسی آزاد به داده‌های باکیفیت هستند، در حالی که نهادهای نظارتی و جامعه مدنی بر حق افراد برای کنترل داده‌های شخصی خود تأکید دارند.

این مقاله با پذیرش این تعارض به عنوان یک واقعیت مسلم در عصر حاضر، در پی پاسخ به این پرسش اصلی است: «مقررات بین‌المللی حفاظت از داده چگونه می‌توانند تعارض میان توسعه مدل‌های هوش مصنوعی و حریم خصوصی داده‌ها را مدیریت کنند و چه خلأها و چالش‌هایی در این مسیر وجود دارد؟» برای پاسخ به این پرسش، ساختار مقاله به شرح زیر است: پس از مقدمه، در بخش دوم به مرور ادبیات و مبانی نظری پرداخته می‌شود. بخش سوم چارچوب‌های مقرراتی بین‌المللی را بررسی می‌کند. بخش چهارم، تحلیل تعارض و چالش‌های خاص را ارائه می‌دهد. بخش پنجم به راهکارهای فنی و سیاستی می‌پردازد. بخش ششم شامل بحث و نتیجه‌گیری نهایی است.

۲. مرور ادبیات و مبانی نظری

۱.۲. توسعه مدل‌های هوش مصنوعی و نیاز داده‌ای

هوش مصنوعی مدرن، به‌ویژه در حوزه یادگیری ماشین و یادگیری عمیق، بر پایه الگوریتم‌هایی است که با داده‌های عظیم آموزش می‌بینند. عملکرد مدل‌هایی مانند ترانسفورمرها در پردازش زبان طبیعی یا شبکه‌های پیچشی در بینایی کامپیوتر، مستقیماً با حجم و کیفیت داده‌های آموزشی مرتبط است (al et Goodfellow, ۲۰۱۶). برای مثال، مدل GPT-۴ شرکت OpenAI با استفاده از تریلیون‌ها کلمه از متن‌های موجود در اینترنت آموزش دیده است که طبیعتاً حاوی اطلاعات شخصی بی‌شماری است (OpenAI, ۲۰۲۳). این مدل‌ها نه تنها برای آموزش اولیه، بلکه برای تنظیم دقیق و ارزیابی نیز به داده‌های جدید نیاز دارند. این نیاز مداوم، یک موتور محرک برای جمع‌آوری داده در مقیاس کلان ایجاد کرده است.

۲.۲. مفهوم حریم خصوصی داده‌ها در عصر دیجیتال

حریم خصوصی داده‌ها به حق فرد برای کنترل چگونگی جمع‌آوری، استفاده، پردازش، افشا و نگهداری اطلاعات شخصی خود اشاره دارد. با دیجیتالی شدن زندگی، دامنه اطلاعات شخصی از نام و آدرس به داده‌های موقعیت مکانی، رفتارهای آنلاین، داده‌های بیومتریک و حتی الگوهای عصبی گسترش یافته است (Solove, ۲۰۲۱). نظریه پردازان حریم خصوصی استدلال می‌کنند که حفظ حریم خصوصی برای استقلال فردی، کرامت انسانی و عملکرد سالم دموکراسی ضروری است. نقض حریم خصوصی می‌تواند به تبعیض، نظارت گسترده، دستکاری در انتخاب‌ها و آسیب‌های روانی منجر شود.

۳.۲. چارچوب‌های اخلاقی برای هوش مصنوعی

پیش از ورود به مباحث حقوقی، چارچوب‌های اخلاقی متعددی برای توسعه مسئولانه هوش مصنوعی ارائه شده‌اند. اصولی مانند شفافیت، عدالت، عدم تبعیض، مسئولیت‌پذیری و احترام به حریم خصوصی، ستون‌های اصلی این چارچوب‌ها هستند (al et Jobin, ۲۰۱۹). برای مثال، «رهنمودهای اخلاقی برای هوش مصنوعی قابل اعتماد» منتشرشده توسط کمیسیون اروپا، حریم خصوصی را به عنوان یک عنصر کلیدی برمی‌شمارد. این اصول اخلاقی، اغلب مبنای شکل‌گیری مقررات حقوقی بعدی بوده‌اند.

۳. مقررات بین‌المللی حفاظت از داده: تحلیل مقایسه‌ای

در این بخش، مهم‌ترین مقررات بین‌المللی و منطقه‌ای حفاظت از داده که بر توسعه هوش مصنوعی تأثیر می‌گذارند، بررسی می‌شوند.

