

اخلاق هوش مصنوعی در پژوهش‌های زیست‌پزشکی

از انتخاب ابزار و ورود داده تا بررسی خروجی و پذیرش مسئولیت علمی



دکتر آرش نجیمی

- ❑ در پایان این ارائه انتظار می‌رود بتوانیم:
- ❑ منابع اصلی اخلاق هوش مصنوعی را بشناسیم.
- ❑ اصول اخلاقی را در استفاده روزمره پژوهشی به کار ببریم.
- ❑ خطرهای مرتبط با داده، خروجی و تصمیم‌گیری را تشخیص دهیم.
- ❑ برای پذیرش، اصلاح یا کنار گذاشتن خروجی AI تصمیم مستدل بگیریم.



یک موقعیت چالش‌برانگیز در پژوهش



یک پژوهشگر قصد دارد عوامل مرتبط با کنترل نامناسب دیابت را بررسی کند. پرسش آغازین: پیش از پذیرش این خروجی، چه تصمیم‌های اخلاقی باید گرفته شود؟ نام بیماران را از فایل حذف و داده‌ها را در یک چت‌بات بارگذاری می‌کند. ابزار، روش تحلیل، کد، تفسیر نتایج و چند منبع پیشنهاد می‌دهد. پژوهشگر پاسخ را دقیق و قانع‌کننده می‌بیند و از آن استفاده می‌کند.

سناریوی آموزشی فرضی

❑ در این ارائه، AI ابزاری برای کمک به فعالیتهای پژوهشی است:

❑ ایده پردازی و صورت بندی سؤال

❑ جست و جو و خلاصه سازی منابع

❑ استخراج اطلاعات و تولید کد

❑ کمک به تحلیل و تفسیر

❑ نگارش، ترجمه و مدیریت اطلاعات

تمرکز ما مسئولیت پژوهشگر هنگام استفاده از این قابلیت‌هاست.



- ❑ سرعت بیشتر، فرصت انجام کارهای بیشتری را فراهم می‌کند.
- ❑ دسترسی آسان به پاسخ، نیاز به بررسی اعتبار آن را افزایش می‌دهد.
- ❑ تولید متن و کد، مسئولیت کنترل صحت را از پژوهشگر نمی‌گیرد.
- ❑ خودکارسازی فعالیت‌ها باید با نظارت متناسب همراه باشد.

هر قابلیت تازه، یک تصمیم تازه درباره حدود اعتماد و مسئولیت ایجاد می‌کند.



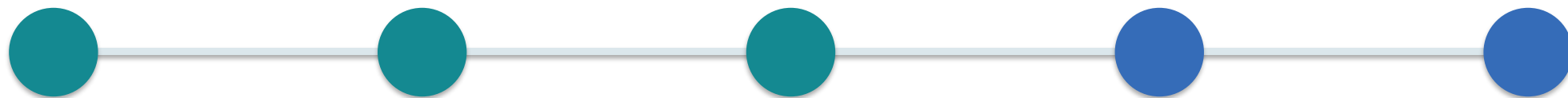
- ☐ آیا این ابزار برای هدف پژوهش مناسب است؟
- ☐ چه اطلاعاتی می‌توان وارد آن کرد؟
- ☐ چه مقدار از خروجی نیازمند بررسی مستقل است؟
- ☐ چه کسی درباره پذیرش نتیجه تصمیم می‌گیرد؟
- ☐ اگر خطایی رخ دهد، چه کسی آن را اصلاح می‌کند؟

اخلاق استفاده، در سراسر فرایند پژوهش حضور دارد.

نقشه تحول اخلاق هوش مصنوعی

7

چند نقطه شاخص در تحول راهنماها:



۲۰۱۹: اصول OECD، راهنمای اروپایی و سند IEEE

۲۰۲۳-۲۰۲۴: چارچوب NIST و راهنمای اختصاصی AI مولد

۲۰۱۷-۲۰۱۸: اصول آسیلومار و اعلامیه مونترآل

۲۰۲۱: توصیه‌نامه یونسکو و راهنمای WHO

۲۰۲۶: ویرایش سوم راهنمای اروپایی استفاده از AI مولد در پژوهش



- ☐ اصل اخلاقی: ارزش راهنما، مانند عدالت یا پاسخ‌گویی
- ☐ راهنما: توصیه درباره نحوه عمل در موقعیت مشخص
- ☐ چارچوب اجرایی: روش شناسایی، سنجش و مدیریت خطر
- ☐ چک‌لیست: پرسش‌هایی برای بررسی رعایت توصیه‌ها

وجود یک راهنما به‌تنهایی به معنای الزام قانونی یکسان در همه کشورها نیست.

آسیلومار و مونترآل: آغاز یک گفت‌وگوی جهانی

- ❑ برداشت کاربردی برای پژوهشگر: ارزش استفاده از AI با اثر آن بر انسان‌ها و جامعه سنجیده می‌شود.
- ❑ اصول آسیلومار بر سودمندی AI، ایمنی، کنترل انسانی و مسئولیت تأکید می‌کنند.
- ❑ اعلامیه مونترآل از گفت‌وگوی دانشگاهیان، شهروندان و دیگر ذی‌نفعان شکل گرفت.
- ❑ این اسناد، پیامدهای اجتماعی AI را در کنار قابلیت‌های فنی قرار دادند.

