



مرور قوانین، استانداردها و الزامات استفاده از هوش مصنوعی در پژوهش های سلامت

از استفاده روزمره پژوهشگران تا توسعه و استقرار سامانه های هوش مصنوعی پزشکی

حسین ربانی

رئیس مرکز تحقیقات پردازش تصویر و سیگنال پزشکی
دانشگاه علوم پزشکی اصفهان

Rabbani.h@ieee.org

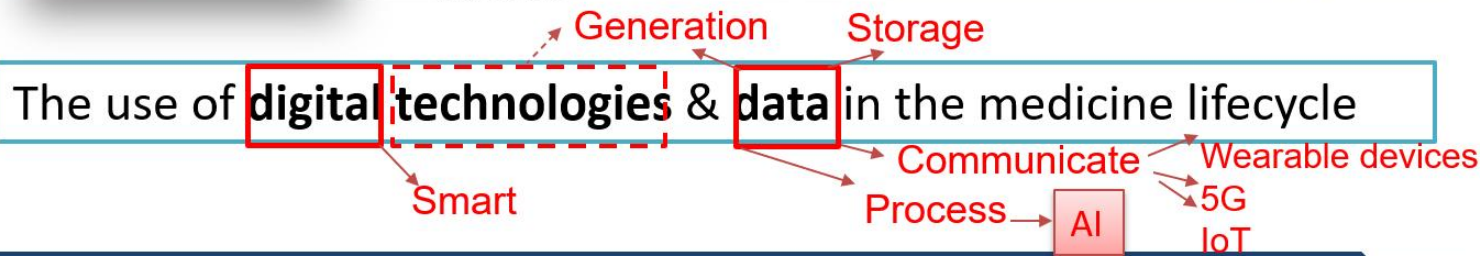
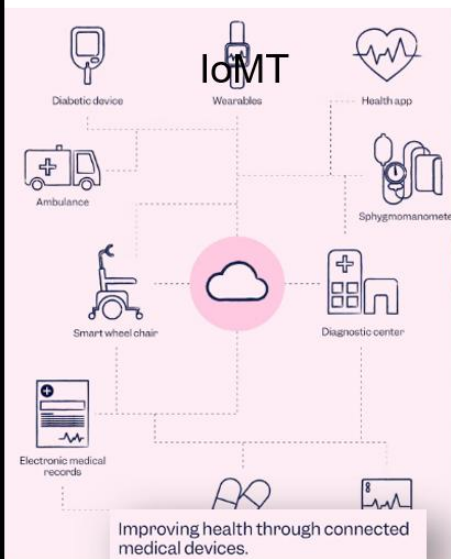
Profiles.mui.ac.ir/hosseini-rabbani

مرداد ۱۴۰۵





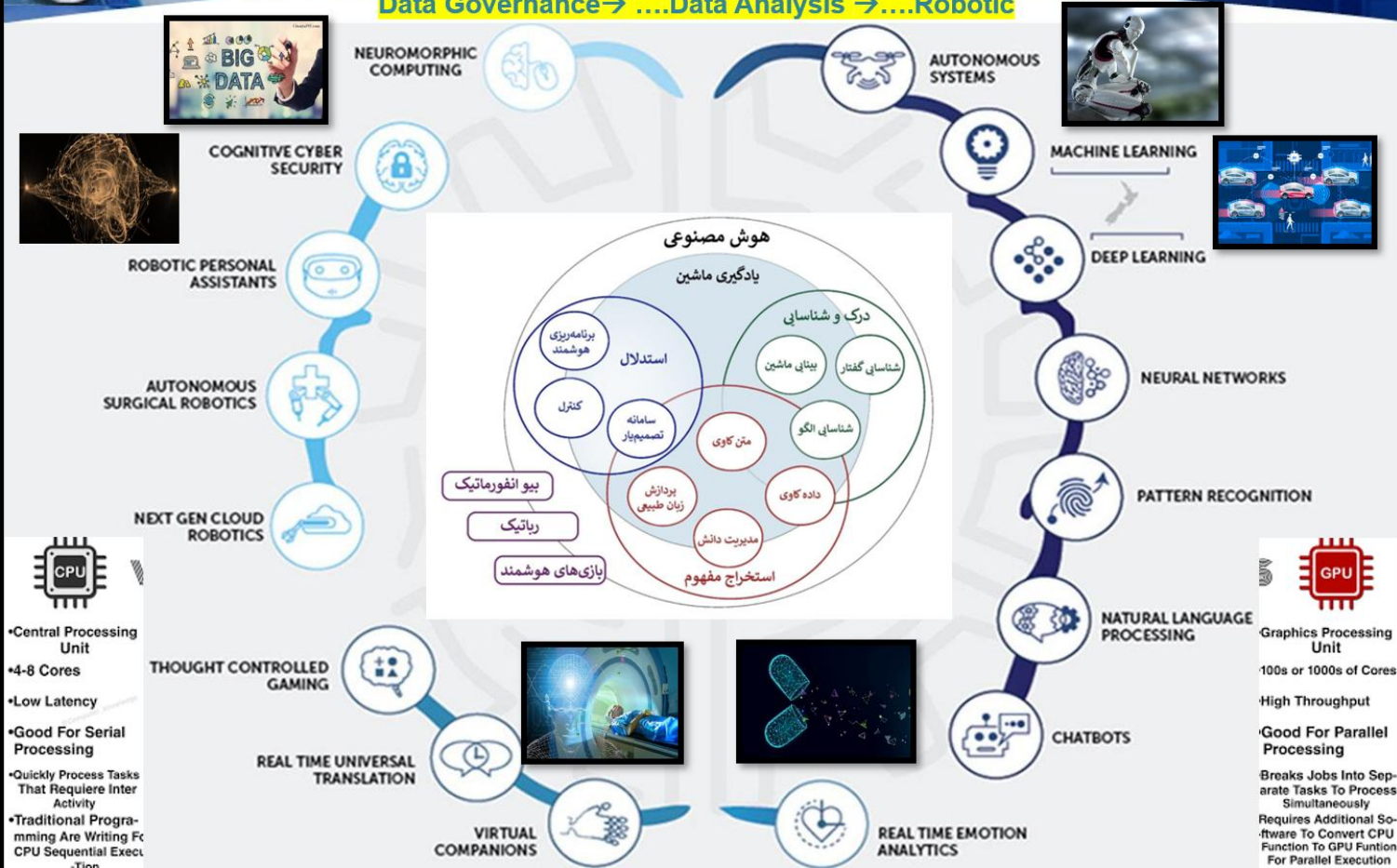
Why Smart Health is Important Today?





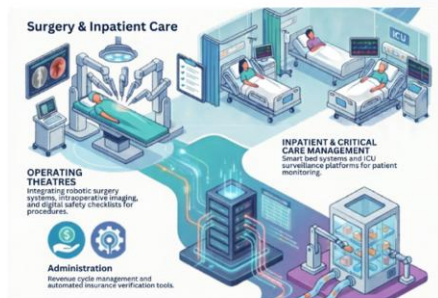
Artificial Intelligence

Data Governance → ...Data Analysis → ...Robotics

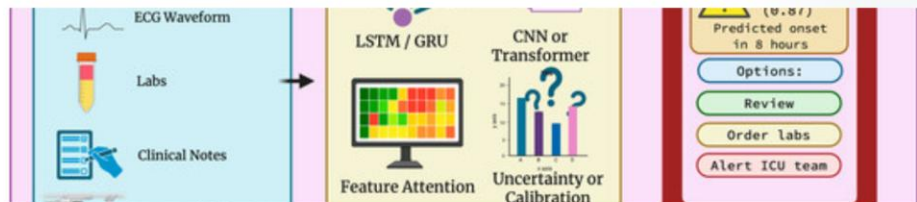


CLINICAL ECOSYSTEM

- ✓ Radiology: Automated Image Triaging.
- ✓ Pathology: Digital Slide Analysis.
- ✓ Surgery: Robotic Navigation.
- ✓ ICU: Predictive Warning Systems.
- ✓ Wearables: Continuous Patient Monitoring.



THE CLINICAL WORKFLOW



Data Input

Patient vitals, images, or lab results.

AI Processing

Risk scores, heatmap generation.

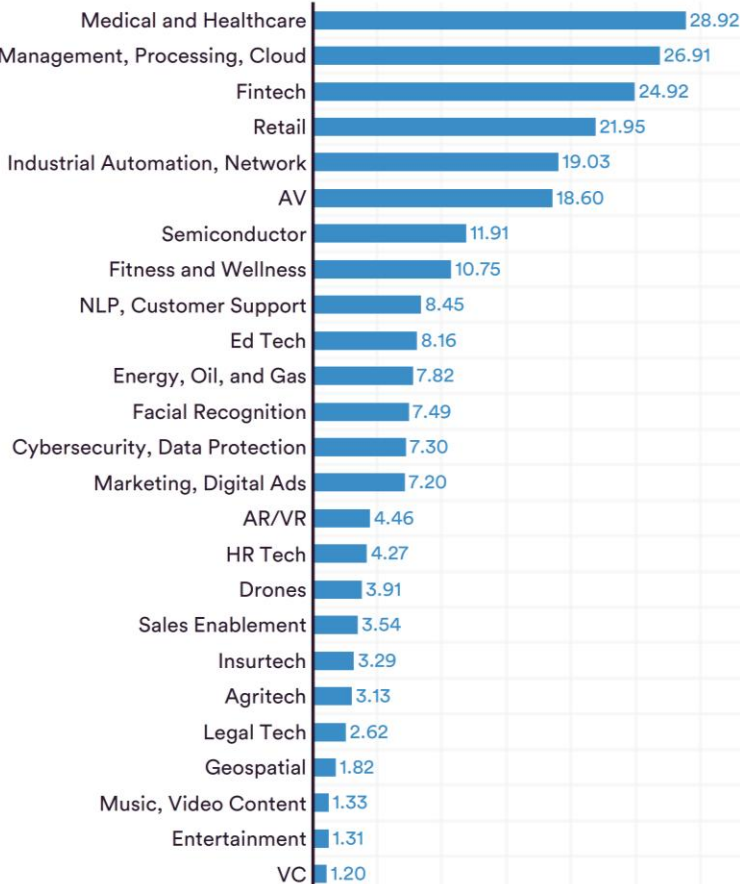
Human Choice

Final diagnostic decision & treatment.

AI ASSISTS, IT DOES NOT REPLACE.

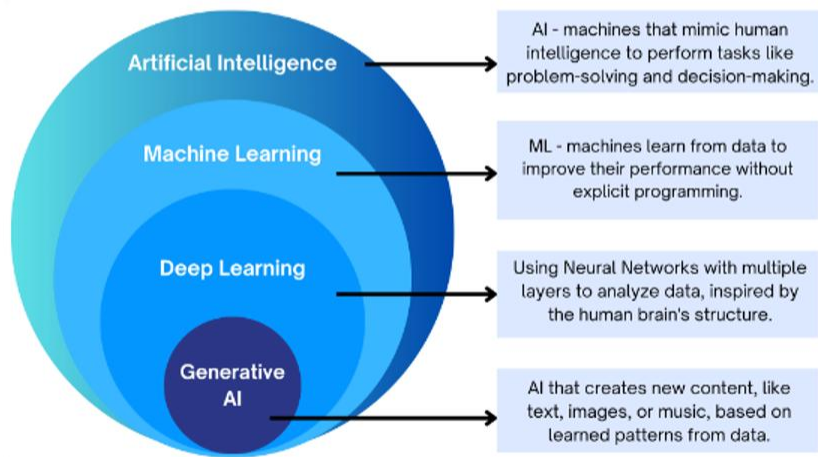
PRIVATE INVESTMENT in AI by FOCUS AREA, 2017–21 (SUM)

Source: NetBase Quid, 2021 | Chart: 2022 AI Index Report

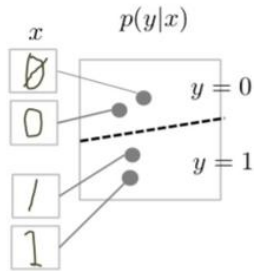


Total Investment (in billions of U.S. Dollars)

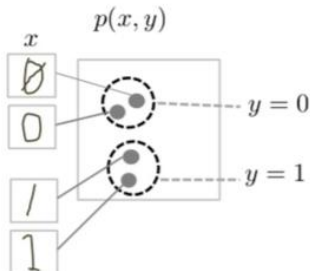
DEFINING THE TERMS



• Discriminative Model



• Generative Model



- ✓ **Artificial Intelligence**: Broad field of machines mimicking intelligence.
- ✓ **Machine Learning**: Systems that learn patterns from data.
- ✓ **Deep Learning**: Multi-layered neural networks (Bio-inspired).
- ✓ **Generative AI**: AI that creates new content (LLMs).

آیا برای استفاده از هوش مصنوعی در پژوهش‌های سلامت قانون مشخصی وجود دارد؟

نوشتن مقاله؟

داوری مقاله؟

تولید شکل؟

استفاده از داده بیماران؟

توسعه سامانه AI؟

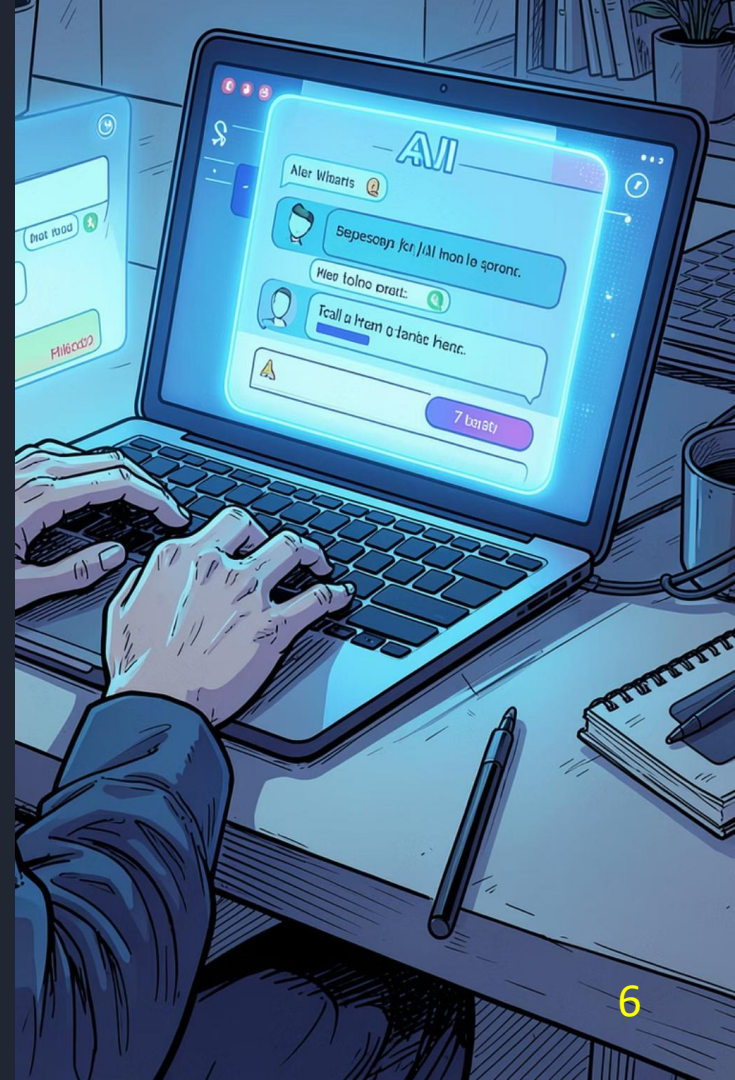
استقرار بالینی؟



همه اینها امروز قانون دارند.