۱.۳. آیین‌نامه عمومی حفاظت از داده (GDPR) / اتحادیه اروپا

GDPR سنگ‌بنای مدرن قوانین حریم خصوصی است و بر اساس اصل «حفاظت از داده از طریق طراحی و پیش‌فرض» عمل می‌کند. اصول کلیدی آن که مستقیماً بر توسعه AI تأثیر می‌گذارند، عبارتند از:

- **قانونمندی، انصاف و شفافیت:** پردازش داده باید دارای مبانی قانونی (مانند رضایت آگاهانه، قرارداد، منافع مشروع) باشد.
- **محدودیت هدف:** داده‌ها تنها برای اهداف مشخص، صریح و مشروع جمع‌آوری می‌شوند.
- **حداقل سازی داده:** تنها داده‌های ضروری و مرتبط با هدف قابل پردازش هستند.
- **دقت:** داده‌ها باید دقیق و به‌روز باشند.
- **محدودیت ذخیره‌سازی:** داده‌ها تنها برای مدت لازم نگهداری می‌شوند.
- **حقوق داده‌اشخاص:** شامل حق دسترسی، حق اصلاح، حق فراموش شدن، حق محدود کردن پردازش، حق اعتراض و حق عدم تبعیض.
- **حفاظت از داده‌ها در زمان انتقال بین‌المللی.**

برای توسعه هوش مصنوعی، GDPR چالش‌های خاصی ایجاد می‌کند. تأکید بر «توضیح‌پذیری» تصمیمات خودکار می‌تواند با ماهیت «جعبه سیاه» برخی مدل‌های پیچیده یادگیری عمیق در تناقض باشد (Powles & Selbst, ۲۰۱۷). همچنین، حق فراموش شدن با نیاز مدل‌های از پیش‌آموزش‌دیده به حذف داده‌های خاص یک فرد بدون تأثیر بر عملکرد کلی مدل، در تعارض است.

۲.۳. قانون حریم خصوصی مصرف‌کننده کالیفرنیا (CCPA/CPRA)

قانون CCPA و نسخه تقویت‌شده آن CPRA، چارچوبی مشابه اما متمایز در ایالات متحده ارائه می‌دهند. تمرکز اصلی بر «حق دانستن» مصرف‌کننده درباره داده‌های جمع‌آوری‌شده، «حق حذف» و «حق عدم تبعیض» در صورت اعمال حقوق حریم خصوصی است. CPRA همچنین یک مقام حفاظت از داده جدید ایجاد کرده و پردازش داده‌های «حساس» را محدود می‌کند (Pardau, ۲۰۱۸). این قوانین برای شرکت‌های فناوری بزرگ واقع در کالیفرنیا که پیشگام توسعه AI هستند، بسیار مرتبط است.

۳.۳. سایر مقررات مهم

- **قانون حفاظت از اطلاعات شخصی (PIPL) چین:** که در سال ۲۰۲۱ به اجرا درآمد، ترکیبی از عناصر GDPR و ویژگی‌های خاص چین است. این قانون بر رضایت، حداقل‌سازی داده و حق انتقال تأکید دارد، اما در عین حال، پردازش داده برای امنیت ملی و منافع عمومی را تسهیل می‌کند (Creemers, ۲۰۲۱). این دوگانگی می‌تواند در توسعه AI توسط دولت و شرکت‌های چینی تأثیرگذار باشد.
- **قانون حفاظت از داده‌های شخصی هند:** پیش‌نویس این قانون که انتظار می‌رود به زودی تصویب شود، بر حاکمیت داده و ذخیره‌سازی محلی داده‌ها تأکید دارد.
- **کنوانسیون ۱۰۸+ شورای اروپا:** به عنوان تنها معاهده بین‌المللی الزام‌آور در زمینه حفاظت از داده، استاندارد جهانی را ترویج می‌دهد.

جدول ۱: مقایسه مقررات کلیدی حفاظت از داده و تأثیر آن‌ها بر توسعه هوش مصنوعی

چالش اصلی برای توسعه AI	حقوق مرتبط با AI	مبانی قانونمندی کلیدی	حوزه اجرایی	مقررات
تطبیق با اصول حداقل‌سازی داده و توضیح‌پذیری مدل‌های پیچیده	حق توضیح، حق فراموش شدن، حق اعتراض	رضایت، قرارداد، منافع مشروع	اتحادیه اروپا (و هر شرکتی که داده اروپایی پردازش کند)	GDPR
مدیریت درخواست‌های حذف در مقیاس بزرگ برای مدل‌های AI	حق حذف، حق عدم تبعیض	اطلاع‌رسانی شفاف و حق انتخاب مصرف‌کننده	کالیفرنیا، ایالات متحده	CPRA
همسویی با الزامات امنیت ملی و منافع عمومی تعریف‌شده توسط دولت	حق انتقال، حق حذف	رضایت، ضرورت برای قرارداد/قانون	جمهوری خلق چین	PIPL چین
اجرا و نظارت در یک اکوسیستم فناوری در حال رشد	مشابه GDPR	مشابه GDPR	برزیل	LGPD برزیل