Asilomar AI Principles

[View original resource](#)

Summary

The Asilomar AI Principles represent a historic moment when 90+ AI researchers, ethicists, and thought leaders gathered in Asilomar, California to establish shared guidelines for beneficial AI development. These 23 principles span three critical timeframes: immediate research priorities, near-term ethical considerations, and long-term safety challenges as AI systems become more capable. Unlike top-down regulatory frameworks, these principles emerged from the AI community itself, creating a foundation that has influenced countless AI ethics initiatives, corporate policies, and regulatory discussions worldwide.

OECD > Topics > AI principles

AI principles

The OECD AI Principles are the first intergovernmental standard on AI. They promote innovative, trustworthy AI that respects human rights and democratic values. Adopted in 2019 and updated in 2024, they are composed of five values-based principles and five recommendations that provide practical and flexible guidance for policymakers and AI actors.

Policy sub-issue

OECD.AI Policy Observatory

Key links

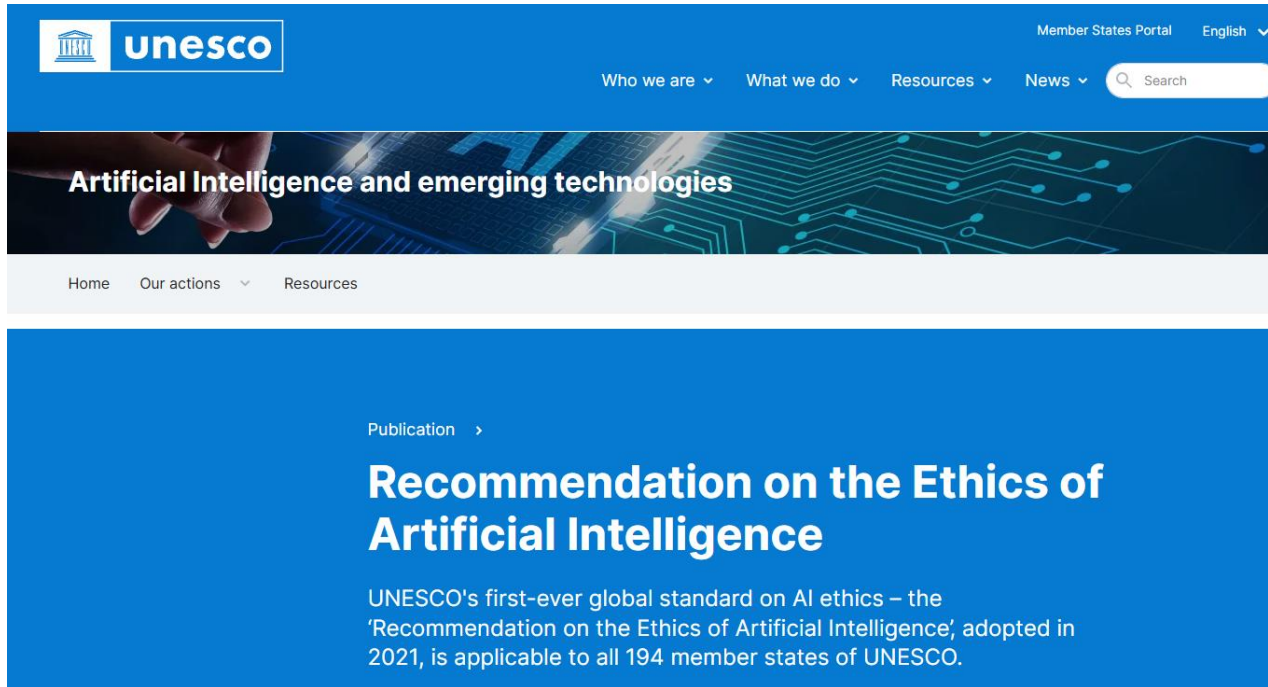
[The state of implementation of the OECD AI Principles >](#)

[The OECD AI Wonk Blog >](#)

پنج محور ارزش‌محور OECD

- ☐ رشد فراگیر، توسعه پایدار و رفاه
- ☐ حقوق انسانی و ارزش‌های دموکراتیک، شامل عدالت و حریم خصوصی
- ☐ شفافیت و توضیح‌پذیری
- ☐ استحکام، امنیت و ایمنی
- ☐ پاسخ‌گویی

بازنگری ۲۰۲۴ به تحولات AI عمومی و مولد توجه بیشتری دارد.



توصیه‌نامه یونسکو بر چهار ارزش بنیادی استوار است:

- ❑ کرامت و حقوق انسانی
- ❑ جوامع عادلانه، صلح‌آمیز و به‌هم‌پیوسته
- ❑ تنوع و شمول
- ❑ حفاظت از محیط‌زیست و زیست‌بوم‌ها

تناسب استفاده، نظارت انسانی، حفاظت از داده و پاسخ‌گویی از اصول مهم این سند هستند.

کاربرد برای پژوهشگر: این الزامات می‌توانند معیار پرسشگری هنگام انتخاب و استفاده از ابزار باشند.