استفاده از AI در پژوهش

فعالیت	مجاز؟
اصلاح زبان	✓
خلاصه سازی	✓
جستجوی منابع	✓
تحلیل داده	✓
تولید شکل شماتیک	✓
تولید داده جعلی	✗
جعل رفرنس	✗
نویسندگی مقاله	✗





ناشران بزرگ چه می گویند؟

همه تقریباً روی سه اصل توافق دارند:

Transparency ✓

Human Responsibility ✓

AI در داوری مقاله

Can reviewers use ChatGPT?

بهبود نگارش Review ✓

Upload manuscript to public AI ✗

افشای اطلاعات محرمانه ✗

AI در تولید شکل و تصویر علمی

Disclosure ✓

AI-generated Images in Health Research

Illustrative Scientific Graphics	AI-assisted Medical Images	Synthetic Scientific Evidence
Example: Graphical Abstract, Workflow Diagram, Educational Illustration, Concept Figure Examples (Medical Focus): - Medical Graphical Abstract for a publication. - AI lifecycle diagram Guidelines: ✓ Generally acceptable ✓ Disclosure recommended	Example: Image enhancement, Segmentation, Denoising, Super-resolution Examples (Medical Focus): - AI-based OCT denoising, CT enhancement. Guidelines: ⚠ Requires validation ⚠ Should preserve clinical information	Example: Fake pathology image, Fake microscopy, Fabricated radiology, Generated experimental results Examples (Medical Focus): - Fake pathology slide, fabricated CT scan to prove a scientific result. Guidelines: ⚠ Scientifically unacceptable ⚠ Research misconduct
Illustration	Medical Support	Scientific Evidence
Principle: AI may assist scientific communication, but it must never fabricate scientific evidence.		

AI-generated Images in Health Research

● Illustrative Scientific Graphics

Example: Graphical Abstract, Workflow Diagram, Educational Illustration, Concept Figure

Examples (Medical Focus):



Medical Graphical Abstract for a publication, AI lifecycle diagram

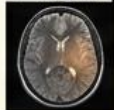
Guidelines:

- ✓ Generally acceptable
- ✓ Disclosure recommended

● AI-assisted Medical Images

Example: Image enhancement, Segmentation, Denoising, Super-resolution

Examples (Medical Focus):



AI-based OCT denoising, CT enhancement.

Guidelines:

- ⚠ Requires validation
- ⚠ Should preserve clinical information

● Synthetic Scientific Evidence

Example: Fake pathology image, Fake microscopy, Fabricated radiology, Generated experimental results

Examples (Medical Focus):



Fake pathology slide, fabricated CT scan to prove a scientific result.

Guidelines:

- ✗ Scientifically unacceptable
- ✗ Research misconduct



Illustration



Medical Support



Scientific Evidence

Principle: AI may assist scientific communication, but it must never fabricate scientific evidence.

ناشران بزرگ چه می گویند؟

همه تقریباً روی سه اصل توافق دارند:

Transparency ✓

Human Responsibility ✓

Disclosure ✓

Can reviewers use ChatGPT?

بهبود نگارش Review ✓

Upload manuscript to public AI ✗

افشای اطلاعات محرمانه ✗

Confidential manuscripts should **not** be uploaded to public AI systems unless explicitly permitted by the journal or institution (see the last slides).

AI-generated Images in Health Research

Illustrative Scientific Graphics	AI-assisted Medical Images	Synthetic Scientific Evidence
Example: Graphical Abstract, Workflow Diagram, Educational Illustration, Concept Figure	Example: Image enhancement, Segmentation, Denoising, Super-resolution	Example: Fake pathology image, Fake microscopy, Fabricated radiology, Generated experimental results
Examples (Medical Focus): Medical Graphical Abstract for a publication, AI lifecycle diagram	Examples (Medical Focus): AI-based OCT denoising, CT enhancement	Examples (Medical Focus): Fake pathology slide, fabricated CT scan to prove a scientific result.
Guidelines: ✓ Generally acceptable ✓ Disclosure recommended	Guidelines: ⚠ Requires validation ⚠ Should preserve clinical information	Guidelines: ⚠ Scientifically unacceptable ⚠ Research misconduct
Illustration	Medical Support	Scientific Evidence
Principle: AI may assist scientific communication, but it must never fabricate scientific evidence.		

AI Assistant

≠

AI Author

مسئولیت همیشه با پژوهشگر است

International Principles
(WHO, UNESCO, OECD)

Country / Government
(EU, Australia, UK, Iran, ...)

Funding Agencies
(NIH, Horizon Europe, ...)

University / Institution
(Harvard, Oxford, ...)

Journal / Publisher
(Nature, IEEE, Elsevier...)

Researcher

Who Defines the Rules?

Different organizations regulate different parts of AI research

موضوع	مرجع
قانون	Government
اخلاق پژوهش	University
الزامات گرنت	Funding Agency
انتشار مقاله	Journal
تصویب مطالعه	IRB / Ethics Committee
محصول پزشکی	Regulatory Agency

The researcher is responsible for complying with all applicable regulations.



Framework for Review of Clinical Research Involving Artificial Intelligence



نمونه‌های قوانین کشورها

پژوهشگر	EU Living Guideline
کمیته بررسی Clinical Research	Harvard Framework
کمیته اخلاق پژوهش	Australia Guideline



Artificial intelligence (AI) in research

The rise of Artificial Intelligence (AI) is producing transformational change in almost all sectors of society. The use of AI tools in research is rapidly increasing to the point where they will soon become ubiquitous. Currently, many of the most used tools are Large Language Models (LLM), a machine learning model trained on vast amounts of text. AI tools, including LLMs, are increasingly built into software that is routinely used by researchers, such as search engines (e.g. Google). Many researchers use widely available LLM interfaces, referred to as chatbots (e.g. ChatGPT) to assist in a range of tasks including transcription of audio data or drafting text (e.g. for development of participant information sheets, interview questions, or publications). Others use LLMs for reviewing literature, hypothesis generation, and data analysis. In many of these uses, AI is taking on a role analogous to a research assistant.



وزارت بهداشت، درمان و آموزش پزشکی
معاونت تحقیقات و فناوری



شماره: ۱۶۰۳/۱۳۴۰
تاریخ: ۱۴۰۳/۰۳/۰۵
پوست: دارد
جهش تولید با مشارکت مردم
(مقام معظم رهبری)

- فاقد اعتبار قانونی است پس نمی تواند حق چاپ یک اثر علمی را داشته باشد.
- در صورتی که مورد شکایت قرار گیرد، نمی تواند اثر علمی-پژوهشی را تایید نماید.

۴- داوران نباید از فناوری‌های هوش مصنوعی (AI) یا (AI-assisted) برای کمک به بازبینی علمی یک دست‌نوشته استفاده کنند، زیرا:

- تفکر انتقادی و ارزیابی اصلی مورد نیاز برای بررسی همتایان خارج از محدوده این فناوری است و این خطر وجود دارد که این فناوری نتایج نادرست، ناقص و ... در مورد دست‌نوشته ایجاد کند.
- ممکن است محرمانگی و حقوق مالکیت نویسندگان را نقض کند و در مواردی که دست‌نوشته حاوی اطلاعات قابل شناسایی شخصی باشد، ممکن است حقوق حریم خصوصی داده‌ها را رعایت نکرده و باعث انتشار عمدی/سپوی اطلاعات حساسی گردد.
- در نهایت، داوران هم‌اکنون مسئول و پاسخگوی اصالت، صحت و درستی نظرات خود هستند.

۵- اطلاعات محرمانه و داده های شخصی شرکت کنندگان در پژوهش/بیماران نباید با ابزارهای هوش مصنوعی به اشتراک گذاشته شود.

۶- ابزارهای مبتنی بر هوش مصنوعی صرفاً می توانند در مراحل مختلف طرح تحقیقاتی/ پایان نامه به عنوان یک دستیار کمکی بوده و هرگز نمی توانند جایگزین دانش و بینش فرد گردند. پس کلیه نویسندگان در قبال مطالبی که توسط هوش مصنوعی در اثر پژوهشی ارائه می‌شود(از جمله: دقت مطالب ارائه شده، عدم سرقت ادبی، ذکر منابع، اعتبار منابع و...) مسئولیت دارند. به عبارتی، مسئولیت اطمینان از درستی و مطابقت یک اثر پژوهشی با هنجارهای اخلاقی (حتی قسمت‌هایی که توسط هوش مصنوعی نوشته شده است) با نویسندگان است.

مستدعی است دستور فرمایید مراتب به نحو مقتضی به اطلاع کلیه افراد ذینفع رسانده شود.

دکتر محبی پارسا
مدیر کلیه کشوری اخلاق
در پژوهش های زیست پزشکی



وزارت بهداشت، درمان و آموزش پزشکی
معاونت تحقیقات و فناوری



شماره: ۱۶۰۳/۱۳۴۰
تاریخ: ۱۴۰۳/۰۳/۰۵
پوست: دارد
جهش تولید با مشارکت مردم
(مقام معظم رهبری)

کلیه معاونین محترم تحقیقات و فناوری دانشگاه / دانشکده های علوم پزشکی

معاون محترم علوم پزشکی دانشگاه آزاد اسلامی

معاون محترم تحقیقات، فناوری و آموزش انستیتو پاستور ایران

معاون محترم پژوهشی انستیتو آموزشی، تحقیقاتی قلب و عروق شهید رجایی

معاون محترم پژوهش و فناوری جهاد دانشگاهی

با سلام و احترام

پیرو بخشنامه شماره ۷۰۰/۷۸۹۹ تاریخ ۱۴۰۲/۰۸/۰۷ در خصوص ضرورت رعایت ملاحظات اخلاقی- علمی در استفاده از هوش مصنوعی جهت نگارش آثار پژوهشی تا زمان تدوین و ابلاغ راهنمای کامل "ملاحظات اخلاقی نحوه استفاده از هوش مصنوعی در تهیه آثار پژوهشی"، موارد زیر را به استحضار می‌رساند:

۱- فناوری هوش مصنوعی شامل سیستم‌هایی است که بر اساس برنامه‌های زبان برنامه‌نویسی پیشرفته بنا شده‌اند و می‌توانند برای سازماندهی و خلاصه‌نویسی منابع پژوهشی، تولید نمودار و گرافیک‌های پژوهشی، تحلیل داده، تولید متن، تصاویر، داده‌های مصنوعی، شبیه‌سازی، کد نویسی و سایر موارد استفاده شوند.

۲- استفاده از هوش مصنوعی در تهیه اثر پژوهشی، هیچ مشکل اخلاقی ذاتی ندارد، مشروط بر آن‌که به طور مناسب و اخلاقی استفاده شود و نویسندگان ماهیت هرگونه تعامل و نحوه استفاده از آن را به صورت شفاف با درج دقیق منبع آن در طرحنامه، مقاله و... مشخص کنند. از جمله اینکه:

- در صورت استفاده از ابزارهای هوش مصنوعی جهت تهیه محتوا، ضروری است نویسنده هم در متن (ارجاع درون‌متنی) و هم در لیست فهرست منابع انتهای طرح تحقیقاتی/پایان نامه (ارجاع پایان‌متنی) ملاحظات صحیح در ارجاع‌دهی را رعایت نماید.
- نویسندگان باید مشخص نمایند که کدام مدل از هوش مصنوعی، در چه زمانی و توسط چه کسی استفاده شده است.
- نویسندگان باید نوع ابزارهای هوش مصنوعی و ابزارهای ماشینی، مشخصات فنی مانند: نام کامل ابزار، نسخه و مدل آن که در تهیه اثر پژوهشی خود استفاده نموده‌اند را، مشخص نمایند.
- نویسندگان باید نحوه استفاده از ابزار هوش مصنوعی و این که کدام قسمت از محتوای تولید شده مقاله توسط هوش مصنوعی نوشته شده است را شرح دهند. به عنوان مثال، اگر از هوش مصنوعی برای کمک به نوشتن استفاده شده، این مورد را در بخش "تقدیر و تشکر" توضیح دهند، یا اگر از هوش مصنوعی برای جمع‌آوری داده‌ها، تجزیه و تحلیل یا تولید شکل استفاده شده، نویسندگان باید این کاربرد را در روش‌ها بصورت شفاف توضیح دهند.

۳- ابزار هوش مصنوعی نباید به عنوان نویسنده در نظر گرفته شود، زیرا:

دو کاربرد مختلف استفاده از AI در پژوهشهای سلامت

1. استفاده از AI برای انجام پژوهش (Research Assistant)

2. پژوهش روی خود AI و توسعه سامانه AI پزشکی (AI as Medical Technology)