۴. تحلیل تعارض و چالش‌های کلیدی

تعامل توسعه AI و قوانین حریم خصوصی، چندین چالش ساختاری و عملی ایجاد می‌کند.

۱.۴. چالش‌های ناشی از ویژگی‌های فنی AI

- **حافظه‌سازی در مدل‌ها (Memorization):** تحقیقات نشان داده‌اند که مدل‌های زبان بزرگ می‌توانند بخش‌هایی از داده‌های آموزشی حساس (مانند شماره تلفن یا آدرس ایمیل) را به خاطر بسپارند و در خروجی خود تکرار کنند (Carlini et al., ۲۰۲۱). این امر می‌تواند نقض آشکار اصول محرمانگی باشد، حتی اگر داده‌ها در ابتدا به صورت ناشناس جمع‌آوری شده باشند.
- **استنتاج اطلاعات حساس (Attacks Inference):** حمله‌کنندگان می‌توانند با پرس‌وژو از یک مدل آموزش‌دیده، در مورد ویژگی‌های حساس داده‌های آموزشی (مانند تشخیص بیماری یک فرد در یک مجموعه داده پزشکی) استنتاج کنند (al et Shokri, ۲۰۱۷).
- **ناشناس‌سازی و بازشناسایی:** تکنیک‌های سنتی ناشناس‌سازی اغلب در برابر الگوریتم‌های پیشرفته بازشناسایی آسیب‌پذیر هستند. هوش مصنوعی خود می‌تواند ابزاری قدرتمند برای شکستن ناشناس‌سازی باشد، که چرخه معیوبی ایجاد می‌کند (al et Rocher, ۲۰۱۹).

۲.۴. چالش‌های ناشی از الزامات حقوقی

- **تضاد با اصل حداقل‌سازی داده:** فرآیند توسعه AI، به‌ویژه در مرحله اکتشاف و بهینه‌سازی، اغلب نیازمند دسترسی به داده‌های گسترده و نامحدود است. محدود کردن داده‌ها به حداقل ضرورت ممکن، می‌تواند نوآوری و دقت مدل را محدود کند.
- **حق فراموش شدن در مقابل مدل‌های ثابت:** حذف داده‌های یک فرد از یک مدل آموزش‌دیده، برخلاف حذف از یک پایگاه داده ساده، فرآیندی فنی بسیار پیچیده است. اغلب نیاز به بازآموزی کامل مدل با هزینه‌های محاسباتی گزاف دارد (al et Ginart, ۲۰۱۹).
- **توضیح‌پذیری و شفافیت:** ماده ۲۲ GDPR حق فرد برای «عدم پذیرش تصمیم‌گیری منحصرأ خودکار» را به رسمیت می‌شناسد. حتی زمانی که چنین تصمیمی مجاز است، فرد حق دریافت «توضیح معنادار» دارد. تفسیر نحوه رسیدن یک مدل پیچیده به یک نتیجه خاص، یک مسئله تحقیقاتی باز در حوزه هوش مصنوعی قابل تفسیر (XAI) است.

۳.۴. چالش حاکمیتی و بین‌المللی

- **تکه‌تکه‌سازی نظارتی (Fragmentation Regulatory):** وجود قوانین متفاوت در مناطق مختلف، بار انطباق را برای شرکت‌های بین‌المللی توسعه‌دهنده AI افزایش می‌دهد و می‌تواند مانع همکاری‌های پژوهشی جهانی شود.
- **رقابت ژئوپلیتیکی:** رقابت بین ایالات متحده، چین و اروپا برای رهبری در هوش مصنوعی، می‌تواند مانع از شکل‌گیری یک استاندارد جهانی واحد شود. رویکرد آمریکایی مبتنی بر بخش خصوصی-محور و انعطاف‌پذیر، با رویکرد مبتنی بر حقوق بنیادین اروپا و رویکرد دولتی-هدایت‌شده چین در تعارض است (al et Roberts, ۲۰۲۱).

