

دوفصلنامه پژوهش‌های مدیریت و فرماندهی نظامی  
دانشگاه افسری و تربیت پاسداری امام حسین (علیه السلام) - دانشگاه امیرالمؤمنین (علیه السلام) نرسا  
سال بیستم، شماره ۵۷، بهار و تابستان ۱۴۰۴، صص ۱۱۵-۱۳۸

## اصول و کاربرد هوش مصنوعی در شناسایی حملات N.B.C

مجتبی محمدی<sup>۱</sup>، حسن میرحاج<sup>۲</sup>، روح‌اله عبدالهی<sup>۳</sup>

تاریخ دریافت: ۱۴۰۳/۱۲/۱۱ تاریخ پذیرش: ۱۴۰۳/۰۶/۲۳

### چکیده

در دنیای امروز، استفاده از فناوری‌های نوین به سرعت در حال تغییر الگوهای کلاسیک جنگ و استراتژی‌های نظامی هستند. یکی از تأثیرگذارترین این فناوری‌ها، هوش مصنوعی (AI) است که به عنوان یک عامل کلیدی در تحول روش‌های نبرد و شناسایی مدرن در زمینه‌های هسته‌ای، بیولوژیکی و شیمیایی (N.B.C) شناخته می‌شود. یکی از مشکلات اساسی در شناسایی حملات جنگ نوین برای نیروهای نظامی، دشواری در شناسایی سریع و دقیق نوع تهدید است. در این حملات، تهدیدات ممکن است به سرعت و به شکل‌های مختلف منتشر شوند و بسیاری از مواد شیمیایی، بیولوژیکی و هسته‌ای به سرعت قابل شناسایی نیستند یا برای شناسایی نیاز به تجهیزات خاص و روش خاص دارند که ممکن است در شرایط جنگی به سرعت در دسترس نباشند. هوش مصنوعی نه تنها قادر است به تحلیل داده‌ها و شبیه‌سازی سناریوهای مختلف بپردازد، بلکه می‌تواند در تصمیم‌گیری‌های سریع و کارآمد در مواقع بحران‌های و حملات جنگ نوین نقش بسزایی ایفا کند. در این مقاله مروری به اصول و کاربرد هوش مصنوعی در راستای استفاده در شناسایی سریع و دقیق در حملات و یا خرابکاری نوین مورد بحث قرار می‌گیرد.

واژگان کلیدی: هوش مصنوعی، N.B.C، الگوریتم، AI

<sup>۱</sup> پژوهشگر، مجتبی دانش بنیان امام حسین (علیه السلام)، سازمان تحقیقات و جهاد خودکفایی

<sup>۲</sup> گروه جنگ نوین، دانشگاه امیرالمؤمنین (علیه السلام)

<sup>۳</sup> گروه جنگ نوین، دانشگاه امیرالمؤمنین (علیه السلام)

## مقدمه

اساس پاسخ به یک سانحه میکروبی، شیمیایی و یا پرتوی، مشابه هر سانحه‌ای است که در آن مواد زیان‌آور دخالت داشته باشد. بنابراین، لازم است شرایطی فراهم گردد تا خطرات احتمالی شناسایی شده و نیروهای امنیتی و کارکنان اضطراری از اقدامات لازم آگاه و صورت گیرد (IAEA, 2003). در یک سانحه، داشتن اطلاعاتی کافی و قابل اعتماد در ساعات اولیه سانحه برای مدیران و تیم‌های عملیاتی و پدافندی بسیار مهم است. همانطور که می‌دانیم آشکار کردن وجود آلودگی اهمیت بسزایی در اقدامات رزمی و رفع آلودگی برای هر یگانی دارد. استفاده از هوش مصنوعی در این زمینه می‌تواند شامل تکنیک‌های مختلفی انجام شود. در اینجا به بررسی هوش مصنوعی در شناسایی و مدیریت چنین تهدیداتی اشاره شده است:

- **شناسایی و پیش‌بینی حملات:** از الگوریتم‌های یادگیری ماشین و تحلیل داده برای شناسایی الگوها و پیش‌بینی وقوع حوادث مشابه می‌توان استفاده کرد.
- **تشخیص سریع آلودگی‌ها:** استفاده از سنسورها و فناوری‌های هوش مصنوعی برای تشخیص حضور مواد خطرناک در محیط و اعلام هشدار به مقامات مربوطه.
- **شبیه‌سازی و آموزش:** ایجاد شبیه‌سازهای واقعی‌تر برای آموزش نیروهای واکنش سریع و بهبود آمادگی در مقابل حوادث NBC.
- **تصمیم‌گیری و مدیریت بحران:** توسعه سیستم‌های پشتیبانی تصمیم‌گیری مبتنی بر هوش مصنوعی که می‌توانند اطلاعات مختلف را تجزیه و تحلیل کرده و بهترین راهکارها را پیشنهاد دهند.

بکارگیری از هوش مصنوعی در استراتژی‌های جنگ نوین به نظامیان این امکان را می‌دهد که تهدیدات را شناسایی کرده و به طور دقیق پیش‌بینی کنند. از شبیه‌سازی وضعیت‌های بحرانی گرفته تا تجزیه و تحلیل اطلاعات مربوط به تهدیدات هسته‌ای، بیولوژیکی و شیمیایی،

هوش مصنوعی می‌تواند به‌عنوان ابزاری مؤثر در کاهش خطرات و بهینه‌سازی عملیات‌های نظامی مورد استفاده قرار گیرد.

همچنین هوش مصنوعی به ارتقاء آموزش و آمادگی نیروها کمک می‌کند، با ارائه سناریوهای تمرینی متنوع و واقعی‌تر. بنابراین، این فناوری می‌تواند به عنوان یک خط دفاعی قوی در برابر تهدیدات پیچیده و متغیر عمل کند و به نیروهای مسلح توانایی واکنش سریع و مؤثر به موقعیت‌های بحرانی را بدهد. هوش مصنوعی نه تنها ابزاری برای خودکارسازی، بلکه یک تغییر دهنده بازی در حوزه استراتژی‌های نظامی و امنیتی است که به کشورها این امکان را می‌دهد تا با تهدیدات نوظهور به صورت چابک و هوشمندانه‌تر مواجه شوند.