قوانین

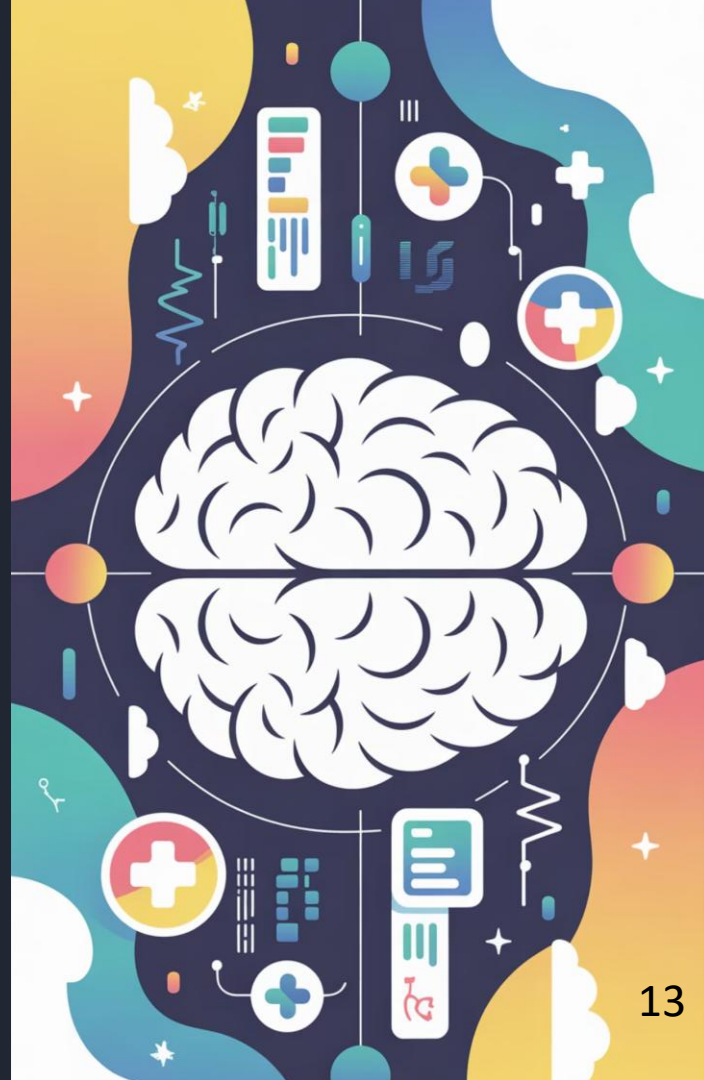
استفاده

Publisher, University, Ethics, Funding Agency

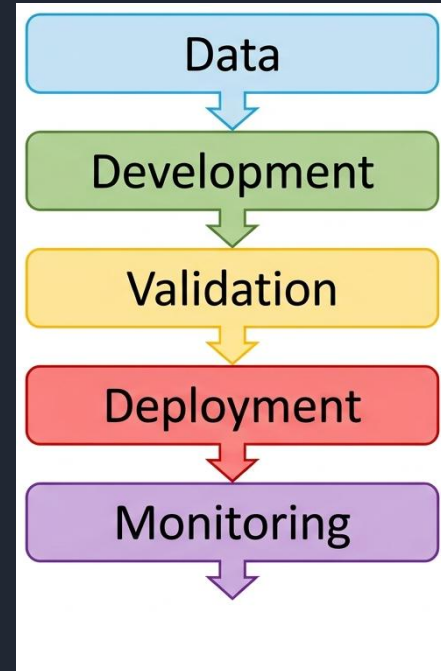
AI as a Research Assistant

FDA, WHO, ISO, IEC, EU AI Act

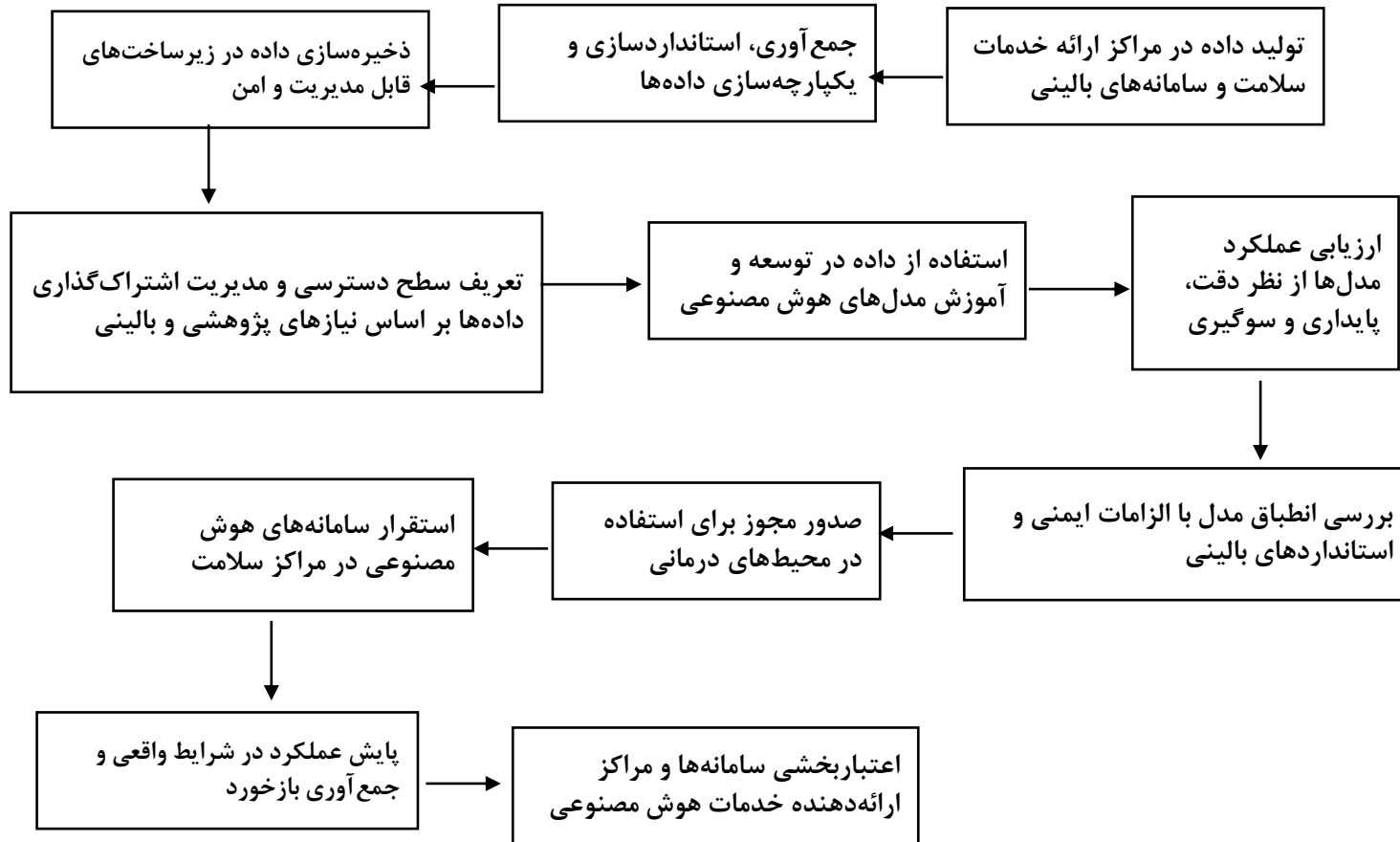
AI as a Medical Product



اگر بخواهیم یک سامانه هوش مصنوعی پزشکی
توسعه دهیم، دنیا چه قوانینی وضع کرده است؟



چرخه عمر داده، از تولید داده تا استفاده از سامانه های هوش مصنوعی



Landscape of AI Standards

استانداردها	حوزه
WHO, OECD	Ethics
ISO/IEC 42001	Governance
ISO/IEC 23894	Risk
FAIR, GDPR, EHDS	Data
ISO/IEC 23053, ISO/IEC 5338	Development
CONSORT-AI, SPIRIT-AI, STARD-AI, TRIPOD+AI, PROBAST+AI	Validation
FDA, EU AI Act	Deployment



Global Ethical Principles for Trustworthy AI



WHO

HEALTHCARE FOCUS

Guidance on Ethics & Governance of AI for Health (2021)

Protect Autonomy

Patient Safety

Transparency

Promote Equity

Accountability



OECD

POLICY & GOVERNANCE

OECD Recommendation on AI Principles (2019)

Human-Centered

Security & Safety

Robustness

Transparency

Accountability



UNESCO

SOCIETY & HUMAN RIGHTS

Recommendation on the Ethics of AI (2021)

Human Rights

Human Dignity

Diversity & Inclusion

Sustainability

Peaceful Society



Trustworthy AI

Global Convergence Core (نقطه همگرایی جهانی)



Human Rights & Oversight



Safety & Robustness



Transparency



Fairness & Equity



Accountability



Strategic Core Takeaway: Although developed by distinct international bodies for Health, Policy, and Society, these frameworks converge on a shared mandate: **AI must remain human-centered, safe, transparent, and accountable.** Standard operational frameworks like **ISO/IEC 42001** and **ISO 23894** act as the direct practical implementation of these ethics.

The Strategic Governance Shift

Transitioning from Technical AI Implementation to Organization-Centric Governance



Traditional Technical AI

Algorithm-Centric Focus

Data Collection



Algorithm / Model Training



Validation



Prediction / Deployment



Limited to software accuracy & technical performance.



EVOLUTION



ISO/IEC 42001 (AIMS)

Organization-Centric
Governance

Leadership & Institutional Context



Risk & Ethical Governance



Managed AI Life Cycle (Data → Model → Deployment)



Continuous Audit & Monitoring



Organizational Improvement



Guarantees institutional compliance, safety, & accountability.



Risk Identification

Risk Analysis

Risk Evaluation

Risk Treatment

Monitoring

ISO/IEC 23894 – AI Risk Management

Data Risks ●

- **Data Quality**: داده‌های ناقص یا کم‌کیفیت، عملکرد AI را کاهش می‌دهند.
- **Bias**: سوگیری موجود در داده به تصمیم‌های ناعادلانه منجر می‌شود.
- **Privacy**: احتمال افشای اطلاعات محرمانه بیماران.

Model Risks ●

- **Hallucination**: تولید پاسخ یا اطلاعات نادرست اما متقاعدکننده.
- **Model Drift**: افت عملکرد مدل پس از استقرار به دلیل تغییر شرایط.
- **Transparency**: دشواری در توضیح نحوه تصمیم‌گیری مدل.

Operational Risks ●

- **Security**: حملات یا دستکاری مدل و داده‌ها.
- **Automation Bias**: اتکای بیش از حد کاربران به خروجی AI.
- **Safety**: احتمال آسیب به بیمار در اثر تصمیم اشتباه.

Governance Risks ●

- **Accountability**: مشخص نبودن مسئولیت در برابر خطاها و پیامدهای AI.



استانداردهای گزارش پژوهش AI

جدول

استاندارد

کاربرد

طراحی مطالعه

SPIRIT-AI

RCT

CONSORT-AI

Prediction

TRIPOD+AI

Bias

PROBAST+AI

Diagnostic

STARD-AI

Imaging AI

CLAIM

Clinical AI

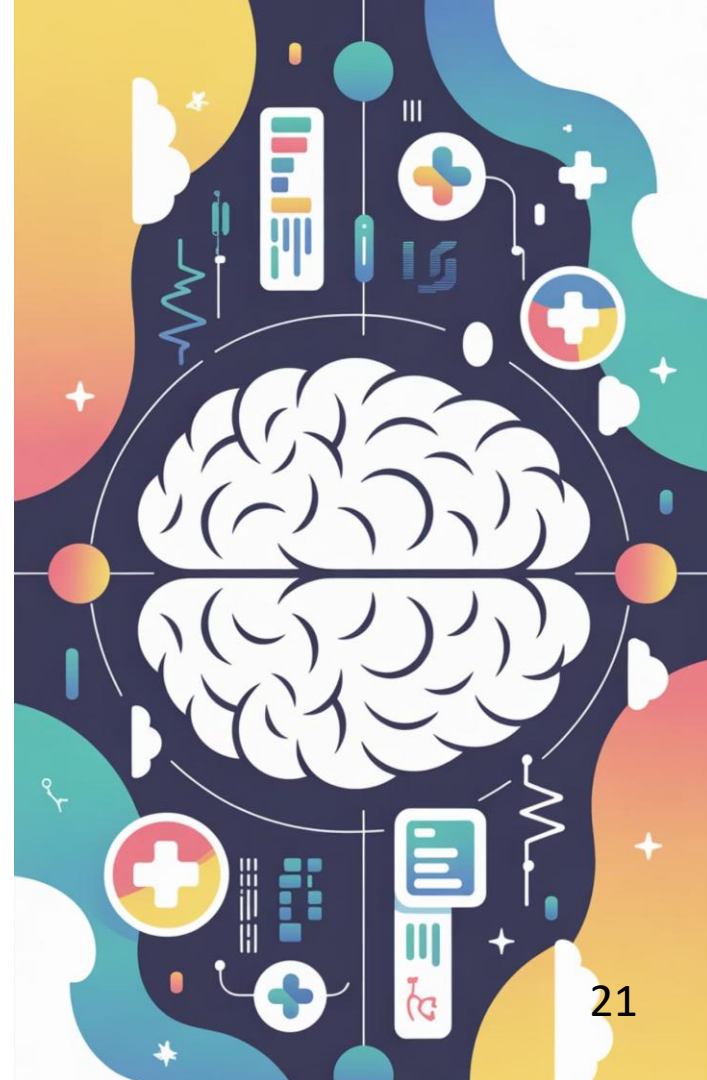
CHART



مطالعات مبتنی بر LLM

TRIPOD-LLM

ابزار ارزیابی سوگیری
نه راهنمای گزارش دهی



GLOBAL REGULATORY AGENCIES



U.S. Food and Drug
Administration



European Medicines
Agency



Medicines and Healthcare
products Regulatory Agency



Health
Canada

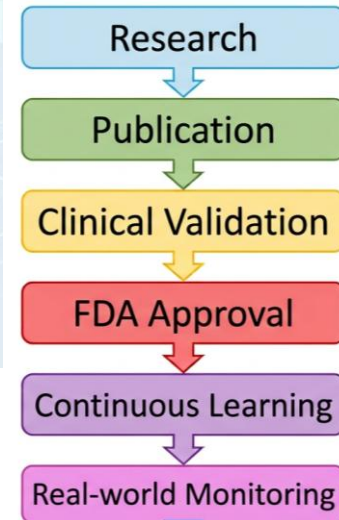


Therapeutic Goods
Administration

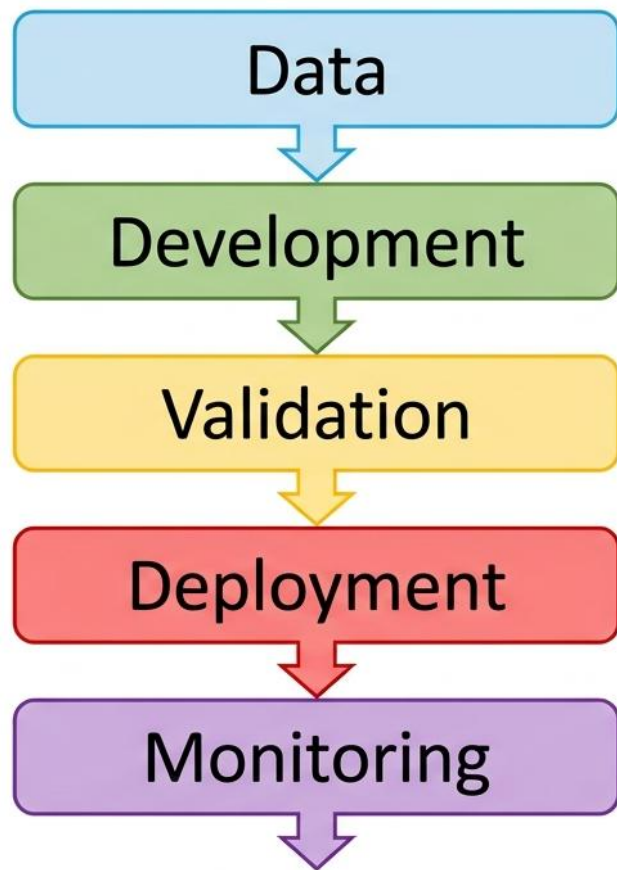
تقریباً تمام رگولاتورها از یک منطق مشترک پیروی می کنند:



FDA



Unlike conventional software, AI systems continue to evolve after deployment. Therefore, regulation must extend beyond market approval.



استاندارد	
FAIR	اصولی برای قابل یافتن، قابل دسترسی، قابل تبادل و قابل استفاده مجدد بودن داده‌های پژوهشی.
GDPR	مقررات اتحادیه اروپا برای حفاظت از حریم خصوصی و حقوق افراد نسبت به داده‌های شخصی.
EHDS	چارچوب اتحادیه اروپا برای اشتراک‌گذاری ایمن و استاندارد داده‌های سلامت در مراقبت، پژوهش و نوآوری.
HL7 FHIR	استاندارد بین‌المللی برای تبادل الکترونیکی اطلاعات سلامت بین سامانه‌های مختلف.
DICOM	استاندارد جهانی ذخیره، انتقال و مدیریت تصاویر پزشکی و اطلاعات مرتبط با آن‌ها.

اقدامات بین‌المللی در حکمرانی AI سلامت

WHO

اصول اخلاقی AI در سلامت

ISO/IEC 42001

مدیریت سامانه‌های AI

ISO/IEC 23894

مدیریت ریسک AI

FDA

چارچوب ارزیابی محصولات AI پزشکی

European AI Act

طبقه‌بندی ریسک سامانه‌های AI

کشورهای پیشرو: آمریکا، اتحادیه اروپا، انگلستان، چین، کره جنوبی، سنگاپور

① تقریباً تمام کشورها به سمت «حکمرانی داده + مدیریت ریسک AI» حرکت کرده‌اند.

اسناد و قوانین موجود در ایران

→ قوانین مجلس و سند ملی هوش مصنوعی

→ اسناد وزارت بهداشت در حوزه سلامت دیجیتال و AI

→ ضوابط کمیته ملی اخلاق در پژوهش‌های زیست‌پزشکی

→ اسناد دانشگاه‌های علوم پزشکی در حوزه AI سلامت

→ سند استانداردها و الزامات اعتباربخشی محصولات AI سلامت،
سند اخلاق AI و سند هوش مصنوعی در نظام سلامت

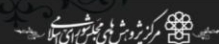
اسناد بالادستی وجود دارند، اما ابزار اجرایی یکپارچه برای مدیریت پژوهش‌ها و داده‌های AI سلامت هنوز محدود است.