☐ اختیار و نظارت انسانی

☐ استحکام فنی و ایمنی

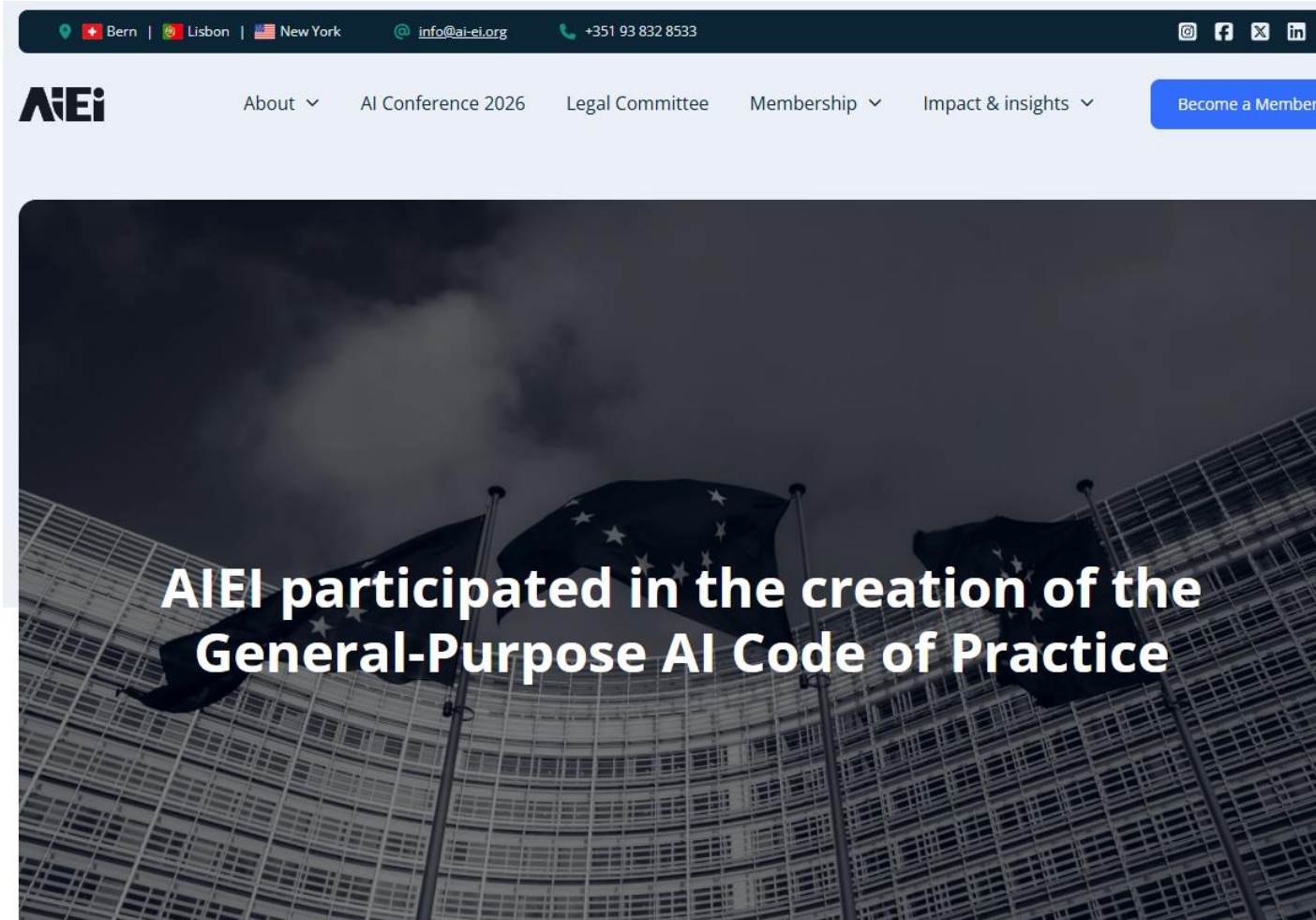
☐ حریم خصوصی و حکمرانی داده

☐ شفافیت

☐ تنوع، عدم تبعیض و عدالت

☐ رفاه اجتماعی و محیط‌زیستی

☐ پاسخ‌گویی



سند طراحی همسو با اخلاق IEEE، در کنار مسئولیت سازندگان، به نقش استفاده‌کنندگان نیز توجه دارد.

□ پیام کاربردی: دسترسی به ابزار، به‌تنهایی شایستگی استفاده از آن را ایجاد نمی‌کند.

□ استفاده‌کننده باید دانش و مهارت لازم را داشته باشد.

□ اثربخشی ابزار باید با هدف کاربرد تناسب داشته باشد.

□ خطرهای سوءاستفاده باید شناخته شوند.

□ اختیار افراد بر داده‌هایشان اهمیت دارد.



- ☐ حفاظت از خودمختاری انسان
- ☐ ارتقای رفاه، ایمنی و منفعت عمومی
- ☐ شفافیت، توضیح‌پذیری و فهم‌پذیری
- ☐ مسئولیت و پاسخ‌گویی
- ☐ شمول و عدالت
- ☐ هوش مصنوعی پاسخ‌گو به نیازها و پایدار

این اصول، پایه مهمی برای استفاده از AI در پژوهش‌های زیست‌پزشکی هستند.