امروزه یکی از مهم‌ترین زمینه‌های تحقیقاتی در اکثر نقاط جهان، بررسی و ارائه روش‌های مداخله برای جلوگیری از ایجاد خسارت‌های مختلف حاصل از حوادث N.B.C است (راسل، ۲۰۱۶). هوش مصنوعی (AI) که با نام انقلاب صنعتی ۴،۰ نیز شناخته می‌شود، گام‌های عمیقی در نوآوری‌های علمی و فناوری در زمینه‌های مختلف برداشته است. این می‌تواند تحولات قابل توجهی را در نحوه انجام فعالیت‌های غیرنظامی و عملیات نظامی ایجاد کند. هوش مصنوعی (AI) به توانایی ماشین‌ها برای تقلید از هوش انسانی و انجام وظایفی که معمولاً به هوش انسانی نیاز دارند، مانند یادگیری، حل مسئله، تصمیم‌گیری و درک زبان طبیعی اشاره دارد (شارما، ۲۰۲۱). شکل ۲ فناوری‌های هوش مصنوعی از جمله یادگیری ماشینی، پردازش زبان طبیعی، رباتیک و بینایی رایانه را نشان می‌دهد. یادگیری ماشینی زیرمجموعه‌ای از هوش مصنوعی است که شامل آموزش الگوریتم‌های رایانه‌ای برای یادگیری الگوهای موجود در داده‌ها و پیش‌بینی یا تصمیم‌گیری بر اساس داده‌ها می‌شود (آگیار و همکاران، ۲۰۲۱). یادگیری عمیق نوعی از یادگیری ماشینی است که از شبکه‌های عصبی با لایه‌های متعدد برای پردازش داده‌های پیچیده مانند تصاویر یا گفتار استفاده می‌کند

(هرشی، ۲۰۲۳). پردازش زبان طبیعی توانایی رایانها برای درک، تفسیر و تولید زبان انسانی از جمله گفتار و متن است (دولین و همکاران، ۲۰۱۹). بینایی رایانه‌ی توانایی رایانها برای تجزیه و تحلیل و تفسیر اطلاعات بصری مانند تصاویر و ویدئوها است (گامز و همکاران، ۲۰۲۲). ما در این مقاله به شناسایی و ارائه انواع داده‌های هوش مصنوعی که بتواند برای بکارگیری آن برای شناسایی N.B.C استفاده نمود، ارائه شده است.

قلب این تحول داده‌ها هستند که برای آموزش و آزمایش مدل‌های هوش مصنوعی ضروری هستند. مدل‌های هوش مصنوعی برای شناسایی الگوها و روندهایی که تشخیص آن‌ها با استفاده از روش‌های سنتی تجزیه و تحلیل داده‌ها دشوار است، به مجموعه داده‌های بزرگ متکی هستند. این به آنها اجازه می‌دهد تا بر اساس داده‌هایی که بر اساس آنها آموزش دیده اند، یاد بگیرند و پیش‌بینی کنند.

#### ۱. عوامل های N.B.C : امروزه استفاده از عوامل شیمیایی، بیولوژیک و هسته‌ای (N.B.C)

به منظور آسیب رساندن از نظر جسمی و روحی-روانی به جامعه برای موفقیت‌های سیاسی و یا اجتماعی دولت‌ها و گروه‌های تروریستی یک موضوع بسیار مهم تلقی می‌شود. این حملات می‌تواند به شکل‌های مختلف از قبیل هسته‌ای (خرابکاری، انتشار مواد رادیواکتیو و غیره)، بیولوژیکی (باکتری‌ها، ویروس‌ها)، شیمیایی (گازهای سمی) انتشار یابد. چون حتی اگر در حد وسیع استفاده نشود، تأثیر ابعاد روحی روانی این موضوع با توجه به شبکه‌ها و ارتباطات امروزی بسیار بیشتر از گذشته است. لذا شناخت راههای مقابله با این عوامل و آمادگی سریع برای مقابله با چنین پیشامدهایی امری ضروری به نظر می‌رسد. آزادسازی یا انتشار عمدی عوامل شیمیایی، میکروبی و هسته‌ای در قالب جنگ‌افزارهای نظامی همچون موشک و بمب با هدف آسیب و کشتار غیر نظامیان و نظامیان نسبت به جنگ‌افزارهای کلاسیک پتانسیل و توان بیشتری را برای

دولت‌ها یا گروه‌ها ایجاد می‌کند. دکترین‌های نظامی در سطوح مختلف (راهبرد، عملیات و تاکتیک) نقش بسیار مهمی را در توصیف و تبیین رویکرد نظامی کشورها در نحوه مدیریت و کنترل خرابکاری و بحران‌های N.B.C و همچنین تبیین بهترین شیوه‌های هدایت نیروها و رفع بحران ایفاء می‌کنند. آگاهی از نوع و نحوه این حملات امکان شناسایی و برنامه‌ریزی در هوش مصنوعی جهت مقابله با آنها را فراهم مینماید. از این رو تهیه برنامه‌های جامع که در شرایط مختلف قادر به پاسخگویی باشد و بتواند بحران رابه موقع مهار کند، بسیار حائز اهمیت است. در کشورهای جهان نیز استفاده از سلاح‌های N.B.C در چندین نقطه و در مواقع مختلف به وقوع پیوسته است. در برخی از معروف‌ترین موارد حملات N.B.C مورد استفاده شده است را اشاره می‌نماییم:

- **جنگ جهانی اول (۱۹۱۴-۱۹۱۸):** استفاده از گازهای شیمیایی در جنگ جهانی اول به عنوان یکی از اولین مورد استفاده از سلاح‌های شیمیایی در تاریخ ثبت شده است. گازهای مختلفی مانند کلر، فسژن و خردل برای حمله خطوط جبهه دشمن استفاده شده است.

- **جنگ ایران و عراق (۱۹۸۰-۱۹۸۸):** یکی از نمونه‌های مشهور حملات شیمیایی در سطح بین‌المللی، حملات شیمیایی در جنگ ایران و عراق است که در آن نیروهای عراقی از سلاح‌های شیمیایی علیه نیروهای ایرانی و غیر نظامیان استفاده کردند. این حملات شامل گازهای خردل، سارین بود.

- **حمله شیمیایی به حلبجه، عراق (۱۹۸۸):** یکی از دیگر حملات معروف در زمان رژیم صدام حسین و در جریان جنگ عراق و ایران بود. نیروهای عراقی از گازهای شیمیایی استفاده کردند و این حمله منجر به مرگ مردم غیر نظامی شد و آثار آن هنوز بر محیط زیست و مردم آن منطقه باقی مانده است.

- جنگ سوریه (۲۰۱۱-تاکنون): در جنگ داخلی سوریه، استفاده از سلاح سرجنگی هسته‌ای و گازهای شیمیایی بارها گزارش شده است، که منجر به مرگ صدها نفر و آلوده شدن بسیاری از غیر نظامیان با گاز مانند سارین شد.

۲. مدل‌سازی بحران N.B.C: مدل‌سازی بحران خرابکاری<sup>۱</sup>، به تحلیل و شبیه‌سازی سناریوهای خرابکاری در زمینه‌های مختلف، از جمله صنعتی، نظامی و امنیتی می‌پردازد. مدل‌سازی بحران‌های مرتبط با هوش مصنوعی نقش مهمی در تهدیدات N.B.C ایفا کند. این مدل‌سازی می‌تواند شامل استفاده از تکنیک‌های هوش مصنوعی، نظریه بازی‌ها و تحلیل داده برای پیش‌بینی و مدیریت بحران‌های ناشی از خرابکاری باشد. در زیر به برخی از جنبه‌های کلیدی این موضوع می‌پردازیم:

## ۲.۱. شناسایی و ارزیابی تهدیدات

- تحلیل آسیب‌پذیری: شناسایی نقاط ضعف در سیستم‌ها و فرآیندها که می‌توانند توسط خرابکاران هدف قرار گیرند.

- تجزیه و تحلیل اطلاعات: استفاده از داده‌های تاریخی و فعلی برای شناسایی الگوهای رفتار خرابکارانه.