معاونت علمی، فناوری و اقتصاد دانش بنیان رئیس جمهور
وزارت ارتباطات و فناوری اطلاعات - وزارت علوم، تحقیقات و فناوری - وزارت آموزش و پرورش
وزارت صنعت، معدن و تجارت - وزارت دفاع و پشتیبانی نیروهای مسلح
وزارت بهداشت، درمان و آموزش پزشکی - وزارت اطلاعات - وزارت جهاد کشاورزی
سازمان برنامه و بودجه کشور - سازمان اداری و استخدامی کشور
سازمان صدا و سیما جمهوری اسلامی ایران - ستاد کل نیروهای مسلح
دبیرخانه شورای عالی انقلاب فرهنگی - دبیرخانه شورای عالی فضای مجازی
سندوق نوآوری و شکوفایی

هیئت وزیران در جلسه ۱۴۰۲/۲/۲۴ به پیشنهاد معاونت علمی، فناوری و اقتصاد دانش بنیان
رئیس جمهور و به استناد اصل یکصد و سی و هشتم قانون اساسی جمهوری اسلامی ایران،
سند ستاد توسعه فناوری و کاربرد هوش مصنوعی را به شرح زیر تصویب کرد:

سند ستاد توسعه فناوری و کاربرد هوش مصنوعی



شماره ۱۴۰۳/۴/۱۰

شماره ۱۴۰۳/۶/۱۴ د.ش

مضوبه تشکیل شورای ملی راهبردی و تأسیس سازمان ملی هوش مصنوعی

(مضوب جلسه ۹۰۱ مورخ ۱۴۰۳/۳/۲۹ شورای عالی انقلاب فرهنگی)

نهاد ریاست جمهوری - معاونت علمی، فناوری و اقتصاد دانش بنیان ریاست جمهوری

وزارت ارتباطات و فناوری اطلاعات - مرکز ملی فضای مجازی

سازمان برنامه و بودجه کشور - سازمان اداری و استخدامی کشور

ستاد علم و فناوری دبیرخانه شورای عالی انقلاب فرهنگی

ماده واحده «تشکیل شورای ملی راهبردی و تأسیس سازمان ملی هوش مصنوعی» که در جلسه ۹۰۱ مورخ ۱۴۰۳/۳/۲۹ شورای عالی انقلاب فرهنگی به تصویب رسیده است؛
به شرح ذیل برای اجرا ابلاغ می گردد:

«ماده واحده - ضمن تأیید کلیات سند ملی هوش مصنوعی و واگذاری تصویب مفاد آن به شورای معین شورای عالی انقلاب فرهنگی؛ به منظور برنامه ریزی، راهبردی، هماهنگی و نظارت بر حسن اجرای این سند شامل: تصویب نقشه راه اجرایی سازی سند، تصویب آیین نامه ها و دستورالعمل های لازم جهت تسهیل اجرای سند، تصویب ره نگاشت کاربرست هوش مصنوعی در دستگاه های اجرایی مستقل ملی و پیشنهاد روزآمدسازی سند؛ «شورای ملی راهبردی هوش مصنوعی»، با ترکیب ذیل تشکیل می گردد.



جمهوری اسلامی ایران
رئیس جمهور

شماره ۳۶۶۲۹
تاریخ ۱۴۰۴/۲/۲۱
تصویب نام ریاست وزیران

تیسره ۲- رئیس دبیرخانه ستاد به پیشنهاد دبیر ستاد و با حکم معاون اول رئیس جمهور منصوب می شود و وظایف اجرایی ستاد را در چهارچوب این سند به عهده دارد.
تیسره ۳- جلسات ستاد با حضور بیش از نیمی از اعضا رسمیت می یابد. مصوبات این ستاد برای تمامی دستگاه های اجرایی الزام آور است.

تیسره ۴- شیوه نامه رسیدگی به موضوعات در ستاد با پیشنهاد دبیرخانه به تصویب ستاد می رسد.
تیسره ۵- دبیرخانه ستاد موظف است گزارش عملکرد ستاد را در مقاطع زمانی شش ماهه به هیئت وزیران ارائه نماید.

۲-۵ وظایف دبیرخانه ستاد

۱-۲-۵ پیشنهاد روزآمدسازی سند بر اساس تحولات ملی و بین المللی در حوزه هوش مصنوعی جهت ارائه به ستاد و تصویب در هیئت وزیران.

۲-۲-۵ تدوین نقشه راه فعالیت های ستاد و ارائه آن به ستاد.

۳-۲-۵ تدوین برنامه های اجرایی در چهارچوب این سند.

۴-۲-۵ بررسی و تأیید برنامه های عملیاتی هوش مصنوعی دستگاه های اجرایی در چهارچوب مصوبات ستاد و برنامه ملی توسعه هوش مصنوعی و ارزیابی و پایش عملکرد دستگاه های اجرایی.

۵-۲-۵ تهیه، تدوین و پیشنهاد مقررات، ضوابط و لوایح مرتبط با هوش مصنوعی به مراجع ذی ربط و پیشنهاد اصلاح و به روزرسانی قوانین و مقررات مرتبط.

۶-۲-۵ شناسایی رویه ها و مقررات مانع توسعه هوش مصنوعی و ارائه پیشنهاد اصلاح و رفع موانع و پیگیری آن ها در مراجع ذی صلاح.

محمد رضا عارف
معاون اول رئیس جمهور

رونوشت به دفتر مقام معظم رهبری، دفتر رئیس جمهور، دفتر رئیس مجلس شورای اسلامی، دفتر رئیس قوه قضاییه، دفتر معاون اول رئیس جمهور، دبیرخانه مجمع تشخیص مصلحت نظام، دیوان محاسبات کشور، دیوان عدالت اداری، سازمان بازرسی کل کشور، معاونت های قوانین و نظارت مجلس شورای اسلامی، امور تدوین، تنقیح و انتشار قوانین و مقررات، کلیه وزارتخانه ها، سازمان ها و مؤسسات دولتی، معاونت های رئیس جمهور، نهادهای انقلاب اسلامی، روزنامه رسمی جمهوری اسلامی ایران، دبیرخانه شورای اطلاع رسانی دولت و دفتر هیئت دولت. (۱۹۸۹۶۳۰)

دستورالعمل استفاده از اطلاعات و داده‌های سلامت گردآوری شده توسط واحدهای ستادی وزارت بهداشت و دانشگاه های علوم پزشکی

این دستورالعمل به عنوان سند تکمیلی سایر ضوابط و دستورالعمل‌های کشوری مصوب، برای انجام پژوهش‌های علوم پزشکی و به منظور شفاف‌سازی و تسهیل فرایند دسترسی به اطلاعات و داده‌های سلامت با اهداف پژوهشی تعریف گردیده است.

1- کلیه اطلاعات و داده‌های مربوط به حوزه سلامت که با استفاده از بودجه عمومی دولت در قالب برنامه‌های جاری معلومت‌ها و ادارات ذیربط در ستاد وزارت بهداشت، درمان و آموزش پزشکی و واحدهای متناظر در دانشگاه‌های علوم پزشکی جمع‌آوری و ذخیره می‌گردد، می‌بایست در اختیار پژوهشگران اصلی که با پروپوزال تصویب شده در یکی از کمیته‌های اخلاق در پژوهش معتبر که اعتبار آن به دبیر کمیته ملی اخلاق در پژوهش های زیست پزشکی رسیده و شناسه ملی دریافت کرده باشد، مقتضای دریافت اطلاعات مذکور هستند؛ قرار داده شود. در مواردی که از ارائه اطلاعات خودداری می‌گردد لازم است که علت یا ذکر دلیل و به صورت مکتوب به مقتضای اعلام گردد.

تبصره: داده‌هایی که در قالب پژوهش های مصوب جمع آوری می‌شوند از شمول این دستورالعمل خارج بوده و انتشار داده‌ها بر اساس توافق حامیان و مجریان پژوهش خواهد بود.

2- حضور افراد به عنوان مدیر، کارشناس و ... در ادارات و نهادهای مذکور حتی مبتنی بر درج اساسی در زمره نویسندگان مقالات و یا سایر آثار علمی ایجاد نمی‌کند. هر گونه تمهید یا اجبار پژوهشگران برای قرار دادن نام یا نام‌هایی در فهرست نویسندگان که دارای معیارهای نویسندگی بر اساس آیین نامه اخلاق در انتشار آثار پژوهشی نباشند، سوءاستفاده از موقعیت شغلی بوده و مطلقاً ممنوع است و قابل پیگیری خواهد بود.

3- لازم است تمام اطلاعات مذکور به نحوی که به شکل غیرقابل برگشت بی‌نام شده باشند، در اختیار محققین قرار گیرد.

4- این دستورالعمل شامل اطلاعات محرمانه نمی‌باشد. طبقه‌بندی اطلاعات به عنوان "اطلاعات محرمانه" باید در زمان جمع‌آوری آن‌ها و توسط یک مرجع دارای صلاحیت و یا براساس قانون و مقررات باشد. در غیر این صورت ادعای محرمانه بودن اطلاعات نمی‌تواند مانعی از ارائه اطلاعات به پژوهشگران گردد.

قانون انتشار و دسترسی آزاد به اطلاعات
مصوب 1387/11/6
با اصلاحات و الحاقات بعدی

قانون مدیریت داده‌ها و اطلاعات ملی

مصوب 1401/6/30

پیوست 1- دستورالعمل یک‌صفحه‌ای: استفاده اخلاقی و ایمن از هوش مصنوعی در آموزش علوم پزشکی

پیوست 2- دستورالعمل یک‌صفحه‌ای: استفاده اخلاقی و ایمن از هوش مصنوعی در پژوهش های علوم پزشکی

پیوست 3- دستورالعمل یک‌صفحه‌ای: استفاده اخلاقی و ایمن حاکمیتی از هوش مصنوعی در نظام سلامت

شورای راهبری و سیاستگذاری هوش مصنوعی در نظام سلامت

فهرست و ارسی شاخص‌های اعتباربخشی محصولات مبتنی بر هوش مصنوعی در حوزه سلامت

شورای راهبری و سیاستگذاری هوش مصنوعی در نظام سلامت

سند ملی هوش مصنوعی در نظام سلامت جمهوری اسلامی ایران

مقایسه مدل‌های حکمرانی داده و هوش مصنوعی سلامت در کشورهای منتخب

کشور / منطقه	چارچوب‌های اصلی	نهادهای مسئول	رویکرد تنظیم‌گری	تمرکز اصلی	ویژگی‌های شاخص
اتحادیه اروپا	General Data Protection Regulation (GDPR), EU AI Act, European Health Data Space (EHDS)	European AI Office, Health Data Access Bodies	قانون محور و حقوق محور، مبتنی بر ریسک	حفاظت از حقوق فردی، استانداردسازی داده، توسعه فناوری	تفکیک تنظیم‌گری داده (EHDS) و الگوریتم (AI Act)، الزام به شفافیت، توضیح‌پذیری و نظارت انسانی
ایالات متحده آمریکا	Food and Drug Administration (FDA), Predetermined Change Control Plan (PCCP), Good Machine Learning Practice (GMLP)	FDA	تنظیم‌گری نهادی و مبتنی بر ریسک	ایمنی بیمار و عملکرد محصول	نظارت بر چرخه عمر کامل محصول، پایش پس از استقرار، امکان به‌روزرسانی کنترل‌شده مدل‌ها از طریق PCCP
کانادا	Personal Information Protection and Electronic Documents Act (PIPEDA), مقررات تجهیزات و نرم‌افزارهای پزشکی	Office of the Privacy Commissioner, Health Canada	ترکیب حفاظت از داده و تنظیم‌گری سلامت	حفظ حریم خصوصی و ارزیابی ایمنی نرم‌افزارهای پزشکی	تأکید بر مسئولیت‌پذیری، شفافیت، دسترسی افراد به داده‌های خود و ارزیابی نرم‌افزارهای پزشکی مبتنی بر هوش مصنوعی
استرالیا	privacy Act 1988, Australian Privacy Principles, استراتژی سلامت دیجیتال، اصول اخلاقی هوش مصنوعی	دولت استرالیا و نهادهای سلامت دیجیتال	ترکیبی از قانون، راهبرد ملی و اصول اخلاقی	حفاظت از داده، توسعه زیرساخت سلامت دیجیتال، استفاده مسئولانه از هوش مصنوعی	تأکید بر عدالت، شفافیت، مسئولیت‌پذیری، قابلیت تعامل داده‌ها و توسعه زیرساخت ملی
چین	Personal Information Protection Law (PIPL), برنامه Healthy China, سیاست‌های National Health Commission	National Health Commission	متمرکز و دولت‌محور	توسعه زیرساخت سلامت دیجیتال و کنترل داده‌ها	ترکیب سیاست‌گذاری متمرکز سلامت با قوانین حفاظت از داده، تأکید بر دیجیتالی‌سازی خدمات سلامت

جدول ۲. مقایسه وضعیت حکمرانی داده و هوش مصنوعی سلامت در کشورهای منتخب و ایران

کشور / منطقه	قانون حفاظت از داده	قانون اختصاصی هوش مصنوعی	فضای ملی اشتراک داده سلامت	تنظیم‌گر تخصصی AI سلامت	نظارت پس از استقرار
اتحادیه اروپا	GDPR	EU AI Act	EHDS	European AI Office	✓
ایالات متحده آمریکا	محدود و پراکنده	GMLP, PCCP	X	FDA	✓ بسیار قوی
کانادا	PIPEDA	X	محدود	نسبی	✓
استرالیا	Privacy Act	X (اصول اخلاقی)	در حال توسعه	محدود	نسبی
چین	PIPL	X	محدود	National Health Commission	محدود
ایران	در حال توسعه	X	X	X	X

خلاً موجود

پراکندگی داده‌ها

پراکنده بودن داده‌های سلامت و نبود استاندارد واحد جمع‌آوری داده

نبود سازوکار مشخص

نبود سازوکار مشخص برای استفاده پژوهشی از داده‌ها و نبود فرم تخصصی ارزیابی AI-health

ضعف در ممیزی

نبود نظام ثبت چرخه عمر مدل‌ها و ضعف در ممیزی و گزارش‌دهی

ابهام حقوقی

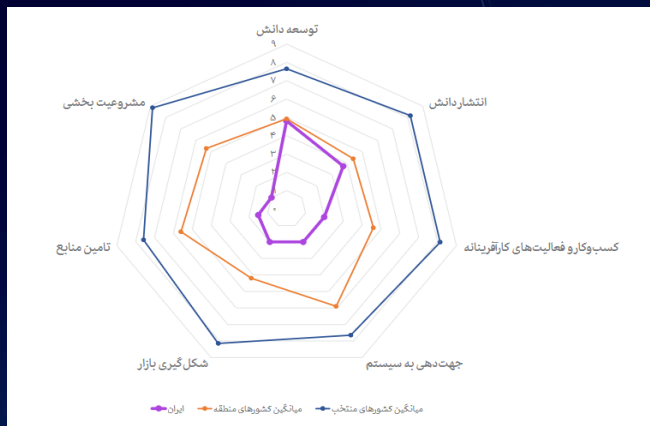
ابهام در مسئولیت حقوقی و مالکیت داده و مدل

⚠ نیاز به یک چارچوب اجرایی ملی وجود دارد.