۵. راهکارها و راه‌حل‌های پیشنهادی

برای تعدیل تعارض بین توسعه AI و حریم خصوصی، مجموعه‌ای از راهکارهای فنی و سیاستی لازم است.

۱.۵. راهکارهای فنی حفظ حریم خصوصی

- **یادگیری فدرال (Learning Federated):** در این پارادایم، مدل به جای جمع‌آوری داده در یک سرور مرکزی، روی دستگاه‌های کاربران (مانند تلفن‌های همراه) آموزش می‌بیند و تنها به‌روزرسانی‌های مدل (گرادیان‌ها) به سرور ارسال می‌شوند. این امر داده‌های خام را در مبدا نگه می‌دارد (al et Kairouz, ۲۰۲۱).
- **تفاضل خصوصی (Privacy Differential):** یک چارچوب ریاضی قوی که با تزریق نویز حساب‌شده به داده‌ها یا فرآیندهای محاسباتی، تضمین می‌کند که خروجی تحلیل، اطلاعات قابل تشخیص درباره هیچ فرد خاصی را فاش نمی‌کند (Roth & Dwork, ۲۰۱۴). شرکت‌هایی مانند اپل و گوگل از این تکنیک در محصولات خود استفاده می‌کنند.
- **محاسبات امن چندجانبه (Computation Multi-Party Secure):** این تکنیک اجازه می‌دهد تا چندین طرف بدون افشای داده‌های ورودی خود، محاسباتی را بر روی داده‌های ترکیبی انجام دهند.
- **حذف مؤثر از مدل‌های یادگیری ماشین:** توسعه الگوریتم‌هایی که بتوانند تأثیر داده‌های یک فرد را از یک مدل آموزش‌دیده بدون بازآموزی کامل، به‌طور کارآمد حذف کنند.

۲.۵. راهکارهای سیاستی و حقوقی

- **حفظ حریم خصوصی از طریق طراحی (Design by Privacy):** ادغام اصول حریم خصوصی در تمام مراحل طراحی سیستم‌های AI، از مرحله مفهومی تا توسعه و استقرار.
- **توسعه چارچوب‌های حکمرانی داده:** ایجاد ساختارهای داخلی در سازمان‌ها برای نظارت بر انطباق، مانند انتصاب مسئولان حفاظت از داده (DPO).
- **توسعه استانداردها و گواهینامه‌های بین‌المللی:** ایجاد نشان‌های اعتماد (Trustmarks) برای سیستم‌های AI که با استانداردهای حریم خصوصی مطابقت دارند. سازمان‌هایی مانند ISO و IEEE در حال کار بر روی این استانداردها هستند.
- **همکاری‌های بین‌المللی:** ترویج گفت‌وگو بین مقامات نظارتی مناطق مختلف (مانند GDPR و CCPA) برای هماهنگی بیشتر و کاهش تکه‌تکه‌سازی.
- **تخصیص منابع نظارتی:** تقویت ظرفیت و تخصص مقامات نظارتی حفاظت از داده برای درک و نظارت بر پیچیدگی‌های فنی سیستم‌های AI.

۶. بحث و نتیجه‌گیری

توسعه مدل‌های هوش مصنوعی و حفاظت از حریم خصوصی داده‌ها در یک رابطه دیالکتیکی پیچیده قرار دارند. از یک سو، هوش مصنوعی برای پیشرفت به داده نیاز دارد و می‌تواند منافع اجتماعی عظیمی (در سلامت، آب‌وهوا، آموزش و ...) ایجاد کند. از سوی دیگر، حریم خصوصی یک حق بنیادین انسانی و ضرورتی برای حفظ اعتماد در جامعه دیجیتال است. مقررات بین‌المللی حفاظت از داده،

به‌ویژه GDPR، تلاشی مهم برای ایجاد تعادل و مهار قدرت فناوری‌های داده‌محور هستند. این مقررات با تعیین اصولی مانند حداقل‌سازی، شفافیت و پاسخگویی، چارچوبی الزام‌آور برای توسعه مسئولانه AI ایجاد کرده‌اند.

با این حال، همانطور که این مقاله نشان داد، تطبیق کامل توسعه AI با این مقررات با چالش‌های فنی عمیقی روبرو است. ماهیت آماری و ظرفیت حافظه‌سازی مدل‌های پیچیده، با منطق دترمینیستی و فردمحور قوانین حریم خصوصی در تنش است. راه‌حل این تعارض، نه در توقف توسعه AI یا تضعیف قوانین حفاظتی، بلکه در نوآوری است. از یک طرف، نوآوری‌های فنی مانند تفاضل خصوصی و یادگیری فدرال ابزارهای قدرتمندی برای محافظت از حریم خصوصی در حین استفاده از داده ارائه می‌دهند. از طرف دیگر، نیازمند نوآوری در حاکمیت، سیاست‌گذاری و همکاری بین‌المللی هستیم.