## ۲.۲. مدل‌سازی سناریوها

- شبیه‌سازی سناریوهای مختلف: ایجاد مدل‌های قابل فهم برای سناریوهای محتمل خرابکاری، مانند نفوذ به سیستم‌های امنیتی، خرابی تجهیزات یا حملات سایبری.

- استفاده از نظریه بازی‌ها: تحلیل رفتار خرابکاران و مسئولان از طریق مدل‌های تصمیم‌گیری و پیش‌بینی واکنش‌ها.

---

<sup>۱</sup> Crisis Sabotage Modeling

### ۲,۳. توسعه استراتژی‌های پاسخ

- برنامه‌ریزی پاسخ به بحران: ایجاد طرح‌هایی برای مقابله با خرابکاری‌ها شامل شناسایی سریع، اولویت‌بندی واکنش‌ها، و جابه‌جایی منابع.
- آموزش و شبیه‌سازی: تمرین سناریوهای خرابکاری برای بهبود توانایی لشکرها و سازمان‌ها در پاسخ‌گویی به بحران.

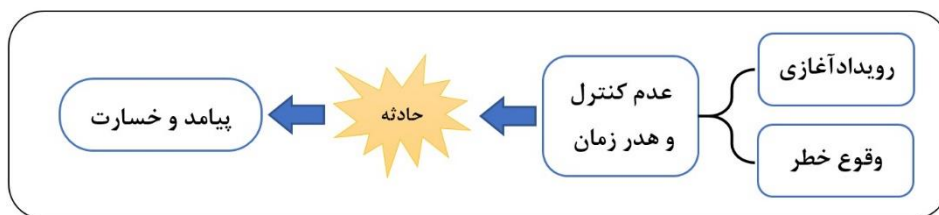
### ۲,۴. تحلیل و بهبود مستمر

- جمع‌آوری و تجزیه و تحلیل داده‌ها پس از بحران: ارزیابی عملکرد سیستم‌ها و بروزرسانی سیاست‌ها و استراتژی‌ها بر اساس تجربیات به دست آمده.
- بازنگری پروتکل‌ها: بهبود روش‌ها و فرایندها بر اساس داده‌های تازه برای کاهش خطرات آینده.

### ۲,۵. استفاده از فناوری‌های نوین

- هوش مصنوعی و یادگیری ماشین: برای پیش‌بینی رفتار خرابکارانه و شناسایی الگوهای غیرعادی در سیستم‌ها.
- سیستم‌های نظارتی و حسگرها: به افزایش آگاهی از وضعیت و شناسایی سریع تهدیدات.

بر طبق استاندارد OHSAS18001 ریسک به احتمال به وجود آمدن آسیب و یا خرابکاری ناشی از یک خطر گفته می‌شود؛ خطر عاملی است که در صورت وقوع پتانسیل ایجاد آسیب را دارد. خطر به عنوان منبع یا شرایط بالقوه ایجاد صدمه به انسان و آسیب به دارایی‌ها، محیط و یا ترکیبی از آنها است. به همین دلیل حادثه رویدادی است که منجر به خسارت و آسیب به محیط، جامعه و حتی به تاسیسات می‌گردد. در شکل ۱ زیر نحوه وقوع یک حادثه یا خرابکاری را نشان داده شده است.



شکل ۱: مکانیزم وقوع یک حادثه یا خرابکاری

۳. داده‌های هوش مصنوعی: داده‌ها برای هوش مصنوعی حیاتی هستند زیرا پایه‌ای هستند که الگوریتم‌های یادگیری ماشینی بر اساس آن یاد می‌گیرند، پیش‌بینی می‌کنند و عملکرد خود را در طول زمان بهبود می‌بخشند. برای آموزش یک مدل هوش مصنوعی، مقادیر زیادی داده مورد نیاز است تا مدل بتواند الگوها را تشخیص دهد، پیش‌بینی کند و عملکرد خود را در طول زمان بهبود بخشد (میکلف و همکاران، ۲۰۲۲). پارامترهای مهم و حیاتی که برای استفاده در تجهیزات شناسایی حملات عوامل N.B.C می‌توان به موارد زیر اشاره نمود:

### ۳.۱. رویکردهای یادگیری داده‌ها

الگوریتم‌های هوش مصنوعی برای یادگیری الگوها و پیش‌بینی یا تصمیم‌گیری بر اساس داده‌ها به داده نیاز دارند. تکنیک‌های یادگیری ماشینی هوش مصنوعی الگوریتم‌هایی هستند که به ماشین‌ها اجازه می‌دهند بدون برنامه‌نویسی صریح الگوها را یاد بگیرند و از داده‌ها پیش‌بینی کنند (ضاحیه و همکاران، ۲۰۲۲). این تکنیک‌ها به طور گسترده در کاربردهای مختلفی مانند پردازش زبان طبیعی، تشخیص تصویر و گفتار و سیستم‌های توصیه استفاده می‌شوند. به طور کلی، هرچه داده‌های بیشتری برای یادگیری الگوریتم هوش مصنوعی در دسترس باشد، پیش‌بینی‌ها یا تصمیم‌گیری‌های آن دقیق‌تر خواهد بود. چندین رویکرد یادگیری داده برای ساختن سیستم‌های هوش مصنوعی وجود دارد

(ماربی، ۲۰۱۸).

### ۳,۲. هوش مصنوعی داده محور

داده محوری و داده محور دو مفهوم مرتبط اما متمایز در دنیای تحلیل داده‌ها و تصمیم‌گیری هستند. با استفاده از داده‌ها، سازمان‌ها می‌توانند درک عمیق‌تری از عملیات، پایش و کنترل خود به دست آورند و بر اساس بینش‌های مبتنی بر داده، تصمیمات آگاهانه‌تری بگیرند. رویکردهای داده محور معمولاً در صنایعی مانند مالی، مراقبت‌های بهداشتی و خرده فروشی استفاده می‌شود، جایی که داده‌های دقیق و به‌موقع برای تصمیم‌گیری حیاتی هستند. به‌عنوان مثال، در پایش عوامل جنگ نوین، از رویکردهای داده محور برای تجزیه و تحلیل داده‌های بحران برای بهبود نتایج، شناسایی الگوهای نقاط و بهینه‌سازی برنامه‌های کنترل و شناسایی استفاده می‌شود. رویکردهای داده محور و داده محور دو رویکرد برای ساختن سیستم‌های هوش مصنوعی هستند که بر داده‌ها متکی هستند.

#### رویکرد داده محور:

این به رویکردی اشاره دارد که در آن داده‌ها کانون مرکزی یک سیستم یا فرآیند هستند. یک رویکرد داده محور شامل یک مدل نسبتاً ثابت است که جمع‌آوری، ذخیره‌سازی و تجزیه و تحلیل داده‌های باکیفیت را برای آموزش الگوریتم‌های هوش مصنوعی، بهبود عملکرد آنها و استفاده از داده‌ها برای اطلاع‌رسانی به فرآیندهای تصمیم‌گیری و حل مسئله در اولویت قرار می‌دهد. این رویکرد اغلب شامل استفاده از تجزیه و تحلیل‌های پیشرفته مانند یادگیری ماشین یا هوش مصنوعی برای کشف الگوها، روندها یا بینش‌هایی است که ممکن است بلافاصله از داده‌ها آشکار نشوند. رویکرد داده محور بر ایجاد زیرساخت‌های داده قوی و قابل اعتماد که می‌تواند طیف گسترده‌ای از برنامه‌های کاربردی هوش مصنوعی را پشتیبانی کند، تمرکز

دارد. هدف ایجاد یک مخزن داده متمرکز است که می‌تواند به عنوان یک منبع واحد از حقیقت برای همه برنامه‌های کاربردی هوش مصنوعی در یک سازمان عمل کند (جراحی، ۲۰۲۲). این رویکرد به ویژه زمانی مفید است که حجم زیادی از داده‌ها از منابع مختلف وجود داشته باشد، یا زمانی که داده‌ها پیچیده و کار با آنها دشوار است.