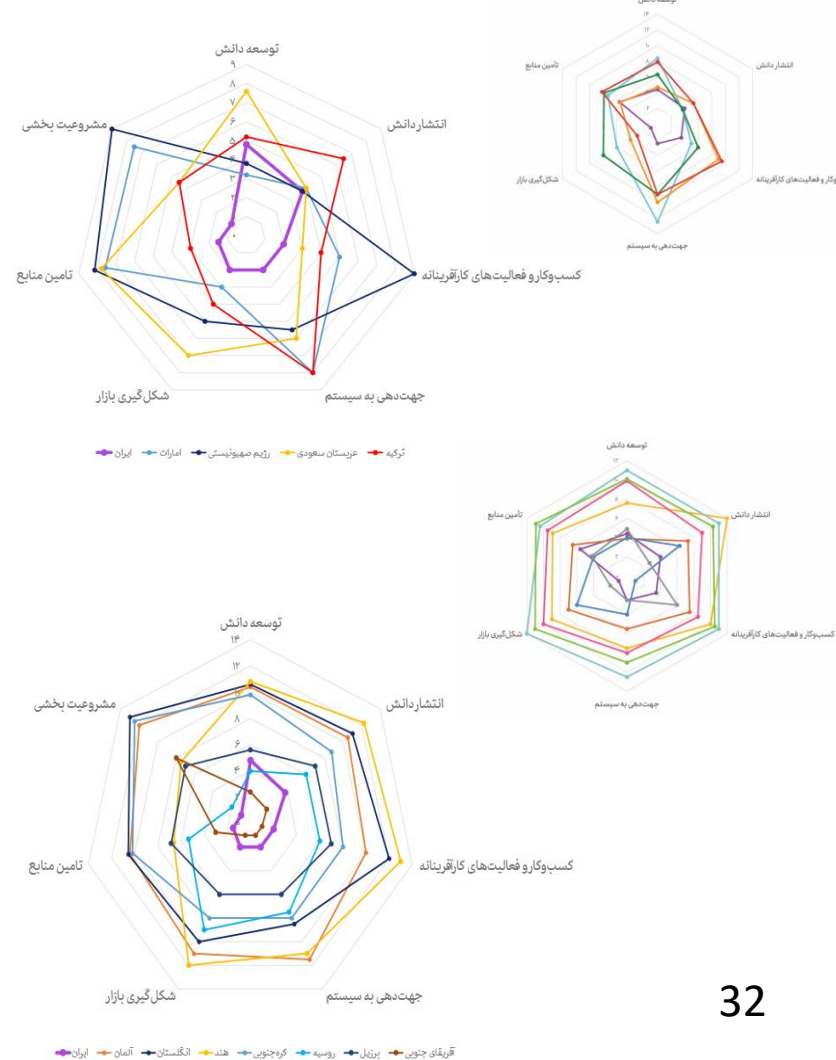
شاخص هوش مصنوعی ایران ۱۴۰۳

Iran AI Index 2024

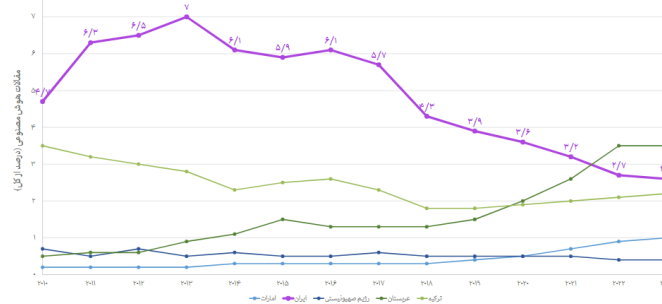


شاخص هوش مصنوعی ایران ۱۴۰۴

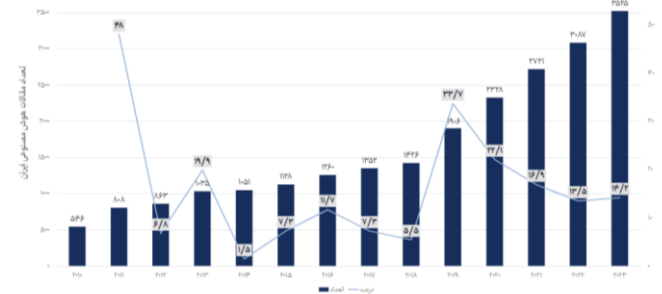
Iran AI Index 2025





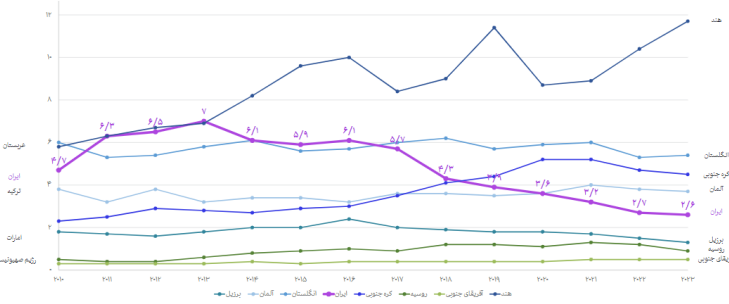
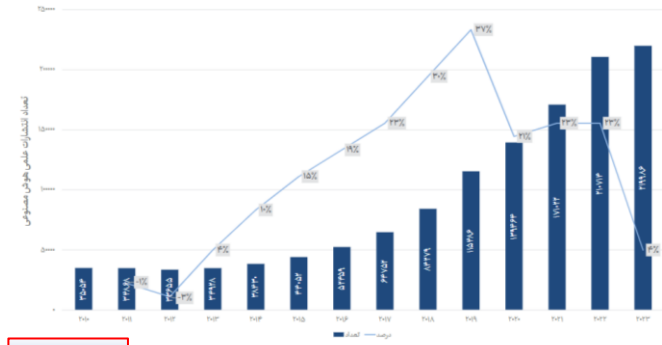


شکل ۷. مقالات هوش مصنوعی (درصد از کل) به تفکیک کشورهای منطقه، در بازه زمانی ۲۰۱۰ تا ۲۰۲۳.



Scopus

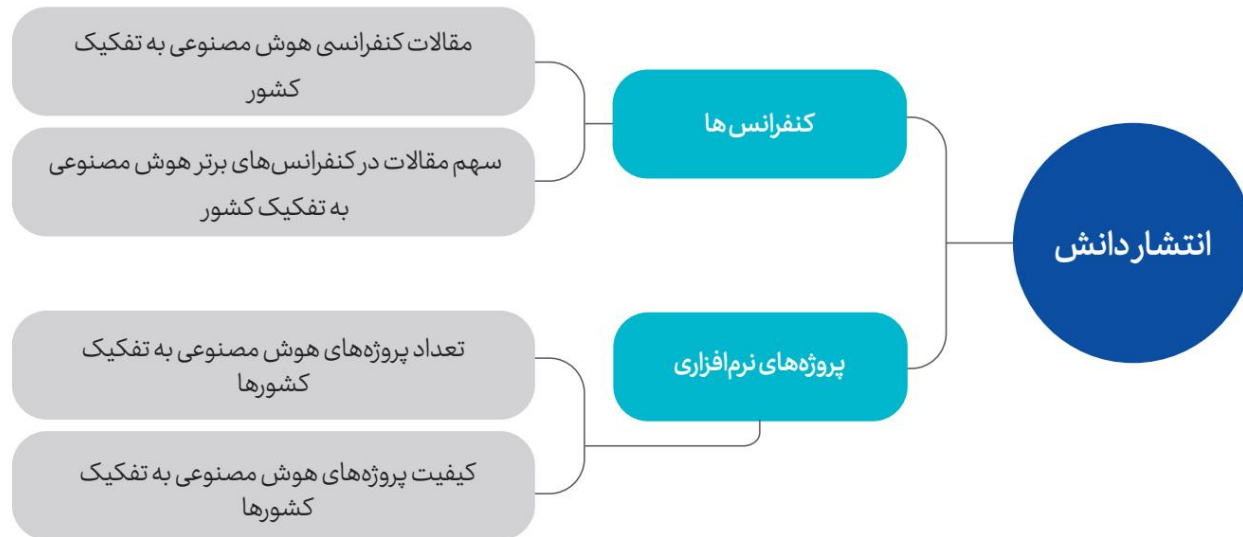
شکل ۴. تعداد و نرخ رشد انتشارات علمی هوش مصنوعی در جهان، در بازه زمانی ۲۰۱۰ تا ۲۰۲۳.

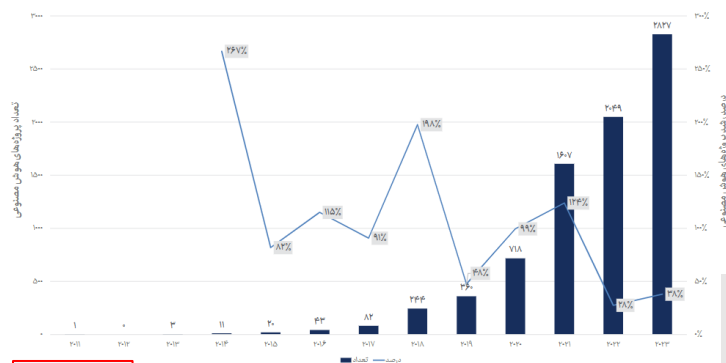


شکل ۸. مقالات هوش مصنوعی (درصد از کل) به تفکیک کشورهای منتخب، در بازه زمانی ۲۰۱۰ تا ۲۰۲۳.

از نظر علمی، ایران در سال‌های اخیر دچار افت جدی نسبت به کشورهای منطقه و جهان شده است. در سال ۲۰۲۳ ایران از نظر تعداد و کیفیت مقالات مرتبط با هوش مصنوعی در جایگاه دوم منطقه پس از عربستان سعودی قرار داشته است و با ادامه روند فعلی در ۲ الی ۳ سال آینده جایگاه دوم را به ترکیه واگذار خواهد کرد. همچنین تعداد مقالات کنفرانسی ایران در کنفرانس‌های برتر AI روندی نزولی داشته و هم اکنون در جایگاهی پس از رژیم صهیونیستی، عربستان سعودی، ترکیه و امارات قرار دارد که نشان از ارتباط اندک علمی ایران با جهان رو به پیشرفت هوش مصنوعی است.

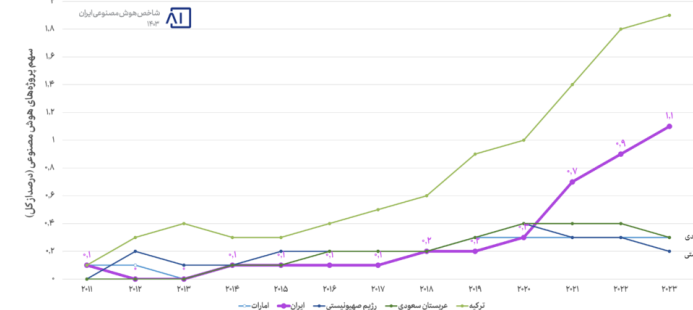
شکل ۸ نشان می‌دهد که ایران در سال ۲۰۱۳ در هوش مصنوعی، سهمی فراتر از کشورهای پیشرفته‌ای چون انگلستان، کره جنوبی و آلمان داشته است، اما سهم ایران در سال‌های اخیر به شدت در حال کاهش است. ادامه کاهش توان علمی ایران می‌تواند منجر به از دست رفتن این جایگاه نسبت به برزیل و روسیه در آینده کوتاه‌مدت شود. اگرچه در سال ۲۰۲۳ ایران حدود دو برابر برزیل و روسیه و تقریباً شش برابر آفریقای جنوبی تولید مقاله هوش مصنوعی داشته است، اما از آلمان و انگلستان عقب‌تر است. آلمان و انگلستان به ترتیب حدود ۱۴۰۰ و ۲۰۰۰ مقاله بیشتر از ایران در سال ۲۰۲۳ منتشر کرده‌اند. هند نیز با فاصله زیادی از ایران جلوتر است به طوری که در سال ۲۰۲۳، بیش از چهار برابر ایران مقاله منتشر کرده است.



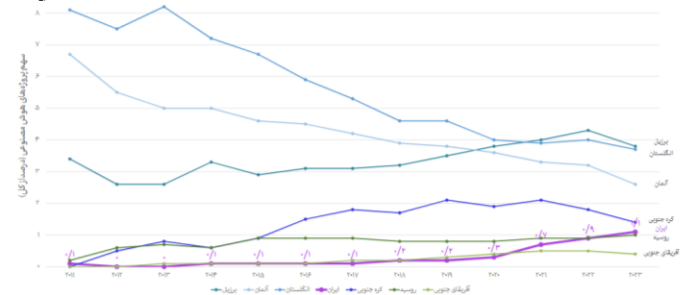


GitHub

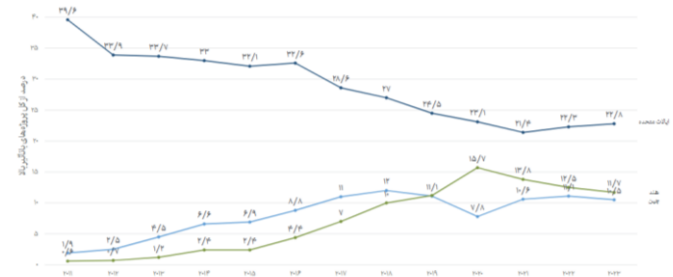
شکل ۴۳. تعداد و نرخ رشد پروژه‌های نرم‌افزاری هوش مصنوعی در ایران، در بازه زمانی ۱۴۰۱ تا ۱۳۹۳.



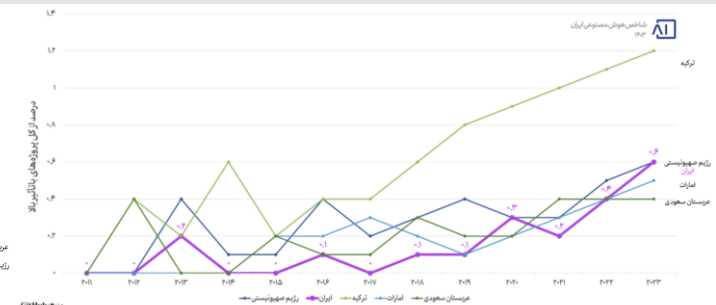
منبع: GitHub



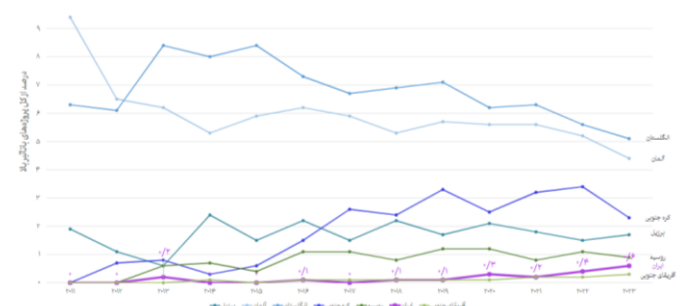
شکل ۴۵. سهم پروژه‌های نرم‌افزاری هوش مصنوعی کشورهای منتخب، در بازه زمانی ۱۴۰۱ تا ۱۳۹۳.