آینده مطلوب، توسعه «هوش مصنوعی با حریم خصوصی» است. تحقق این امر مستلزم تلاش مستمر پژوهشگران، مهندسان، حقوقدانان، مقامات نظارتی و سیاست‌گذاران است. پیشنهاد می‌شود تحقیقات آتی بر روی ارزیابی عملی تأثیر قوانین مختلف بر نوآوری AI، توسعه استانداردهای قابل‌اندازه‌گیری برای حریم خصوصی در مدل‌های AI و تحلیل پیامدهای اخلاقی-اجتماعی راهکارهای فنی متمرکز شوند. تنها از طریق یک رویکرد بین‌رشته‌ای و مشارکتی است که می‌توان از مزایای هوش مصنوعی بهره برد و در عین حال، کرامت و حقوق بنیادین افراد در عصر دیجیتال را حفظ کرد.

منابع

۱. Carlini, N., Tramer, F., Wallace, E., Jagielski, M., Herbert-Voss, A., Lee, K., ... & Raff, E. (2021). Extracting training data from large language models. *USENIX Security Symposium*, ۲۶۳۳-۲۶۵۰.
۲. Creemers, R. (2021). China's Personal Information Protection Law: A Primer. *DigiChina Project*.
۳. Dwork, C., & Roth, A. (2014). The algorithmic foundations of differential privacy. *Foundations and Trends® in Theoretical Computer Science*, ۹(۳-۴), (۲۱۱-۴۰۷).
۴. Ginart, A., Guan, M., Valiant, G., & Zou, J. Y. (2019). Making AI forget you: Data deletion in machine learning. *Advances in Neural Information Processing Systems*, ۳۲.
۵. Goodfellow, I., Bengio, Y., & Courville, A. (2016). *Deep learning*. MIT press.
۶. Jobin, A., Ienca, M., & Vayena, E. (2019). The global landscape of AI ethics guidelines. *Nature Machine Intelligence*, ۱(۱), ۳۸۹-۳۹۹.
۷. Kairouz, P., McMahan, H. B., Avent, B., Bellet, A., Bennis, M., Bhagoji, A. N., ... & Zhao, S. (2021). Advances and open problems in federated learning. *Foundations and Trends® in Machine Learning*, ۱۴(۱-۲), (۱-۲۱۰).

۸. OpenAI). ۲۰۲۳. (GPT-۴ Technical Report .*arXiv preprint arXiv:۲۳۰۳.۰۸۷۷۴*.
۹. Pardo, S. L. (۲۰۱۸). (The California Consumer Privacy Act :Towards a European-style privacy regime in the United States *Journal of Technology Law & Policy*, ۲۳, ۶۸.
۱۰. Roberts, H., Cowls, J., Morley, J., Taddeo, M., Wang, V., & Floridi, L. (۲۰۲۱). (The Chinese approach to artificial intelligence :An analysis of policy ,ethics ,and regulation *AI & Society*, ۳۶(۱), (۵۹-۷۷).
۱۱. Rocher, L., Hendrickx, J. M., & de Montjoye, Y. A. (۲۰۱۹). (Estimating the success of re-identifications in incomplete datasets using generative models *Nature Communications*, ۱۰(۱), (۳۰۶۹).
۱۲. Selbst, A. D., & Powles, J. (۲۰۱۷). (Meaningful information and the right to explanation . *International Data Privacy Law*, ۷(۴), (۲۳۳-۲۴۲).
۱۳. Shokri, R., Stronati, M., Song, C., & Shmatikov, V. (۲۰۱۷). (Membership inference attacks against machine learning models *IEEE Symposium on Security and Privacy*, ۳-۱۸.
۱۴. Solove, D. J. (۲۰۲۱). *The digital person :Technology and privacy in the information age* . NYU Press.
۱۵. Voigt, P., & Von dem Bussche, A. (۲۰۱۷). *(The EU General Data Protection Regulation)GDPR :A Practical Guide* .Springer.
۱۶. Zhao, Y., Li, J., & Zhang, Y. (۲۰۲۳). (Data-centric AI :Perspectives and challenges *ACM Computing Surveys*.