به طور کلی، یک رویکرد داده محور برای تصمیم‌گیری موثر هوش مصنوعی و حل مشکل ضروری است. با قرار دادن داده‌ها در مرکز سیستم هوش مصنوعی و پیروی از بهترین شیوه‌ها برای کیفیت، پردازش، حاکمیت و یکپارچگی داده‌ها، سازمان‌ها می‌توانند پتانسیل کامل هوش مصنوعی را باز کنند و نتایج بهتری را به دنبال داشته باشند. این رویکرد بر ساخت مدل‌های هوش مصنوعی متمرکز است که به طور خاص برای پیش‌بینی یا تصمیم‌گیری بر اساس داده‌ها طراحی شده‌اند. این رویکرد بر انتخاب، پردازش و تجزیه و تحلیل داده‌ها برای شناسایی الگوها، روابط و بینش‌هایی که می‌تواند برای بهبود دقت و عملکرد یک مدل هوش مصنوعی استفاده شود، تأکید می‌کند.

هدف توسعه یک مدل هوش مصنوعی بود که می‌تواند داده‌های جدید را بیاموزد و با آن‌ها سازگار شود، بدون اینکه توسط مجموعه‌ای از قوانین یا مفروضات از پیش تعریف‌شده محدود شود. این رویکرد به ویژه زمانی مفید است که داده‌ها نسبتاً همگن هستند یا زمانی که هدف خودکار کردن یک فرآیند تصمیم‌گیری خاص است (مازومدر، و همکاران، ۲۰۲۲). یک رویکرد مبتنی بر داده برای هوش مصنوعی شامل استفاده از داده‌ها به عنوان منبع اولیه اطلاعات برای آموزش و بهبود مدل‌های هوش مصنوعی است. در این رویکرد، سیستم هوش مصنوعی به جای اینکه توسط انسان برنامه‌ریزی شود، مستقیماً از داده‌ها یاد می‌گیرد. هوش مصنوعی مبتنی بر داده شامل چندین مرحله کلیدی است که در شکل ۳ نشان داده شده است.

### ۳,۳. آگاهی از موقعیت و داده‌ها

الگوریتم‌های هوش مصنوعی می‌توانند حجم وسیعی از داده‌های بلادرنگ را از منابع متعدد، از جمله رسانه‌های اجتماعی، شبکه‌های حسگر و تصاویر ماهواره‌ای تجزیه و تحلیل کنند تا دیدی جامع از بحران‌ها ارائه دهند. با ادغام داده‌ها از منابع مختلف، سیستم‌های هوش مصنوعی می‌توانند الگوها را شناسایی کنند، ناهنجاری‌ها را شناسایی کنند و گسترش و شدت بحران را پیش‌بینی کنند. این افزایش آگاهی موقعیتی، فرماندهان و مدیران را قادر می‌سازد تا انتخاب‌های آگاهانه داشته باشند و منابع را به طور کارآمد تخصیص دهند.

### ۳,۴. پاکسازی داده‌ها

پاکسازی داده‌ها یک گام اساسی در بهبود کیفیت داده‌ها برای هوش مصنوعی است. این شامل شناسایی و تصحیح خطاها، مقادیر از دست رفته، ناسازگاری‌ها و نقاط پرت در داده‌ها است. تکنیک‌هایی مانند پروفایل داده و اعتبارسنجی می‌تواند به شناسایی مناطقی از داده‌ها که نیاز به تمیز کردن دارند کمک کند.

### ۳,۵. پروفایل داده و آماده سازی داده‌ها

پروفایل داده شامل تجزیه و تحلیل مجموعه داده‌ها برای شناسایی مسائل کیفیت داده مانند داده‌های از دست رفته، سوابق تکراری و مقادیر متناقض است. آماده سازی داده شامل تمیز کردن و تبدیل داده‌های خام به قالب قابل استفاده برای الگوریتم‌های هوش مصنوعی است. این فرایندها برای اطمینان از اینکه داده‌های مورد استفاده برای آموزش مدل‌های هوش مصنوعی دقیق، کامل و سازگار هستند، ضروری هستند.

### ۳,۶. برچسب‌گذاری داده‌ها

برچسب‌گذاری داده‌ها شامل برچسب‌گذاری داده‌ها با ابرداده‌های مرتبط است که

ویژگی‌های آن را توصیف می‌کند، که می‌تواند به اطمینان حاصل شود که مدل‌های هوش مصنوعی با داده‌های با کیفیت بالا آموزش داده شده‌اند. به عنوان مثال، در تشخیص تصویر، برچسب گذاری داده‌ها ممکن است شامل شناسایی اشیاء در تصاویر و افزودن برچسب‌های توصیفی به داده‌ها باشد.

### ۳.۷. تکنیک‌های انتساب داده‌های گمشده

داده‌های از دست رفته یک چالش مهم در تضمین کیفیت داده‌های سیستم‌های هوش مصنوعی است. تکنیک‌های انتساب مختلفی برای رسیدگی به داده‌های از دست رفته پیشنهاد شده‌اند، از جمله انتساب میانگین، انتساب رگرسیون و انتساب چندگانه. تکنیک‌های پیشرفته، مانند روش‌های تکمیل ماتریس، در سال‌های اخیر مورد بررسی قرار گرفته است. هدف این روش‌ها ارائه تخمین‌های معقول از مقادیر از دست رفته و اطمینان از کامل بودن مجموعه داده است.

### ۳.۸. انتخاب ویژگی و مهندسی

انتخاب ویژگی و مهندسی نقش مهمی در پرداختن به چالش‌های کیفیت داده ایفا می‌کند، زیرا شامل شناسایی ویژگی‌های مرتبط و تبدیل داده‌های خام به قالبی مناسب برای تجزیه و تحلیل است. تکنیک‌هایی مانند حذف ویژگی‌های بازگشت (RFE)، LASSO و تجزیه و تحلیل اجزای اصلی را می‌توان برای کاهش ابعاد داده‌ها و حذف نویز به کار برد. علاوه بر این، دانش دامنه را می‌توان برای ایجاد ویژگی‌های جدید که الگوها و روابط زیربنایی را بهتر به تصویر می‌کشد، مورد استفاده قرار داد.