راهنمای معرفی‌شده در ژانویه ۲۰۲۴ به ابزارهایی می‌پردازد که می‌توانند با متن، تصویر و دیگر انواع داده کار کنند.

خطرهای مهم برای استفاده‌کننده:

- ☐ پاسخ نادرست یا ناقص
- ☐ سوگیری در محتوا
- ☐ اعتماد بیش‌ازحد به ابزار
- ☐ تهدید امنیت و حرمانگی اطلاعات
- ☐ نابرابری در دسترسی به ابزارهای مناسب



Ethics and governance of artificial intelligence for health

Guidance on large multi-modal models



NIST: تبدیل اصول به مدیریت خطر ¹⁶

چهار کارکرد چارچوب NIST

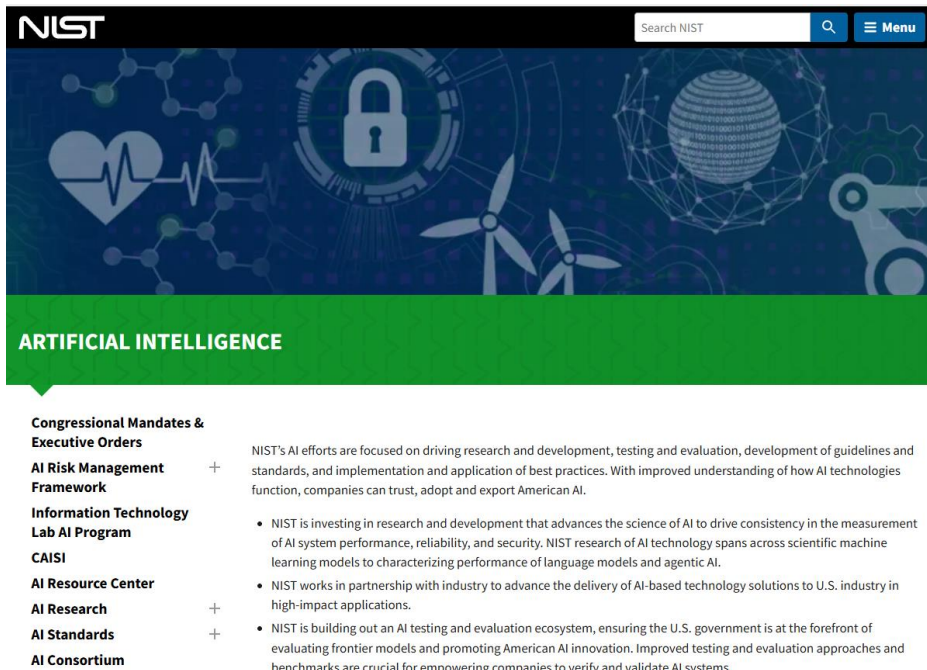
❑ کاربرد پژوهشی: برای هر استفاده مهم، خطر و اقدام کنترلی مشخص شود.

❑ حکمرانی: تعیین مسئولیت‌ها و قواعد استفاده

❑ شناخت زمینه و خطر: روشن کردن کاربرد، داده و پیامدها

❑ سنجش: بررسی کیفیت خروجی و میزان خطر

❑ مدیریت: کاهش خطر و تصمیم درباره ادامه استفاده



راهنمای استفاده مسئولانه از Ai در پژوهش¹⁷

راهنمای اروپایی، ویرایش سوم مه ۲۰۲۶، مستقیماً پژوهشگران را مخاطب قرار می‌دهد.

❑ مسئولیت خروجی با پژوهشگر باقی می‌ماند.

❑ استفاده از پاسخ‌ها باید نقادانه باشد.

❑ محدودیت‌ها، سوگیری‌ها و اطلاعات نادرست باید بررسی شوند.

❑ محرمانگی و حقوق مرتبط با اطلاعات رعایت شوند.

این سند، یکی از منابع مستقیم ارائه حاضر است.



۱. سودمندی و تناسب

Clear and understandable AI operations



۲. اختیار و نظارت انسانی

Reliable and precise AI-generated decisions

۳. صحت و قابلیت اعتماد

۴. عدالت و پرهیز از سوگیری



۵. حریم خصوصی و حفاظت از داده

Promote fairness and eliminate prejudice



۶. شفافیت و قابلیت پیگیری

Trustworthy and explainable AI outputs

۷. مسئولیت و پاسخ‌گویی

۸. پایداری و دسترسی عادلانه



PRIVACY

Respect and safeguard personal data

سودمندی و تناسب استفاده

❑ پرسش راهنما :چرا این فعالیت را به AI می‌سپاریم؟

❑ هدف استفاده از AI باید روشن باشد.

❑ ابزار باید با نوع مسئله تناسب داشته باشد.

❑ حجم داده ورودی و میزان خودکارسازی، بیش از نیاز نباشد.

❑ فایده مورد انتظار باید در کنار خطر و هزینه بررسی شود.



❑ پرسش راهنما: چه کسی می‌تواند با دلیل، پیشنهاد ابزار را رد کند؟

❑ پژوهشگر باید توان فهم و نقد خروجی را حفظ کند.