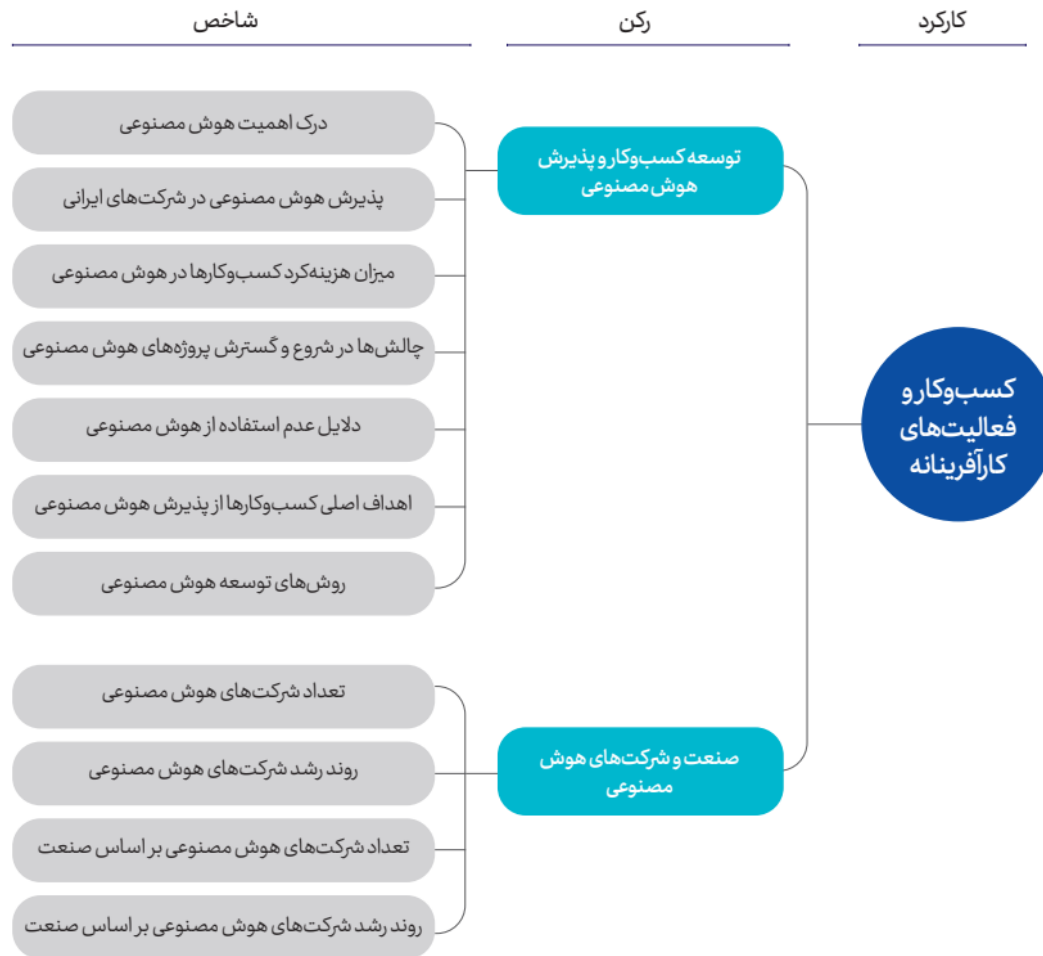
از نظر توسعه دهندگان هوش مصنوعی، ایران پس از ترکیه در جایگاه دوم منطقه قرار دارد. البته، روند مشارکت ایرانیان در توسعه پروژه‌های با کیفیت هوش مصنوعی در سال‌های اخیر صعودی بوده، روندی رو به رشد داشته است. این حوزه، یکی از مزایای کلیدی ایران در رقابت جهانی هوش مصنوعی است.



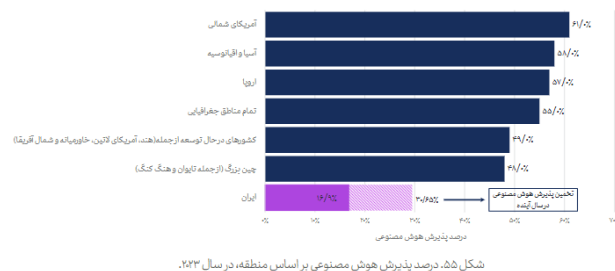
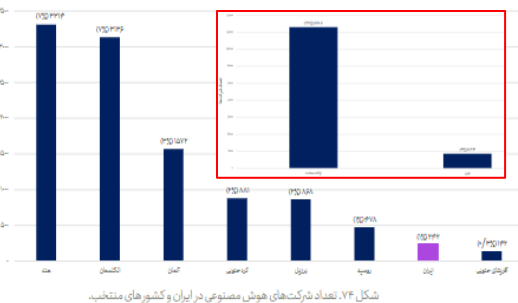
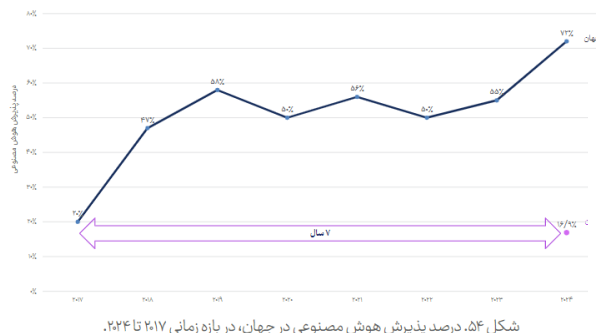
منبع: GitHub



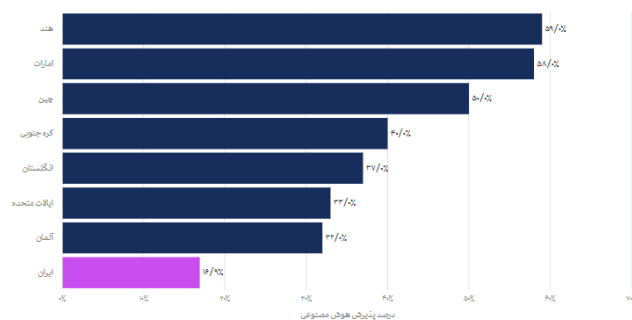
شکل ۴۸. سهم پروژه‌های نرم‌افزاری با تأثیر بالای هوش مصنوعی کشورهای منتخب، در بازه زمانی ۱۴۰۱ تا ۱۳۹۳.



برای مقایسه وضعیت ایران با جهان، باید بدین نکته توجه کرد که طیف وسیعی از تخمین‌ها در ادبیات مربوط به پذیرش هوش مصنوعی وجود دارد، اما به دلیل مشکلات در تعریف و اندازه‌گیری هوش مصنوعی، این تخمین‌ها متفاوت هستند و اغلب نمی‌توان به طور دقیق آن‌ها را با هم مقایسه کرد. با این وجود از گزارش مکنزی به منظور یک مقایسه کلی استفاده شده است. آخرین گزارش مکنزی^۹ در این خصوص نشان می‌دهد که در سال ۲۲،۲۲۴ درصد از سازمان‌های مورد بررسی، هوش مصنوعی را حداقل در یک واحد^{۱۰} پیاده‌سازی کرده‌اند. همان‌طور که در شکل ۵۴ قابل مشاهده است، نرخ پذیرش هوش مصنوعی در ایران به میزان قابل توجهی کمتر از سطح جهانی می‌باشد و با یک اختلاف حدود ۷ سال نسبت به جهان قرار دارد. از طرف دیگر بر اساس گزارش شاخص هوش مصنوعی دانشگاه استنفورد^{۱۱} که در شکل ۵۵ آمده است، میزان پذیرش هوش مصنوعی در ایران در سال ۱۴۳۳ از کشورهای در حال توسعه در سال ۲۰۲۳ نیز کمتر است. حتی اگر درصد کسانی که بیان کرده‌اند می‌خواهند در سال آینده از هوش مصنوعی استفاده کنند، نیز به درصد پذیرش امسال اضافه شود (قسمت هاشور خورده) باز هم در سطحی پایین‌تر از جهان قرار خواهد داشت.



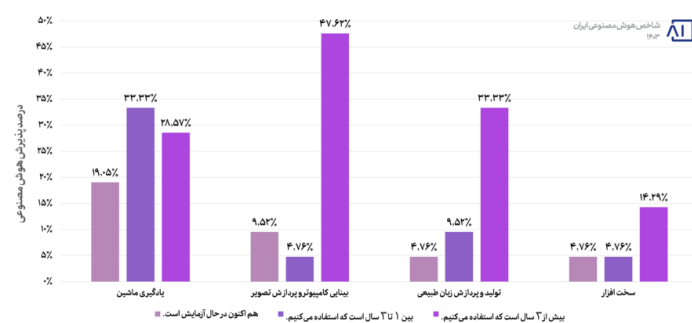
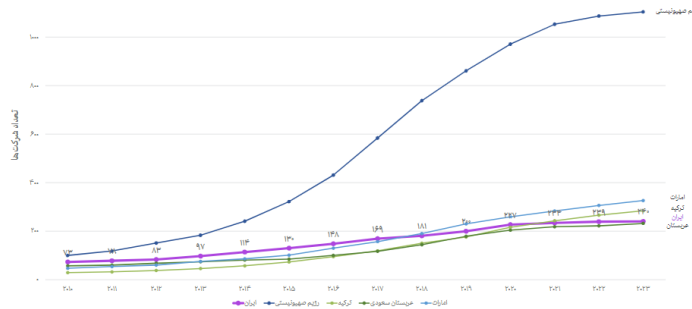
به منظور بررسی دقیق‌تر پذیرش هوش مصنوعی در سطح کشورها، می‌توان از گزارش آی‌بی‌ام^{۱۲} استفاده کرد. بر اساس شکل ۵۶، هند و امارات با ۵۹ و ۵۸ درصد بالاترین سطح پذیرش را در میان کشورهای مورد بررسی دارند. از سوی دیگر، سطح پذیرش هوش مصنوعی در ایران با ۱۶.۹ درصد، حدود نصف آلمان است، که خود به عنوان آخرین کشور، کمترین میزان پذیرش را در میان این کشورها دارد.



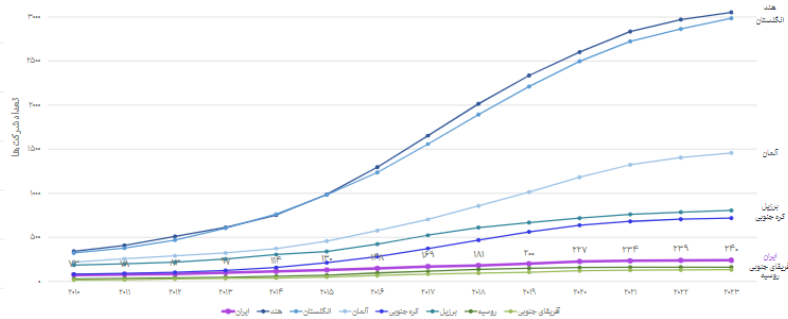
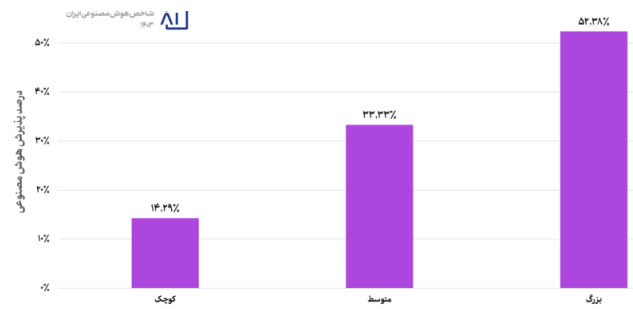
8. Core Business
9. The state of AI in early 2024: Gen AI adoption spikes and starts to generate value
10. Business Unit or Function
11. Artificial Intelligence Index Report 2024

از نظر کاربردی سازی هوش مصنوعی در کسب و کارهای ایرانی، در حال حاضر، ایران حداقل ۷ سال از میانگین جهانی فاصله دارد. بر این اساس، به کارگیری هوش مصنوعی در کسب و کارهای ایرانی در حال حدوداً ۱۷ درصد تخمین زده می‌شود که با میانگین جهانی فاصله زیادی دارد. این یک شاخص کلیدی برای کاربرد هوش مصنوعی در اقتصاد کشور است.

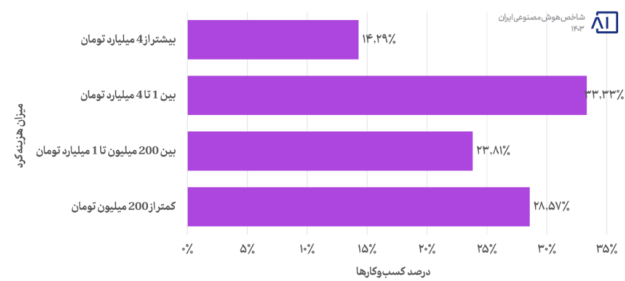
از نظر تعداد شرکت‌های هوش مصنوعی، ایران تا سال ۲۰۱۶ پس از رژیم صهیونیستی جایگاه دوم منطقه را داشته است، ولی در سال ۲۰۲۳ با سقوط ۲ پله‌ای در جایگاه چهارم منطقه قرار گرفت. همچنین روند رشد شرکت‌های ایرانی در مقایسه با سایر کشورهای منطقه و جهانی کاهشی بوده است.



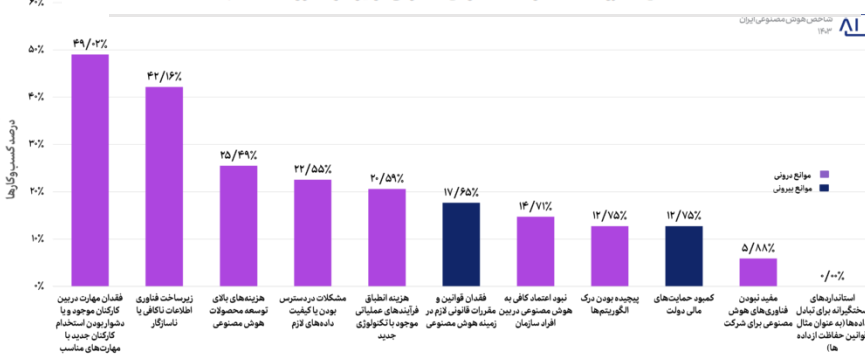
شکل ۷۶. روند تعداد شرکت‌های هوش مصنوعی در ایران و کشورهای منطقه.



منبع: پرسشنامه سنجش بازگویی هوش مصنوعی در شرکت‌های ایران



شکل ۷۷. روند تعداد شرکت‌های هوش مصنوعی در ایران و کشورهای منتخب.