### ۳.۹. افزایش داده‌ها

تکنیک‌های تقویت داده‌ها را می‌توان برای رسیدگی به چالش‌های کیفیت داده مربوط به مجموعه داده‌های محدود یا نامتعادل استفاده کرد. این تکنیک‌ها نمونه‌های داده مصنوعی

را با اعمال تبدیل‌های مختلف مانند چرخش، مقیاس‌گذاری و چرخش به داده‌های اصلی تولید می‌کنند. تقویت داده‌ها به ویژه در بهبود عملکرد مدل‌های یادگیری عمیق در بینایی رایانه و وظایف پردازش زبان طبیعی موفق بوده است.

### ۳,۱۰. یادگیری فعال

یادگیری فعال رویکردی است که می‌تواند با هدایت فرآیند جمع‌آوری داده‌ها به رفع چالش‌های کیفیت داده‌ها کمک کند. در یادگیری فعال، مدل‌های هوش مصنوعی به طور مکرر آموزنده‌ترین نمونه‌ها را برای برچسب‌گذاری انتخاب می‌کنند و به مجموعه آموزشی اضافه می‌کنند، در نتیجه میزان داده‌های برچسب‌گذاری شده مورد نیاز را کاهش می‌دهند و عملکرد مدل را بهبود می‌بخشند. یادگیری فعال در کاربردهای مختلف، از جمله طبقه‌بندی متن و تشخیص اشیاء، امیدوار کننده است.

### ۳,۱۱. اعتبارسنجی و آزمایش داده‌ها

اعتبارسنجی داده‌ها شامل بررسی دقیق و کامل بودن داده‌ها است، در حالی که آزمایش شامل ارزیابی عملکرد مدل‌های هوش مصنوعی با استفاده از معیارهای مختلف است. این فرآیندها می‌توانند به شناسایی و رسیدگی به مسائل مربوط به کیفیت داده‌ها کمک کنند که ممکن است بر دقت و اثربخشی مدل‌های هوش مصنوعی تأثیر بگذارد.

### ۳,۱۲. کاهش سوگیری داده‌ها

سوگیری داده‌ها می‌تواند منجر به پیش‌بینی‌ها و تصمیمات هوش مصنوعی نادرست یا ناعادلانه شود. کاهش سوگیری داده‌ها شامل شناسایی و رسیدگی به سوگیری در داده‌های آموزشی از طریق تکنیک‌هایی مانند متعادل سازی داده‌ها است که شامل تنظیم توزیع داده‌ها برای کاهش سوگیری است.

### ۳،۱۳. نظارت و نگهداری مستمر

مدل‌های هوش مصنوعی باید به طور مداوم نظارت و نگهداری شوند تا اطمینان حاصل شود که دقیق و مؤثر باقی می‌مانند. این شامل بررسی‌های مداوم کیفیت داده‌ها، به‌روزرسانی مدل‌ها با داده‌های جدید و بازآموزی مدل‌ها در صورت نیاز است.

### ۳،۱۴. عناصر چالش برانگیز حجم داده‌ها

چالش حجم داده یکی از جنبه‌های کلیدی تحقیق و کاربرد هوش مصنوعی است. مجموعه داده‌های بزرگ برای آموزش مدل‌های هوش مصنوعی حیاتی هستند و همچنان در اندازه و پیچیدگی رشد می‌کنند. این رشد چالش‌هایی را به همراه دارد که باید برای اطمینان از استفاده مؤثر از هوش مصنوعی در حوزه‌های مختلف مورد توجه قرار گیرد. عناصر چالش برانگیز حجم داده در شکل ۸ ارائه شده است.



شکل ۲: عناصر چالش برانگیز حجم داده‌ها.

### ۳،۱۵. چالش‌های داده

چالش داده‌ها مثل یک شمشیر دو لبه. رشد نمایی داده‌ها نیروی محرکه موفقیت هوش مصنوعی به ویژه در تکنیک‌های یادگیری عمیق است. با این حال، حجم انبوه داده چالش‌های متعددی از جمله در ذخیره‌سازی، پردازش و مدیریت داده‌ها را به همراه دارد.

### ۳،۱۶. چالش‌های ذخیره سازی

حجم عظیمی از داده‌های تولید شده امروزه نیازمند راه حل‌های ذخیره سازی کارآمدتری برای پشتیبانی از برنامه‌های کاربردی هوش مصنوعی است. معماری‌های ذخیره‌سازی سنتی ممکن است نتوانند مقیاس‌پذیری، عملکرد و هزینه‌های مورد نیاز حجم کاری هوش مصنوعی را برآورده کنند. فن‌آوری‌های ذخیره‌سازی جدید، مانند حافظه غیرفرار (NVM) و سیستم‌های ذخیره‌سازی توزیع‌شده، به عنوان راه‌حل‌های ممکن پیشنهاد شده‌اند.

### ۳،۱۷. چالش‌های پردازش

مدل‌های هوش مصنوعی، به ویژه الگوریتم‌های یادگیری عمیق، به منابع محاسباتی عظیمی برای پردازش مجموعه داده‌های بزرگ نیاز دارند (یابلونسکی، ۲۰۱۹). این امر منجر به افزایش نیاز به سخت افزارهای تخصصی مانند GPU و TPU برای تسریع آموزش و استنتاج هوش مصنوعی شده است. علاوه بر این، تکنیک‌های جدیدی مانند فشرده‌سازی مدل، هرس و کوانتیزه کردن برای بهینه‌سازی مدل‌های هوش مصنوعی برای پردازش کارآمدتر مورد بررسی قرار گرفته‌اند.

### ۳،۱۸. چالش‌های مدیریت داده

از منظر حجم کلان داده، مدیریت موثر داده برای سیستم‌های هوش مصنوعی برای مدیریت حجم عظیمی از داده‌ها حیاتی است. این شامل پاکسازی داده‌ها، پیش پردازش،

برچسب گذاری و مراقبت می‌شود. تکنیک‌هایی مانند یادگیری فعال، نظارت ضعیف و یادگیری انتقالی برای کاهش بار حاشیه نویسی دستی داده‌ها پیشنهاد شده است.

### ۳,۱۹. چالش‌های حفظ حریم خصوصی داده‌ها در هوش مصنوعی

جمع‌آوری و اشتراک‌گذاری داده‌ها: سیستم‌های هوش مصنوعی به مقادیر زیادی داده برای آموزش مؤثر نیاز دارند، که اغلب سازمان‌ها را به جمع‌آوری داده‌ها از منابع مختلف هدایت می‌کند و به طور بالقوه اطلاعات حساس کاربر را در معرض دید قرار می‌دهد. توافق‌نامه‌های اشتراک‌گذاری داده‌ها و پروژه‌های تحلیل داده‌های مشترک می‌توانند این نگرانی‌ها را تشدید کنند، به‌ویژه زمانی که داده‌ها در مرزهای بین‌المللی با مقررات مختلف حریم خصوصی به اشتراک گذاشته می‌شوند.

### حملات استنتاج

مدل‌های هوش مصنوعی می‌توانند به طور ناخواسته اطلاعات حساسی را در مورد داده‌های آموزشی نشان دهند، حتی زمانی که داده‌ها ناشناس هستند. به عنوان مثال، مهاجمان می‌توانند از وارونگی مدل یا حملات استنتاج عضویت برای استخراج اطلاعات خصوصی از پیش‌بینی‌های یک مدل استفاده کنند یا یاد بگیرند که آیا یک نقطه داده خاص در مجموعه آموزشی گنجانده شده است.