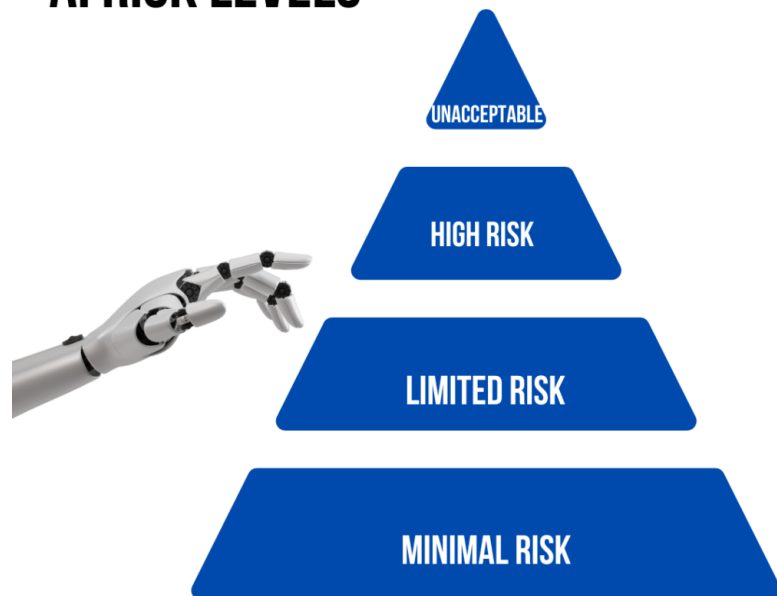
❑ امکان رد، اصلاح و توقف استفاده وجود داشته باشد.

❑ بازیبن باید تخصص، زمان و اختیار کافی داشته باشد.

❑ حضور صوری یک انسان، نظارت مؤثر ایجاد نمی‌کند.



AI RISK LEVELS



SOURCE: EUROPEAN COMMISSION (2023) ILLUSTRATION: MELISSA GUERRA (2024)

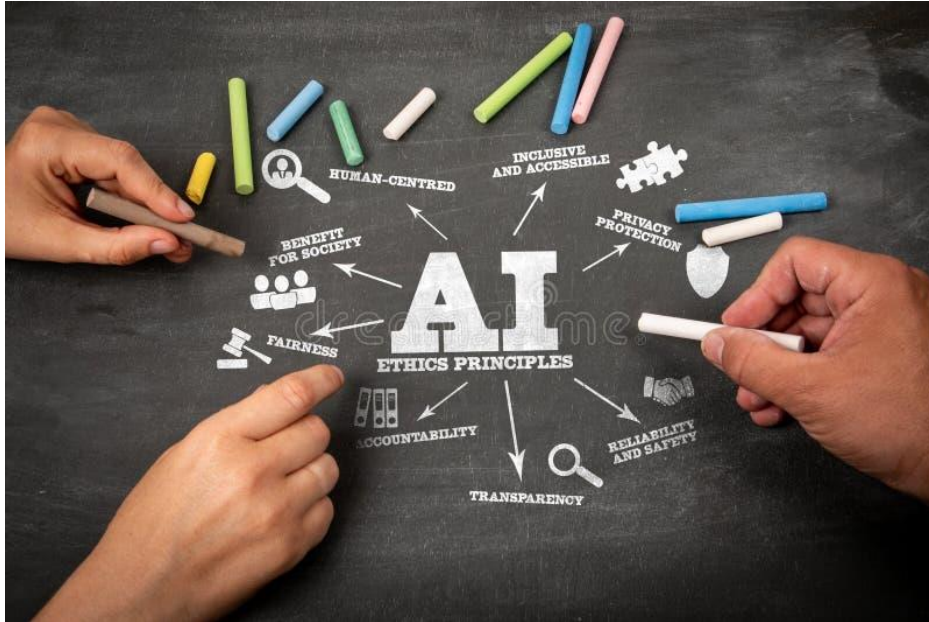
آسیب ناشی از استفاده از AI می‌تواند شامل موارد زیر باشد:

- ❑ نتیجه‌گیری پژوهشی نادرست
- ❑ افشای اطلاعات محرمانه
- ❑ تقویت تبعیض یا برداشت نامناسب از یک گروه
- ❑ اتلاف منابع و هدایت پژوهش به مسیر اشتباه

شدت بررسی باید با احتمال و پیامد خطا تناسب داشته باشد.

- ☐ پرسش راهنما: چه کسی یا چه دیدگاهی در این پاسخ دیده نشده است؟
- ☐ پاسخ ابزار ممکن است برخی زبان‌ها یا جمعیت‌ها را کمتر بازتاب دهد.
- ☐ شواهد در دسترس ابزار، نماینده همه زمینه‌ها نیستند.
- ☐ پرسش جهت‌دار می‌تواند خروجی را به سمت ترجیح پژوهشگر ببرد.
- ☐ کمبود شواهد درباره یک گروه نباید پنهان بماند.