منبع: پرسشنامه سنجش بازگویی هوش مصنوعی در شرکت‌های ایران

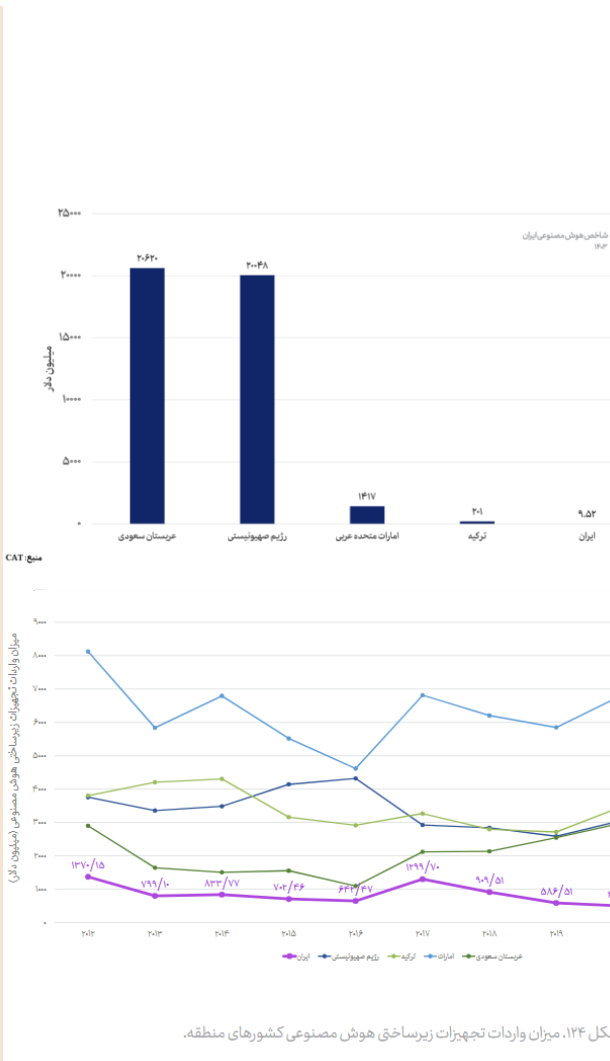


منبع: ATLAS & Crunchbase

از نظر حجم بازار، بازار هوش مصنوعی ایران در سال ۱۴۰۳ حدوداً ۱۶۰۰ میلیارد تومان تخمین زده می‌شود. براساس آمار مالی شرکت‌های هوش مصنوعی در ایران (مجموع زیر ساخت، داده و نرم افزار)، این بازار از سال ۱۳۹۶ در صورت تنزیل نرخ تورم سالانه با شیب اندکی کوچک شده که زنگ خطر جدی برای اکوسیستم هوش مصنوعی است و سهم شرکت‌های کوچک و متوسط از بازار، نسبت به شرکت‌های بزرگ به طور معناداری کاهش یافته است.

سرمایه‌گذاری ایران در تمامی سه عنصر کلیدی در هوش مصنوعی (داده، زیر ساخت، نیروی انسانی) با حدود ۱۰ میلیون دلار از تمامی کشورهای منطقه کمتر بوده است. سرمایه‌گذاری در هوش مصنوعی برای عربستان و رژیم صهیونیستی حدود ۲۰ میلیارد، امارات ۵/۱ میلیارد و ترکیه حدود ۲۰۰ میلیون دلار تخمین زده می‌شود. بدین ترتیب، فاصله معناداری میان سرمایه‌گذاری ایران در هوش مصنوعی با کشورهای منطقه وجود دارد.

حوزه	وضعیت در گزارش ۱۴۰۳	وضعیت در گزارش ۱۴۰۴
سرمایه‌گذاری	پذیرش هوش مصنوعی در کسب‌وکار	درصد ۲۷
فعال	اندازه بازار هوش مصنوعی	۰.۵ میلیارد تومان در ۱۴۰۳
توان پردازشی	تعداد شرکت‌های فعال	حدود ۸۲ میلیون دلار
معال (H100)	AI	شرکت ۳۷
توان پردازشی	جایگاه پنجم منطقه	H100 پردازنده ۱۷۲
میزان واردات تجهیزات زیرساختی هوش مصنوعی (میلیون دلار)	حدود ۱۰ میلیون دلار	۳۷ شرکت
میزان واردات تجهیزات زیرساختی هوش مصنوعی (میلیون دلار)	حدود ۱۰ میلیون دلار	۱۷۲ پردازنده H100
میزان واردات تجهیزات زیرساختی هوش مصنوعی (میلیون دلار)	حدود ۱۰ میلیون دلار	۳۷ شرکت
میزان واردات تجهیزات زیرساختی هوش مصنوعی (میلیون دلار)	حدود ۱۰ میلیون دلار	۱۷۲ پردازنده H100

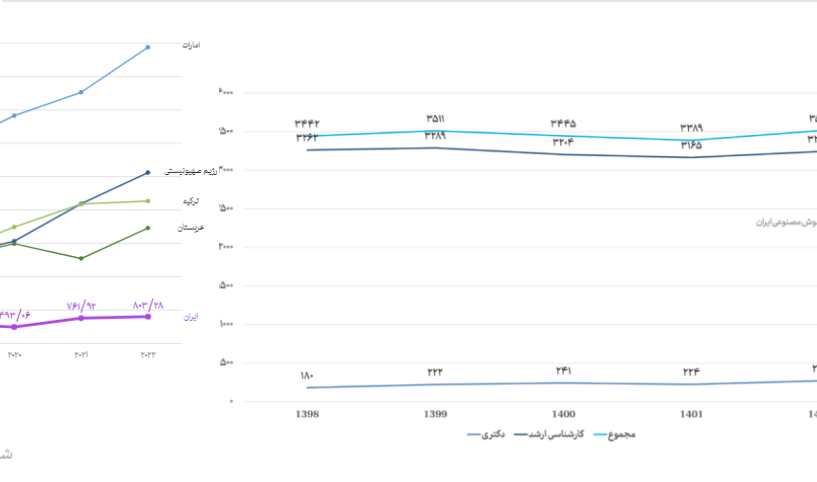


از نظر ظرفیت توان پردازشی داخلی مرتبط با هوش مصنوعی و دسترسی به پردازنده‌های پیشرفته و بروز، ایران در جایگاه پنجم منطقه قرار دارد. دسترسی به موقع و سریع ایران به توان پردازشی مرتبط با هوش مصنوعی در مقایسه با کشورهای منطقه بسیار پایین‌تر از میانگین جهانی است و خدمات ارائه شده در ایران برای پروژه‌های بزرگ مقیاس علمی و صنعتی در سطح جهانی، محدود است.

تعداد کل فارغ التحصیلان کارشناسی ارشد و دکتری مهندسی کامپیوتر مرتبط با هوش مصنوعی در سال حدود ۳۵۰۰ نفر است. این عدد در سال‌های اخیر علی‌رغم همه تحولات به وجود آمده در هوش مصنوعی و موج گسترده مهاجرت تحصیلی و کاری نخبگان از کشور همچنان بدون تغییر بوده است که به معنی کاهش نیروی کار با کیفیت در اکوسیستم هوش مصنوعی ایران است.

از نظر شاخص‌های مرتبط با اخلاق و قانون گذاری هوش مصنوعی، ایران در جایگاه پنجم منطقه قرار دارد که البته با کشورهای دیگر فاصله قابل توجهی ندارد. نداشتن قوانین و رویه‌های مرتبط با حفاظت از داده و چارچوب حکمرانی هوش مصنوعی دو چالش اصلی ایران در این حوزه است.

کمبود نیروی کار با مهارت کافی و نبود زیرساخت‌های مناسب دو عامل کلیدی در عدم توسعه هوش مصنوعی در بنگاه‌های ایرانی است. در میان عوامل کلیدی در عدم توسعه هوش مصنوعی در بنگاه‌های ایرانی، حمایت‌های مالی دولت در میان ده دلیل نخست قرار ندارد و کمبود منابع مالی و سرمایه گذاری در جایگاه سوم اهمیت قرار دارد.



Function

کارکرد

جهت‌دهی
به
سیستم

Pillar

رکن

آمادگی
هوش مصنوعی
دولت

اسناد و اقدامات
راهبردی

Index

شاخص

رتبه بخش دولت

اقدامات و سیاست‌های هوش مصنوعی

GARI
(Governance AI Readiness Index)

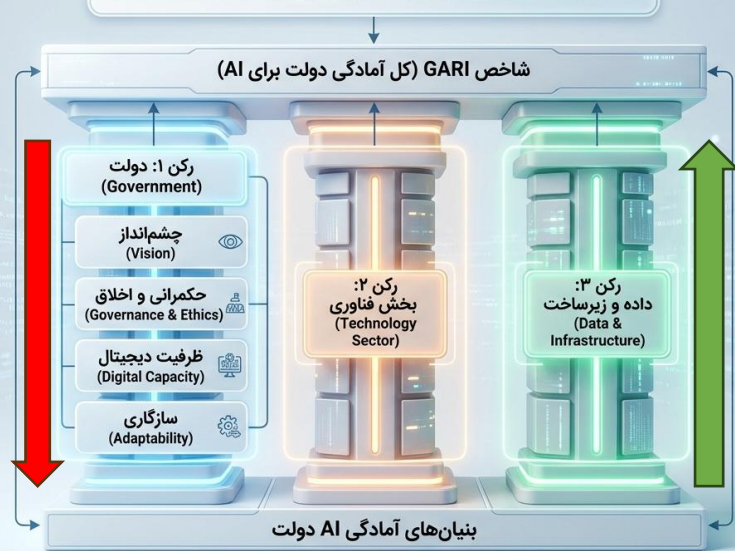
اقدامات و سیاست‌های هوش مصنوعی
به تفکیک حوزه عملکردی

1403 71

1404 91

اقدامات و سیاست‌های هوش مصنوعی
به تفکیک بخش‌ها

ساختار شاخص آمادگی دولت برای هوش مصنوعی (GARI)



GARI (Total)

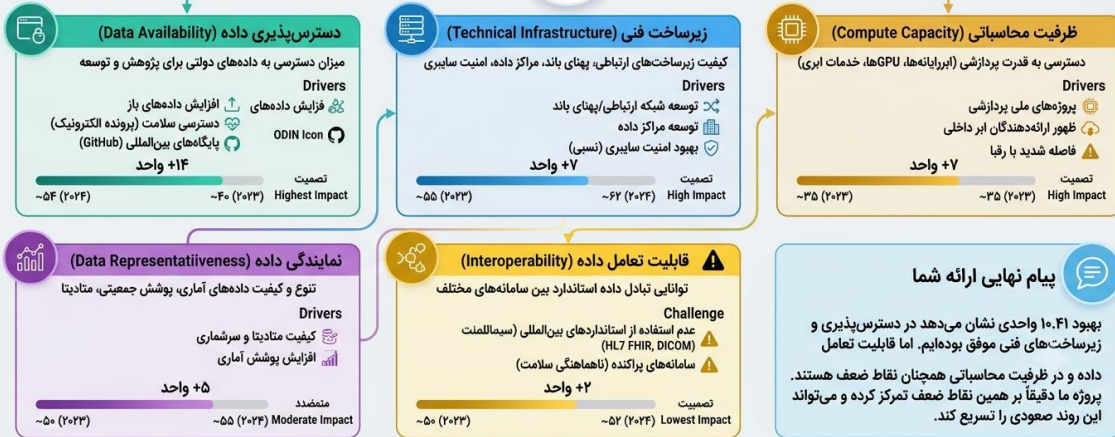
2023 94
2024 91

تحلیل دقیق رشد رکن داده و زیرساخت

۵۵.۸۸
(۲۰۲۳)

+۱۰.۴۱
رشد واحد

۶۶.۲۹
(۲۰۲۴)



پیام نهایی ارائه شما

بهبود ۱۰.۴۱ واحدی نشان می‌دهد در دسترسی‌پذیری و زیرساخت‌های فنی موفق بوده‌ایم. اما قابلیت تعامل داده و در ظرفیت محاسباتی همچنان نقاط ضعف هستند. پروژه ما دقیقاً بر همین نقاط ضعف تمرکز کرده و می‌تواند این روند صعودی را تسریع کند.

رکن فناوری

مهاجرت استعدادها	AI تعداد شرکت‌های	R&D (% GDP) سرمایه‌گذاری	کشور
مثبت	۱,۲۹۹	~۵.۴٪ (بالاترین جهان)	ریژیم صهیونیستی
مثبت (۴.۱)	۵۷۴	~۱.۵٪	امارات متحده عربی
مثبت (۱.۶)	۳۱۸	~۰.۸٪	عربستان سعودی
منفی (-۰.۵)	۵۹۸	~۱.۱٪	ترکیه
منفی (-۸.۹)	۳۸۷	~۰.۸٪	ایران

43

زیرشاخص رکن دولت

تغییر	امتیاز 2024	امتیاز 2023	زیرشاخص رکن دولت
-4.54	38.89	43.43	چشم‌انداز (Vision)
-2.03	18.91	20.94	حکمرانی و اخلاق (Governance & Ethics)
+3.27	24.42	21.15	ظرفیت دیجیتال (Digital Capacity)
-3.09	17.56	20.65	سازگاری (Adaptability)

حوزه	وضعیت		
Publications	خوب	Commercialization	ضعیف
Human Resources	خوب	Governance	ضعیف
Infrastructure	متوسط	Health Data	پراکنده

WHAT MAKES A GOOD MEDICAL AI?

چه چیزی یک هوش مصنوعی پزشکی خوب را می‌سازد؟

GOOD AI $\neq$ GOOD ALGORITHM

GOOD AI =



**SUCCESSFUL
MEDICAL AI**

یک پیشنهاد

طراحی یک چارچوب عملیاتی برای:



مدیریت داده‌های سلامت



مدیریت ریسک و سوگیری الگوریتمی



مستندسازی و ممیزی پروژه‌ها



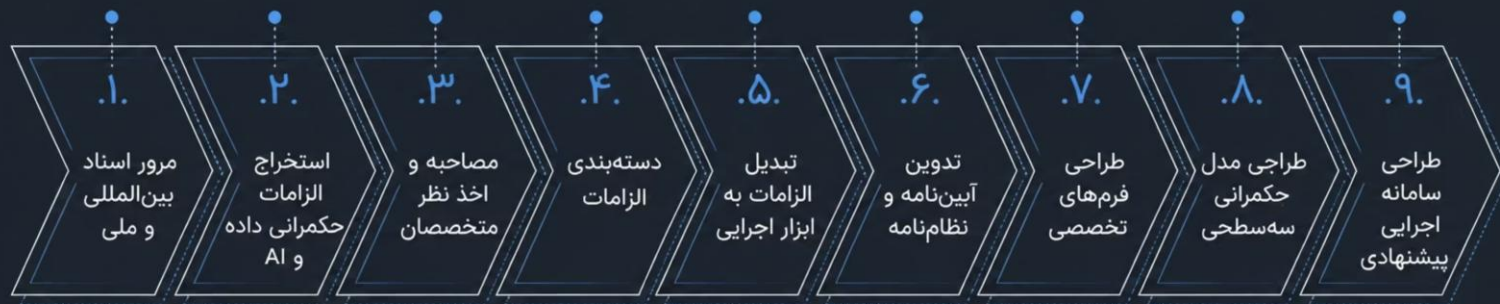
ارزیابی پروژه‌های **AI-health**



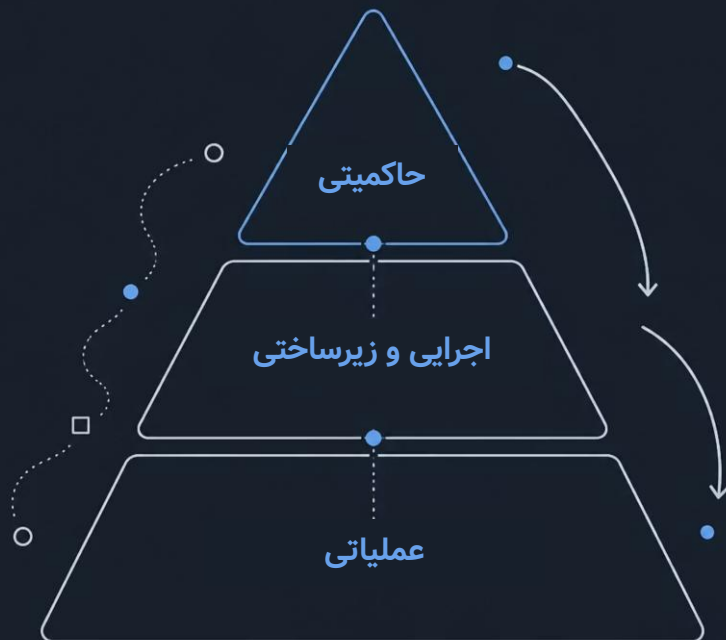
نظارت انسانی و توضیح‌پذیری



ایجاد زیرساخت حکمرانی داده و **AI** سلامت



مدل حکمرانی سه سطحی پیشنهادی، تمام سطوح تصمیم‌گیری، اجرا و بهره‌برداری را پوشش می‌دهد و هماهنگی کامل بین سیاست‌گذاران، مجریان و کاربران نهایی را تضمین می‌کند.