**حریم خصوصی دیفرانسیل:** حریم خصوصی دیفرانسیل (DP) یک رویکرد محبوب برای حفظ حریم خصوصی در طول تجزیه و تحلیل داده‌ها با افزودن نویز کنترل شده به داده‌ها است.

### ۳,۲۰. چالش‌های امنیت داده‌ها در هوش مصنوعی

### ۳,۲۱. حملات خصمانه

حملات خصمانه، که در آن اغتشاشات کوچکی به داده‌های ورودی برای فریب مدل‌های

هوش مصنوعی وارد می‌شود، خطرات امنیتی قابل توجهی را به همراه دارد.. این حملات می‌توانند منجر به پیش‌بینی‌ها یا طبقه‌بندی‌های نادرست شوند و قابلیت اطمینان سیستم‌های هوش مصنوعی را در برنامه‌های کاربردی حیاتی، مانند مراقبت‌های بهداشتی، مالی، و وسایل نقلیه خودران تضعیف کنند.

### ۳,۲۲. مسمومیت داده‌ها

حملات مسمومیت داده شامل دستکاری داده‌های آموزشی برای کاهش عملکرد یک مدل هوش مصنوعی است. شناسایی این حملات می‌تواند دشوار باشد زیرا اغلب زیر مجموعه کوچکی از داده‌های آموزشی را هدف قرار می‌دهند و به حداقل تغییرات در نقاط داده مسموم نیاز دارند.

### ۳,۲۳. دستکاری مدل و داده

مهاجمان همچنین می‌توانند مدل‌ها و داده‌های هوش مصنوعی را مستقیماً با تغییر پارامترهای مدل، وزن‌ها یا خود داده‌ها هدف قرار دهند. تکنیک‌هایی مانند حملات درب پشتی یا شبکه‌های عصبی تروجان می‌توانند رفتار مخرب پنهان را در مدل‌های هوش مصنوعی معرفی کنند و در نتیجه خطرات امنیتی قابل توجهی را به همراه داشته باشند. در حملات درب‌پشتی، مهاجم کدهای مخرب را در طول فرآیند آموزش به مدل تزریق می‌کند که باعث می‌شود مدل خروجی‌های نادرست تولید کند یا رفتار ناخواسته‌ای از خود نشان دهد که توسط یک الگوی ورودی خاص تحریک می‌شود. از سوی دیگر، شبکه‌های عصبی تروجان شامل تعبیه یک ماشه پنهان در مدل است که وقتی توسط یک ورودی خاص فعال می‌شود، باعث می‌شود مدل اقدامات غیرمجاز انجام دهد یا پیش‌بینی‌های نادرستی ارائه دهد.

حملات استخراج مدل شامل دستکاری مدل است که در آن مهاجم به دنبال تکرار مدل

هدف با پرس و جو و یادگیری از پاسخ‌ها است. این می‌تواند منجر به سرقت مالکیت معنوی، یا حتی ایجاد مدل‌های تکراری با نیت مخرب شود.

برای دفاع در برابر دستکاری مدل و داده‌ها، سازمان‌ها می‌توانند از استراتژی‌های مختلفی مانند سخت‌سازی مدل، ذخیره‌سازی امن مدل و اعتبارسنجی ورودی استفاده کنند. تکنیک‌های سخت‌سازی مدل، مانند هرس دقیق، می‌توانند به حذف اجزای مخرب از یک مدل کمک کنند و در عین حال عملکرد کلی آن را حفظ کنند (تالون، ۲۰۱۳). ذخیره‌سازی امن مدل با استفاده از رمزگذاری و کنترل‌های دسترسی می‌تواند یک مدل را از تغییرات غیرمجاز محافظت کند. اعتبارسنجی ورودی را می‌توان برای اطمینان از اینکه فقط ورودی‌های قانونی توسط سیستم هوش مصنوعی پردازش می‌شود، استفاده کرد و خطر ایجاد درهای پشتی مخفی یا شبکه‌های تروجان را کاهش داده است. با اجرای این اقدامات متقابل، سازمان‌ها می‌توانند امنیت و قابلیت اعتماد سیستم‌های هوش مصنوعی خود را در مواجهه با تهدیدات مدل و دستکاری داده‌ها افزایش دهند.

### نتیجه‌گیری

امروزه هوش مصنوعی، صنایع و تجهیزات مدرن را متحول کرده است و ابزارهای قدرتمندی برای افزایش آگاهی موقعیتی، بهبود تصمیم‌گیری، تسهیل هماهنگی و کاهش خطرات در اختیار تصمیم‌گیرندگان قرار می‌دهد. با استفاده از قابلیت‌های هوش مصنوعی، سازمان‌ها و یگان‌های مختلف می‌توانند انعطاف‌پذیری خود را در برابر بحران‌ها تقویت کنند، تأثیر آن‌ها را به حداقل برسانند و از بهبودی سریع و مؤثر اطمینان حاصل کنند. استفاده از الگوریتم‌های هوش مصنوعی می‌تواند به طرز قابل توجهی در بهبود کارایی و کاهش خطرات کمک نماید. بطور خلاصه به چند قابلیت‌های کلیدی در کنترل و شناسایی عوامل N.B.C را اشاره نمود:

### • شناسایی و پایش

- شناسایی تهدیدات و تصویربرداری و تحلیل: به شناسایی زودهنگام خطرات (شیمیایی، بیولوژیکی و هسته‌ای) کمک کند؛ این سیستم‌ها از سنسورهای مختلف و الگوریتم‌های یادگیری ماشین و یا با استفاده از تصاویر هوایی و ماهواره‌ای برای شناسایی تغییرات محیطی برای تحلیل داده‌ها و تشخیص علائم اولیه آلودگی استفاده می‌کنند.

- پیش‌بینی، مدل‌سازی و انتشار: با استفاده از داده‌های تاریخی و مدل‌سازی انتشار مواد خطرناک و تحلیل اثرات آنها بر روی محیطی و سلامت انسان‌ها کمک کند.

• مدیریت سانحه و واکنش

- برنامه‌ریزی و هماهنگی: این سیستم‌ها میتوانند به مدیریت منابع و هماهنگی اقدامات در مواقع بحران کمک کنند، از جمله تخصیص منابع، مدیریت زمانبندی، عملیات، روند انتشار و آلودگی مواد خطرناک

## منابع

آژانس بین‌المللی انرژی اتمی (IAEA)، (۱۳۸۲). مترجم انستیتو پرتو پزشکی نوین، روش ایجاد آمادگی پاسخ اضطراری در سوانح، سند فنی شماره ۹۵۳، ص ۸۰.