حریم خصوصی و اختیار بر داده



- ☐ پرسش راهنما: داده من پس از ارسال چه مسیری طی می‌کند؟
- ☐ بارگذاری فایل، تصمیمی درباره دسترسی به اطلاعات است.
- ☐ نوع داده و شرایط پردازش باید پیش از ارسال بررسی شوند.
- ☐ فقط اطلاعات لازم برای انجام وظیفه وارد ابزار شود.
- ☐ داده بیماران، همکاران و طرح‌های منتشرنشده همگی نیازمند توجه‌اند.

□ نقش AI در انجام فعالیت پژوهشی روشن باشد.

□ قابلیت‌ها و محدودیت‌های ابزار شناخته شوند.

□ بررسی‌ها و اصلاحات انسانی مشخص باشند.

□ عدم قطعیت یا کمبود اطلاعات پنهان نماند.

شفافیت باید امکان فهم و بازبینی فرایند را فراهم کند.

- پرسش راهنما: آیا دلیل پذیرش این پاسخ را مستقل از بیان ابزار می‌دانیم؟
- پژوهشگر باید بداند خروجی برای چه کاری قابل استفاده است.
- ادعاهای مهم باید به شواهد قابل بررسی متصل شوند.
- توضیح روان مدل درباره پاسخ خود، تضمین درستی آن نیست.
- در صورت نبود امکان بررسی، اتکا به خروجی باید محدود شود.

❑ مسئولیت علمی با پژوهشگر و تیم پژوهش باقی می‌ماند.

❑ هر خروجی حساس باید بازبین مشخص داشته باشد.

❑ مسیر اصلاح خطا و اطلاع‌رسانی روشن باشد.

❑ تصمیم پذیرش خروجی باید قابل توضیح باشد.

«ابزار این‌طور گفته بود «دلیل کافی برای دفاع از تصمیم پژوهشی نیست».

- ❑ ابزار برای وظیفه موردنظر باید کیفیت قابل قبول داشته باشد.
- ❑ تغییر نسخه یا شرایط استفاده ممکن است خروجی را تغییر دهد.
- ❑ دسترسی ابزار به فایل‌ها و سامانه‌ها باید محدود و متناسب باشد.
- ❑ خروجی‌های مهم به کنترل مستقل نیاز دارند.

اعتماد باید به کاربرد مشخص و شواهد عملکرد متکی باشد.

پژوهشگر باید بتواند:

- ☐ محدودیت‌های ابزار را تشخیص دهد.
- ☐ ورودی مناسب و روشن فراهم کند.
- ☐ خطاهای مهم خروجی را شناسایی کند.
- ☐ موارد خارج از تخصص خود را به فرد مناسب ارجاع دهد.

توان استفاده از رابط کاربری با توان ارزیابی علمی خروجی متفاوت است.

- ❑ مصرف منابع باید با فایده پژوهشی تناسب داشته باشد.
- ❑ اعضای تیم به آموزش و ابزار مناسب دسترسی داشته باشند.
- ❑ وابستگی به یک خدمت یا حساب شخصی مدیریت شود.
- ❑ هزینه و محدودیت دسترسی، مشارکت علمی را بی دلیل محدود نکند.

پایداری، هم بُعد محیط‌زیستی دارد و هم بُعد سازمانی.

نقشه استفاده از AI در فرایند پژوهش

در هر مرحله، یک پرسش اخلاقی برجسته می‌شود:

□ ایده‌پردازی: آیا ابزار دیدگاه ما را محدود می‌کند؟

□ مرور منابع: آیا شواهد درست و کامل بازتاب یافته‌اند؟

□ تحلیل: آیا روش و کد قابل بررسی‌اند؟

□ نگارش: آیا معنا، عدم قطعیت و منشأ مطالب حفظ شده‌اند؟

اصول اخلاقی ثابت‌اند؛ شکل کاربرد آن‌ها تغییر می‌کند.

- ☐ تناسب ابزار با زبان، رشته و نوع وظیفه بررسی شود.
- ☐ دسترسی به منابع و امکان مشاهده شواهد روشن باشد.
- ☐ شرایط حفاظت از اطلاعات شناخته شود.
- ☐ امکان ذخیره خروجی و بازبینی فرایند وجود داشته باشد.

محبوبیت یا قیمت ابزار، معیار کافی برای انتخاب پژوهشی نیست.

پیش از وارد کردن داده چه بررسییم؟

- ☐ آیا ورود این اطلاعات به این محیط مجاز است؟
- ☐ آیا همه ستون‌ها، صفحات یا جزئیات لازم‌اند؟
- ☐ ابزار داده را چگونه نگهداری و استفاده می‌کند؟
- ☐ آیا می‌توان با داده تجمیعی یا نمونه مصنوعی به هدف رسید؟

تصمیم درباره داده باید پیش از بارگذاری گرفته شود.

مثال فرضی: شرح حال یک بیمار با بیماری بسیار نادر در یک شهر کوچک.

❑ سن، تاریخ مراجعه، محل زندگی و بیماری نادر ممکن است در کنار هم فرد را مشخص کنند.

❑ متن آزاد، تصویر و فراداده فایل می توانند اطلاعات شناسایی کننده داشته باشند.

❑ حذف شناسه مستقیم، خطر بازشناسایی را همیشه از بین نمی برد.