- Aguiar-Pérez, J. M., Pérez-Juárez, M. A., Alonso-Felipe, M., Del-Pozo-Velázquez, J., Rozada-Raneros, S., & Barrio-Conde, M. (2021). [Title of the article]. [Journal Name], Volume, [Page range]. [https://doi.org/\[DOI Number\]](https://doi.org/[DOI Number])
- Alvarez-Coello, D., Wilms, D., Bekan, A., & Gómez, J. M. (2021). Towards a data-centric architecture in the automotive industry. *Procedia Computer Science*, 181, 658–663. <https://doi.org/10.1016/j.procs.2021.01.211>
- Barocas, S., Hardt, M., & Narayanan, A. (2021). *Fairness and machine learning: Limitations and opportunities*. [fairmlbook.org](https://fairmlbook.org). <https://fairmlbook.org>
- Batista, G. E., & Monard, M. C. (2018). Data quality in machine learning: A study in the context of imbalanced data. *Neurocomputing*, 275, 1665–1679. <https://doi.org/10.1016/j.neucom.2017.09.096>
- Bergen, K. J., Johnson, P. A., de Hoop, M. V., & Beroza, G. C. (2019). Machine learning for data-driven discovery in solid Earth geoscience. *Science*, 363(6433), eaau0323. <https://doi.org/10.1126/science.aau0323>
- Broo, D. G., & Jennifer, S. (2020). Towards data-centric decision making for smart infrastructure: Data and its challenges. *\*IFAC-PapersOnLine*, 53\*(2), 352–357. <https://doi.org/10.1016/j.ifacol.2020.12.143>
- Burr, C., & Leslie, D. (2023). Ethical assurance: A practical approach to the responsible design, development, and deployment of data-driven technologies. *AI and Ethics*, 3, 73–98. <https://doi.org/10.1007/s43681-022-00154-8>
- Chen, H., Chiang, R. H., & Storey, V. C. (2012). Business intelligence and analytics: From big data to big impact. *MIS Quarterly*, 36(4), 1165–1188. <https://doi.org/10.2307/41703503>
- Dahiya, N., Sheifali, G., & Sartajvir, S. (2022). A review paper on machine learning applications, advantages, and techniques. *ECS Transactions*, 107(1), 6137. <https://doi.org/10.1149/10701.6137ecst>
- Daries, J. P., Reich, J., Waldo, J., Young, E. M., Whittinghill, J., Ho, A. D., & Chuang, I. (2014). Privacy, anonymity, and big data in the social sciences. *Communications of the ACM*, 57(9), 56–63. <https://doi.org/10.1145/2643132>
- Devlin, J., Chang, M.-W., Lee, K., & Toutanova, K. (2019). BERT: Pre-training of deep bidirectional transformers for language understanding. In *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)* (pp. 4171–4186). Association for Computational Linguistics. <https://doi.org/10.18653/v1/N19-1423>

- Enholm, I. M., Papagiannidis, E., Mikalef, P., & Krogstie, J. (2022). Artificial intelligence and business value: A literature review. *Information Systems Frontiers*, 24, 1709–1734. <https://doi.org/10.1007/s10796-021-10186-w>
- Gandomi, A., & Haider, M. (2015). Beyond the hype: Big data concepts, methods, and analytics. *International Journal of Information Management*, 35(2), 137–144. <https://doi.org/10.1016/j.ijinfomgt.2014.10.007>
- García, S., Luengo, J., & Herrera, F. (2016). *Data preprocessing in data mining*. Springer.
- Géron, A. (2019). *Hands-on machine learning with Scikit-Learn, Keras, and TensorFlow: Concepts, tools, and techniques to build intelligent systems* (2nd ed.). O'Reilly Media.
- Gumbs, A. A., Grasso, V., Bourdel, N., Croner, R., Spolverato, G., Frigerio, I., Illanes, A., Abu Hilal, M., Park, A., & Elyan, E. (2022). The advances in computer vision that are enabling more autonomous actions in surgery: A systematic review of the literature. *Sensors*, 22(12), 4918. <https://doi.org/10.3390/s22134918>
- Guyon, I., Gunn, S., & Ben-Hur, A. (2004). Result analysis of the NIPS 2003 feature selection challenge. *Advances in Neural Information Processing Systems*, 17, 545–552.
- Hajian, S., Bonchi, F., & Castillo, C. (2016). Algorithmic bias: From discrimination discovery to fairness-aware data mining. In *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining* (pp. 2125–2126). Association for Computing Machinery. <https://doi.org/10.1145/2939672.2945386>
- Halevy, A., Korn, F., Noy, N., Olston, C., Polyzotis, N., Roy, S., & Whang, S. (2020). Goods: Organizing Google's datasets. *Communications of the ACM*, 63(7), 50–57. <https://doi.org/10.1145/3342102>
- Hassan, N. U., Asghar, M. Z., Ahmed, S., & Zafar, H. (2021). A survey on data quality issues in big data. *ACM Computing Surveys*, 54(7), 1–37. <https://doi.org/10.1145/3462517>
- Hershey, P. A. (2023). Understanding machine learning concepts. In J. Wang (Ed.), *Encyclopedia of data science and machine learning* (pp. 1007–1022). IGI Global. <https://doi.org/10.4018/978-1-7998-9220-5.ch053>
- Jakubik, J., Vössing, M., Kühn, N., Walk, J., & Satzger, G. (2022). *Data-centric artificial intelligence*. arXiv. <https://arxiv.org/abs/2212.11854>
- Jarrahi, M. H., Ali, M., & Shion, G. (2022). *The principles of data-centric AI (DCAI)*. arXiv. <https://arxiv.org/abs/2211.14611>
- Jöckel, L., & Michael, K. (2019). Increasing trust in data-driven model validation: A framework for probabilistic augmentation of images and meta-data generation using application scope characteristics. In A. Romanovsky, E. Schoitsch, F. Bitsch, & F. Koornneef (Eds.), *Computer safety, reliability, and*