❑ تناسب روش حفاظت باید با نوع داده بررسی شود.

□ پرامپت جهت‌دار:

توضیح بده چرا مداخله ما مؤثر بوده است

پرامپت مناسب‌تر:

با توجه به طراحی و نتایج، تفسیرهای ممکن، محدودیت‌ها و توضیح‌های جایگزین را بررسی کن

□ از ابزار درخواست توجیه نتیجه مطلوب نکنیم.

□ اطلاعات مخالف و عدم قطعیت را در درخواست بگنجانیم.

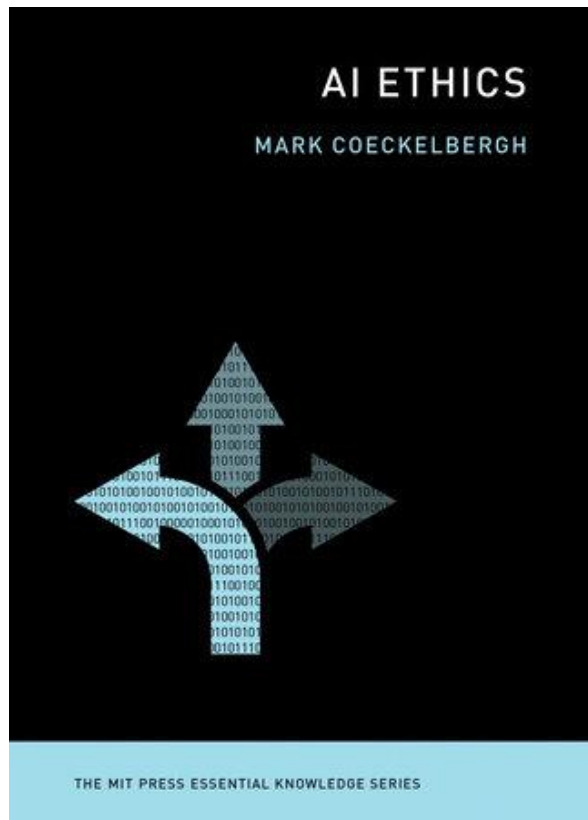
❑ کنترل ضروری: ادعاهای اثرگذار بر تصمیم با منبع یا روش مستقل بررسی شوند.

❑ متن روان ممکن است شامل اطلاعات ساختگی باشد.

❑ عدد، نقل‌قول یا منبع ظاهراً دقیق نیز ممکن است نادرست باشد.

❑ اطمینان در لحن، معادل اطمینان علمی نیست.

❑ خطاهای کوچک می‌توانند نتیجه نهایی را تغییر دهند.



- ❑ پرسش مکمل: چه یافته‌ای می‌تواند تفسیر فعلی ما را تضعیف کند؟
- ❑ پژوهشگر ممکن است پاسخ موافق فرضیه خود را ترجیح دهد.
- ❑ ابزار نیز ممکن است با ترجیح بیان‌شده کاربر همراه شود.
- ❑ تکرار درخواست تا دریافت پاسخ مطلوب، خطر سوگیری را افزایش می‌دهد.
- ❑ بررسی شواهد مخالف باید بخشی از فرایند باشد.

اعتماد بیش از حد و تضعیف قضاوت

نشانه‌های اعتماد نامتناسب:

❑ اقدام پیشنهادی: معیار پذیرش و مسئول بازبینی پیش از دریافت خروجی مشخص شوند.

❑ پذیرش پاسخ بدون خواندن منبع

❑ اجرای کد بدون فهم منطق آن

❑ تغییر تصمیم صرفاً به دلیل اطمینان ابزار

❑ نادیده گرفتن شواهد مخالف پاسخ AI



- ❑ AI می‌تواند در یافتن کلیدواژه و مترادف‌ها کمک کند.
- ❑ پوشش منابع و محدودیت دسترسی ابزار باید شناخته شود.
- ❑ در جست‌وجوهای جامع، اتکا به پاسخ چت‌بات به‌تنهایی کافی نیست.
- ❑ مسیر یافتن و انتخاب منابع باید قابل پیگیری باشد.

معیار پذیرش منبع، ارتباط و اعتبار آن برای سؤال پژوهش است.

خلاصه‌سازی و استخراج اطلاعات

- ☐ جمعیت، طراحی، پیامد و زمان اندازه‌گیری حفظ شوند.
- ☐ عددها و واحدها با متن یا جدول اصلی تطبیق یابند.
- ☐ اطلاعات ناموجود با حدس تکمیل نشوند.
- ☐ نتیجه مقاله از برداشت یا تفسیر ابزار جدا بماند.

استفاده مسئولانه در کدنویسی و تحلیل

❑ نوع متغیرها و روش برخورد با داده گمشده بررسی شوند.

❑ روش آماری باید با سؤال و طراحی پژوهش تناسب داشته باشد.

❑ کد پیش از اجرا روی نسخه کاری بازبینی شود.

❑ نتایج مهم با یک کنترل مستقل بررسی شوند.

تحلیل داده و روش‌های آماری



مسئولیت انتخاب روش تحلیل با تیم پژوهش باقی می‌ماند.