- security. *SAFECOMP 2019*. Lecture Notes in Computer Science, vol 11699. Springer. [https://doi.org/10.1007/978-3-030-26250-1\\_23](https://doi.org/10.1007/978-3-030-26250-1_23)
- Juran, J. M., & Godfrey, A. B. (2018). *Juran's quality handbook: The complete guide to performance excellence* (7th ed.). McGraw-Hill Education.
- Kanter, J. M., Benjamin, S., & Kalyan, V. (2018). *Machine Learning 2.0: Engineering data driven AI products*. arXiv. <https://arxiv.org/abs/1807.00401>
- Karkouch, A., Mousannif, H., Al Moatassime, H., & Noel, T. (2018). Data quality in the Internet of Things: A state-of-the-art survey. *Journal of Network and Computer Applications*, 124, 289–310. <https://doi.org/10.1016/j.jnca.2018.07.002>
- Kelleher, J. D., Mac Namee, B., & D'Arcy, A. (2018). *Fundamentals of machine learning for predictive data analytics: Algorithms, worked examples, and case studies* (2nd ed.). MIT Press.
- Khatri, V., & Brown, C. V. (2010). Designing data governance. *Communications of the ACM*, 53(1), 148–152. <https://doi.org/10.1145/1629175.1629210>
- Li, X.-H., Cao, C. C., Shi, Y., Bai, W., Gao, H., Qiu, L., Wang, C., Gao, Y., Zhang, S., Xue, X., & et al. (2020). A survey of data-driven and knowledge-aware explainable AI. *IEEE Transactions on Knowledge and Data Engineering*, 34(1), 29–49. <https://doi.org/10.1109/TKDE.2020.2983930>
- Li, Y. (2017). *Deep reinforcement learning: An overview*. arXiv. <https://arxiv.org/abs/1701.07274>
- Little, R. J., & Rubin, D. B. (2019). *Statistical analysis with missing data* (3rd ed.). John Wiley & Sons.
- Liu, X., Yoo, C., Xing, F., Oh, H., El Fakhri, G., Kang, J.-W., & Woo, J. (2022). Deep unsupervised domain adaptation: A review of recent advances and perspectives. *APSIPA Transactions on Signal and Information Processing*, 11, e25. <https://doi.org/10.1561/116.00000055>
- Lomas, J., Nirmal, P., & Jodi, F. (2018). Continuous improvement: How systems design can benefit the data-driven design community. In *Proceedings of RSD7, Relating Systems Thinking and Design 7*, Turin, Italy.
- Maranghi, M., Anagnostopoulos, A., Cannistraci, I., Chatzigiannakis, I., Croce, F., Di Teodoro, G., Gentile, M., Grani, G., Lenzerini, M., Leonardi, S., & et al. (2022). *AI-based data preparation and data analytics in healthcare: The case of diabetes*. arXiv. <https://arxiv.org/abs/2206.06182>
- Marr, B. (2018). *Artificial intelligence in practice: How 50 successful companies used AI and machine learning to solve problems*. John Wiley & Sons.
- Mazumder, M., Banbury, C., Yao, X., Karlaš, B., Rojas, W. G., Damos, S., Damos, G., He, L., Kiela, D., Jurado, D., & et al. (2022). *DataPerf: Benchmarks for data-centric AI development*. arXiv. <https://arxiv.org/abs/2207.10062>
- Miranda, L. J. (2021). *Towards data-centric machine learning: A short review*. Retrieved April 15, 2023, from <https://ljevimiranda921.github.io/notebook/2021/07/30/data-centric-ml/>

- Mnih, V., Kavukcuoglu, K., Silver, D., Rusu, A. A., Veness, J., Bellemare, M. G., Petersen, S. (2015). Human-level control through deep reinforcement learning. *Nature*, 518(7540), 529–533. <https://doi.org/10.1038/nature14236>
- Ntoutsis, E., Fafalios, P., Gadiraju, U., Iosifidis, V., Nejdli, W., Vidal, M. E., Ruggieri, S., Turini, F., Papadopoulos, S., Krasanakis, E., & et al. (2020). Bias in data-driven artificial intelligence systems—An introductory survey. *Wiley Interdisciplinary Reviews: Data Mining and Knowledge Discovery*, 10(3), e1356. <https://doi.org/10.1002/widm.1356>
- O’Leary, D. E. (2013). Artificial intelligence and big data. *IEEE Intelligent Systems*, 28(2), 96–99. <https://doi.org/10.1109/MIS.2013.39>
- Otto, B. (2011). Organizing data quality management in enterprises. In *Proceedings of the 17th Americas Conference on Information Systems (AMCIS)*, Detroit, MI, United States.
- Panian, Z. (2010). Some practical experiences in data governance. *World Academy of Science, Engineering and Technology*, 66, 1248–1253.
- Pipino, L. L., Lee, Y. W., & Wang, R. Y. (2018). Data quality assessment. In R. Y. Wang, L. L. Pipino, & Y. W. Lee (Eds.), *Data and information quality: Dimensions, principles and techniques* (pp. 219–253). Springer. [https://doi.org/10.1007/978-3-319-65006-6\\_7](https://doi.org/10.1007/978-3-319-65006-6_7)
- Pouyanfar, S., Sadiq, S., Yan, Y., Tian, H., Tao, Y., Reyes, M. P., Shyu, M.-L., Chen, S.-C., & Iyengar, S. S. (2018). A survey on deep learning: Algorithms, techniques, and applications. *ACM Computing Surveys*, 51(5), 1–36. <https://doi.org/10.1145/3234150>
- Redman, T. C. (1996). *Data quality for the information age*. Artech House, Inc.
- Russell, S. J., & Norvig, P. (2016). *Artificial intelligence: A modern approach* (3rd ed.). Pearson Education Limited.
- Sharma, L., & Garg, P. K. (2021). *Artificial intelligence: Technologies, applications, and challenges*. Taylor & Francis.
- Sun, X., Liu, Y., & Liu, J. (2018). Ensemble learning for multi-source remote sensing data classification based on different feature extraction methods. *IEEE Access*, 6, 50861–50869. <https://doi.org/10.1109/ACCESS.2018.2869578>
- Tallon, P. P. (2013). Corporate governance of big data: Perspectives on value, risk, and cost. *IEEE Computer*, 46(6), 32–38. <https://doi.org/10.1109/MC.2013.155>
- Uddin, M. F., & Navarun, G. (2014). Seven V’s of Big Data understanding Big Data to extract value. In *Proceedings of the 2014 Zone 1 Conference of the American Society for Engineering Education*, Bridgeport, CT, USA. <https://doi.org/10.1109/ASEEZone1.2014.6820689>
- Wang, Z., Li, M., Lu, J., & Cheng, X. (2022). Business Innovation based on artificial intelligence and Blockchain technology. *Information Processing & Management*, 59(1), 102759. <https://doi.org/10.1016/j.ipm.2021.102759>

- Weill, P., & Ross, J. W. (2004). *IT governance: How top performers manage IT decision rights for superior results*. Harvard Business School Press.
- Xu, K., Li, Y., Liu, C., Liu, X., Hao, X., Gao, J., & Maropoulos, P. G. (2020). Advanced data collection and analysis in data-driven manufacturing process. *Chinese Journal of Mechanical Engineering*, 33(1), 1–21. <https://doi.org/10.1186/s10033-020-00470-2>
- Yablonsky, S. (2019). Multidimensional data-driven artificial intelligence innovation. *Technology Innovation Management Review*, 9(12), 16–28. <https://doi.org/10.22215/timreview/1288>
- Yang, Y., Zheng, L., Zhang, J., Cui, Q., Li, Z., & Yu, P. S. (2018). *TI-CNN: Convolutional neural networks for fake news detection*. arXiv. <https://arxiv.org/abs/1806.00749>
- Zha, D., Bhat, Z. P., Lai, K.-H., Yang, F., & Hu, X. (2023). *Data-centric AI: Perspectives and challenges*. arXiv. <https://arxiv.org/abs/2301.04819>
- Zha, D., Bhat, Z. P., Lai, K. H., Yang, F., Jiang, Z., Zhong, S., & Hu, X. (2023). *Data-centric artificial intelligence: A survey*. arXiv. <https://arxiv.org/abs/2303.10158>
- Zhang, X. D., Wu, H. Y., Wu, D., Wang, Y. Y., Chang, J. H., Zhai, Z. B., & Fan, F. Y. (2009). Toxicologic effects of gold nanoparticles in vivo by different administration routes. *International Journal of Nanomedicine*, 5, 771–781. <https://doi.org/10.2147/IJN.S14912>
- Zhuang, F., Qi, Z., Duan, K., Xi, D., Zhu, Y., Zhu, H., Xiong, H., & He, Q. (2020). A comprehensive survey on transfer learning. *Proceedings of the IEEE*, 109(1), 43–76. <https://doi.org/10.1109/JPROC.2020.3004555>