نگارش و ویرایش مقالات

- | ✓ مجاز | × غیر مجاز |
|-------------------------------|---|
| • ویرایش زبانی و بهبود ساختار | • تولید کامل متن علمی بدون مشارکت فکری |
| • ساده‌سازی متن پیچیده | • تولید بخش‌های کلیدی بدون اتکا به داده واقعی |
| • کمک به نگارش پیش‌نویس | • نگارش نتایج و بحث بدون تحلیل واقعی |

□ متن بازنویسی‌شده از نظر حفظ معنای علمی بررسی شود.

□ اصطلاحات تخصصی و شدت ادعاها کنترل شوند.

□ ابزار نباید اطلاعات تازه و بدون پشتوانه اضافه کند.

□ حقوق مؤلفان و منشأ مطالب رعایت شوند.

بهبود زبان نباید معنای یافته یا میزان اطمینان را تغییر دهد.

برخی ابزارها می‌توانند علاوه بر تولید پاسخ:

- ☐ پیشنهاد اجرایی: دسترسی حداقلی، ثبت اقدامات و تأیید انسانی برای اقدامات مهم در نظر گرفته شود.
- ☐ فایل بخوانند یا تغییر دهند.
- ☐ کد اجرا کنند.
- ☐ در منابع جست‌وجو کنند.
- ☐ اطلاعات را به خدمات دیگر منتقل کنند.



- ❑ طرح پژوهشی منتشرنشده ممکن است محرمانه باشد.
- ❑ متن داوری، مصاحبه و صورت جلسه نیز اطلاعات حساس دارند.
- ❑ ابزارهای ضبط و خلاصه سازی خودکار می توانند اطلاعات دیگران را پردازش کنند.
- ❑ استفاده از AI باید با حدود دسترسی و حقوق صاحبان اطلاعات سازگار باشد.

داده حساس فقط به پرونده بیمار محدود نیست.

- ❑ دامنه اثر خطا بر داده، تحلیل و نتیجه مشخص شود.
- ❑ استفاده از خروجی متأثر تا زمان بررسی محدود شود.
- ❑ نسخه‌های مرتبط اصلاح و افراد لازم آگاه شوند.
- ❑ علت خطا و راه جلوگیری از تکرار ثبت شود.

پاسخ‌گویی با اقدام اصلاحی قابل مشاهده همراه است.

- ❑ ابزار متناسب با پیچیدگی وظیفه انتخاب شود.
- ❑ درخواست‌های تکراری و بی‌هدف کاهش یابند.
- ❑ زمان لازم برای بررسی خروجی در برآورد کار لحاظ شود.
- ❑ راه ادامه فعالیت در صورت قطع خدمت یا دسترسی مشخص باشد.

صرفه‌جویی واقعی، کیفیت بررسی و اصلاح را نیز در نظر می‌گیرد.

وقتی اصول اخلاقی با هم تعارض پیدا می‌کنند

نمونه‌های تعارض:

❑ تصمیم قابل دفاع: گزینه‌ها، دلیل انتخاب و خطر باقی‌مانده روشن باشند.

❑ شفافیت بیشتر در برابر حفاظت از اطلاعات

❑ سرعت بالاتر در برابر فرصت بررسی انسانی

❑ دسترسی آسان‌تر در برابر کنترل داده

❑ کیفیت بیشتر در برابر هزینه و مصرف منابع

- ✓ هدف و حدود وظیفه روشن است؟
- ✓ ابزار با نوع کار و حساسیت داده تناسب دارد؟
- ✓ ورود این اطلاعات به محیط انتخاب شده مجاز است؟
- ✓ امکان بررسی مستقل خروجی وجود دارد؟
- ✓ مسئول بازبینی مشخص است؟

ابهام مهم باید پیش از شروع برطرف شود.

✓ درخواست بی طرفانه و محدود به وظیفه است؟

✓ فقط اطلاعات لازم وارد ابزار می شود؟

✓ منابع، عددها و ادعاهای مهم بررسی می شوند؟

✓ پاسخ های مخالف و عدم قطعیت دیده می شوند؟

✓ اقدامات خودکار تحت نظارت اند؟

پاسخ نامطمئن تا زمان بررسی، مبنای تصمیم نهایی قرار نگیرد.

چکلیست پس از استفاده

- ✓ خروجی نهایی از نظر صحت و سوگیری بازبینی شده است؟
- ✓ اصلاحات و دلیل پذیرش نتیجه قابل پیگیری اند؟
- ✓ نقش ابزار و مسئولیت انسانی روشن است؟
- ✓ نگهداری یا حذف داده با شرایط مجاز سازگار است؟
- ✓ محدودیت‌های مؤثر بر نتیجه بیان شده‌اند؟

این چکلیست باید با سیاست‌های سازمان تطبیق داده شود.

یک تصمیم برای پژوهش بعدی

پیش از استفاده بعدی از AI، این جمله‌ها را کامل کنیم:

از این ابزار برای ... استفاده می‌کنم

مهم‌ترین خطر این کاربرد ... است

خروجی را با ... بررسی می‌کنم

مسئول پذیرش و اصلاح نتیجه ... است

نتیجه مطلوب: استفاده‌ای که دلیل، حدود و مسئول آن مشخص است.