اساسنامه جامع ملی استفاده از داده‌های سلامت در سامانه‌های هوش مصنوعی

چارچوب تنظیم، نظارت، اخلاق و حاکمیت داده سلامت

آیین‌نامه اجرایی طرح‌های پژوهشی هوش مصنوعی در
داده‌های سلامت

فصل اول: هدف و دامنه

ماده ۱ - هدف:

تعیین مراحل قانونی، اخلاقی و اجرایی پروژه‌های پژوهشی مبتنی بر داده‌های سلامت در سامانه‌های هوش مصنوعی شامل: دسترسی، پردازش، امنیت داده‌ها، شفافیت الگوریتم و گزارش‌دهی.

ماده ۲ - حاکمیت داده سلامت

مطابق سیاست کلان ۲ و هدف ۷ از سند ملی هوش مصنوعی در نظام سلامت جمهوری اسلامی ایران کلیه فرآیندهای جمع‌آوری، پردازش، ذخیره، اشتراک‌گذاری و بهره‌برداری از داده‌های سلامت در چارچوب سیاست‌های ملی حاکمیت داده سلامت

مراکز ملی داده سلامت قابل اعتماد (THDC)

مراکز ملی خوانش داده سلامت به عنوان ستون فقرات زیرساخت هوش مصنوعی در حوزه بهداشت و درمان، وظایف کلیدی زیر را بر عهده دارند:

برچسب گذاری داده های پزشکی

کنترل کیفیت داده

تجمیع و استانداردسازی داده ها

نگهداری بانک داده و مدل های AI

ارائه زیرساخت پردازشی

مدیریت دسترسی پژوهشگران

پشتیبانی از پروژه های ملی AI سلامت

پیشنهاد اولیه برای استقرار بصورت منطقه ای 

چارچوب‌های زیربنایی پشتیبان THARF



هدف

چارچوب THARF میزان آمادگی سازمان‌ها را در هشت بُعد کلیدی ارزیابی کرده و مسیر بلوغ آن‌ها را از وضعیت آغازین تا استقرار یک اکوسیستم هوش مصنوعی قابل اعتماد در نظام سلامت هدایت می‌کند.

اصول طراحی چارچوب مرجع



مسیر گام‌به‌گام از طراحی و استقرار تا بلوغ و ارتقای مستمر مرکز داده سلامت قابل اعتماد



توانمندسازهای کلیدی (در همه مراحل)



مرکز داده سلامت قابل اعتماد بالغ (Mature THDC)



این اصول به عنوان راهنمای حاکم بر تمامی مدل‌ها، فرایندها، ساختارها و خدمات مراکز داده سلامت عمل می‌کنند و تضمین می‌کنند که تمامی تصمیم‌گیری‌ها در راستای ایجاد ارزش پایدار برای نظام سلامت انجام شود.

برنامه اشتراک‌گذاری داده‌ها
نوع داده‌ها: بالینی، تصویر، ژنتیکی، داده‌های AI و ترکیبی، فرمت و حجم تقریبی داده‌ها، محدوده اشتراک‌گذاری: کامل یا محدود (با ذکر دلیل محدودیت در صورت وجود)، دامنه استفاده و بازه زمانی استفاده، معیارهای پژوهشگران واجد شرایط برای استفاده مجدد، حفاظت از حریم خصوصی و محرمانگی، ناشناس‌سازی (de-identification)، رضایت آگاهانه، مستندسازی داده‌ها، شرح متغیرها، روش کدگذاری، استانداردهای پیش‌پژوه‌ای، فرمت قابل اشتراک‌گذاری، نحوه دسترسی به داده‌های تولیدشده در پایان پروژه (در صورت انتخاب دسترسی محدود، پژوهشگر باید مشخص کند، چه گروه‌هایی مجاز به استفاده از داده‌ها هستند و محدودیت‌ها شامل چه مواردی می‌شود؛ مثلاً فقط تیم پژوهشی، فقط برای تحقیقات مشابه، محدودیت زمانی و غیره) ☐ منابع باز (Public) ☐ دسترسی محدود – (Restricted) نوع و دامنه محدودیت: ☐ عدم دسترسی (No access) آیا آزمون ریسک بازشناسایی (Re-identification Risk Test) انجام شده است؟ آیا ارزیابی سازگاری استفاده ثانویه (Compatibility Assessment) انجام شده است؟ آیا سطح دسترسی مبتنی بر نقش (Role-based Access) تعریف شده است؟
جمع‌آوری و پردازش داده‌ها
توسیع روش‌های جمع‌آوری و ثبت داده‌ها، روش‌های کنترل کیفیت داده‌ها، روش‌های پردازش و آموزش مدل، نظارت انسانی، مدیریت ریسک و روش‌های مستندسازی، تفکیک داده‌های خام، پیش‌پردازش‌شده و مشتق‌شده، ثبت داده‌ها و استفاده‌های داده، ثبت کامل، جزئیه عمر مدل، ثبت تغییرات پارامتر و استفاده، تعریف شاخص‌های عدالت الگوریتمی و (تاریخ آن، مدت زمان‌های امنیت سایبری احتمالی)
مشخصات فنی سامانه هوش مصنوعی
نقش دقیق هوش مصنوعی در این طرح (به عنوان مثال جهت Training, prediction, decision support و...) ذکر شود. سطح ریسک سامانه هوش مصنوعی (Low / High / Unacceptable). در صورت High Risk بودن، موارد پایش پس از استقرار، ثبت خطا در سامانه ملی، توضیح پذیری (explainability) تقویت شده و نظارت انسانی چند لایه توضیح داده شود. در صورت وجود unacceptable risk، اجرای طرح غیرمجاز است. سازوکار Human-in-the-Loop سازوکار توضیح پذیری Explainability (در صورت الزام‌گذاری بر تصمیم بالینی) شاخص‌های سنجش سوگیری و عدالت الگوریتمی نحوه ثبت لاگ‌ها و قابلیت ممیزی مدل برنامه پایش پس از استقرار (در صورت High Risk بودن) نوع الگوریتم، داده‌های تست و آموزش، روش ارزیابی ریسک، روش مستندسازی فرآیندها، تعهدات حقوقی، نظارت انسانی

بیان مساله و ضرورت اجرای مطالعه
با در نظر گرفتن مطالعات قبلی، ضرورت اجرای مطالعه را تبیین کنید. (حداکثر ۱۵۰۰ کلمه) ۴ برای طرح‌های هوش مصنوعی ضرورت استفاده از هوش مصنوعی در این مطالعه باید به صورت تحلیلی تبیین گردد (شامل بازآزمایی روش‌های مرسوم، ایجادکننده‌ها، معیارهای ارزیابی با نیاز به تصمیم‌گیری پس از بررسی)
روش اجرا
نوع و جهت مطالعه، مراحل اجرای مطالعه، معیارهای ورود و خروج، روش نمونه‌گیری، محاسبه حجم نمونه، روش‌های جمع‌آوری اطلاعات، روایی و پایایی ابزار گردآوری داده‌ها، نحوه آموزش، شرح مداخله یا تجویز دارو، روش‌ها و ابزارهای تجزیه و تحلیل اطلاعات و ... نقش دقیق هوش مصنوعی در این طرح (به عنوان مثال جهت Training, prediction, decision support و...) ذکر شود. پژوهشگر موظف است موارد زیر را به‌طور دقیق در پیوست ۱ و ۲ تبیین نماید: سطح ریسک سامانه هوش مصنوعی (Low / High / Unacceptable) سازوکار Human-in-the-Loop سازوکار توضیح پذیری Explainability (در صورت اثرگذاری بر تصمیم بالینی) شاخص‌های سنجش سوگیری و عدالت الگوریتمی نحوه ثبت لاگ‌ها و قابلیت ممیزی مدل برنامه پایش پس از استقرار (در صورت High Risk بودن)

کد طرح:	
---------	--

داور گرامی:

از اینکه داوری این طرح را پذیرفتید سپاسگزاریم. لطفاً به نکات زیر توجه فرمایید:

دیدگاه‌های شما نقش موثری در تصمیم‌گیری درباره این درخواست دارد.

لطفاً نظرات خود را در این فرم وارد کرده و سپس در سامانه پژوهشیار بارگذاری نمایید. فرم را از راه دیگری ارسال نفرمایید.

هویت داوران برای مجری مشخص نمی‌شود و اطلاعات طرح محرمانه است. پس از ارسال فرم، لطفاً کلیه مستندات را حذف نمایید.

در صورت وجود تضاد منافع یا مجری یا ارتباط شخصی با طرح، از داوری انصراف دهید.

۱. ارزیابی اطلاعات کلی طرح

هدف طرح واضح و شفاف است؟ ☐ ایلی ☐ خیر ☐ تا حدی

نوع داده‌ها با اهداف طرح تناسب دارد؟ ☐ ایلی ☐ خیر ☐ تا حدی

۲. ضرورت اجرا و نقش هوش مصنوعی

نقش هوش مصنوعی (training / prediction / decision support) به‌طور شفاف و به درستی تعریف شده است؟

☐ ایلی ☐ خیر ☐ تا حدی

مسأله سلامت برای کاربرد هوش مصنوعی مناسب است؟ ☐ ایلی ☐ خیر ☐ تا حدی

۳. طبقه‌بندی سطح ریسک سامانه

به نظر شما سطح ریسک سامانه چیست؟

در صورت High Risk بودن، آیا الزامات تکمیلی رعایت شده‌اند؟ ☐ ایلی ☐ خیر ☐ تا حدی

۴. نوآوری طرح

به نظر شما نوآوری طرح در کدام سطح است؟

☐ نوآوری در نوع داده سلامت (مثلاً داده داخلی، داده ترکیبی، داده کپدستری)

☐ نوآوری در الگوریتم یا معماری مدل

☐ نوآوری در کاربرد بالینی/تصمیم‌گیری

☐ نوآوری در Human-in-the-Loop یا مدیریت ریسکتوضیح داور.....:

۵. ارزیابی روش اجرای مطالعه

۴-الف. داده‌های سلامت

نوع داده‌ها دقیق و کافی تعریف شده؟ ☐ ایلی ☐ خیر ☐ تا حدی

دامنه استفاده از داده‌ها محدود و مشخص است؟ ☐ ایلی ☐ خیر ☐ تا حدی

سطح و مدت دسترسی منطقی و متناسب است؟ ☐ ایلی ☐ خیر ☐ تا حدی

روش ناشناس‌سازی و امنیت داده‌ها قابل قبول است؟ ☐ ایلی ☐ خیر ☐ تا حدی

سیاست دسترسی به داده‌ها در پایان طرح شفاف و مناسب است؟ ☐ ایلی ☐ خیر ☐ تا حدی

آزمون ریسک بازنشاسایی انجام شده؟ ☐ ایلی ☐ خیر ☐ تا حدی

ارزیابی سازگاری استفاده ثانویه انجام شده؟ ☐ ایلی ☐ خیر ☐ تا حدی

ثبت متادیتا و نسخه‌بندی داده‌ها پیش‌بینی شده؟ ☐ ایلی ☐ خیر ☐ تا حدی

تثکبک داده خام / مشتق‌شده رعایت شده؟ ☐ ایلی ☐ خیر ☐ تا حدی

۴-ب. سامانه‌های هوش مصنوعی

نوع الگوریتم با هدف پژوهش متناسب است؟ ☐ ایلی ☐ خیر ☐ تا حدی

داده‌های train / test / validation شفاف‌اند؟ ☐ ایلی ☐ خیر ☐ تا حدی

نحوه تصمیم‌گیری مدل و تبدیل خروجی به تصمیم عملیاتی توضیح داده شده؟ ☐ ایلی ☐ خیر ☐ تا حدی

نظرات انسانی (Human-in-the-Loop) مناسب است؟ ☐ ایلی ☐ خیر ☐ تا حدی

روش‌های مدیریت ریسک، سوگیری و خطا مناسب هستند؟ ☐ ایلی ☐ خیر ☐ تا حدی

روش‌های مستندسازی، لاگ‌گیری و گزارش‌دهی درست پیش‌بینی شده؟ ☐ ایلی ☐ خیر ☐ تا حدی

سازوکار Explainability متناسب با سطح ریسک تعریف شده؟ ☐ ایلی ☐ خیر ☐ تا حدی

شاخص‌های عدالت الگوریتمی کمی و قابل سنجش هستند؟ ☐ ایلی ☐ خیر ☐ تا حدی

پایش پس از استقرار (در صورت High Risk بودن) پیش‌بینی شده؟ ☐ ایلی ☐ خیر ☐ تا حدی

سازوکار ثبت خطا و گزارش به سامانه ملی مشخص است؟ ☐ ایلی ☐ خیر ☐ تا حدی

تقسیم مسئولیت حقوقی بین توسعه‌دهنده، مجری و کاربر مشخص است؟ ☐ ایلی ☐ خیر ☐ تا حدی

۶. پیوستگی با پژوهش‌های پیشین

تیم پژوهشی تجربه قبلی در کار با داده سلامت و AI دارد؟ ☐ ایلی ☐ خیر ☐ تا حدی

این طرح بخشی از مسیر تکاملی مدل داده است یا پروژه منفرد؟ ☐ مسیر تکاملی ☐ پروژه منفرد

۷. نگارش و شفافیت پروپوزال

اصطلاحات مرتبط با AI دقیق و غیرمبهم استفاده شده‌اند؟ ☐ ایلی ☐ خیر ☐ تا حدی

تصمیم‌گیری هوش مصنوعی به زبان قابل فهم توضیح داده شده؟ ☐ ایلی ☐ خیر ☐ تا حدی

۸. ملاحظات اخلاقی و دسترسی به داده‌ها

☐ تعهد به عدم استفاده خارج از دامنه

☐ تعهد به عدم استفاده تجاری

☐ شفافیت در گزارش‌دهی و ممیزی

☐ نظرات انسانی بر خروجی‌های تصمیم‌گیر

☐ احقاق اعتراض و بازبینی انسانی برای بیمار لحاظ شده

☐ امکان توقف فوری مدل در صورت بروز خطای بحرانی وجود دارد

☐ توضیحات پالاقوه دارایی فکری (مدل الگوریتم) شفاف بوده و سازوکار اعلام و مدیریت مالکیت آن مطابق آیین‌نامه مشخص شده است.

توضیح داور.....:

۹. بودجه و منابع

هزینه‌های پردازش داده، زیرساخت امن، ذخیره‌سازی، ممیزی و گزارش‌دهی منطقی است؟ ☐ ایلی ☐ خیر ☐ تا حدی

۱۰. جمع‌بندی نهایی

طرح از نظر:

☐ طبایق با آیین‌نامه AI سلامت و اسناد بالادستی

☐ ثقات حکمرانی داده

☐ ثقات مدیریت ریسک

آماده اجرا است؟ ☐ ایلی ☐ خیر ☐ مشروط به اصلاحات AI محور

نیاز به محدودسازی دامنه داده یا مدل وجود دارد؟ ☐ ایلی ☐ خیر

فرم داوری

داوری تخصصی AI-health فراتر از ارزیابی علمی صرف، ابعاد حکمرانی و ایمنی هوش مصنوعی را نیز پوشش می‌دهد:



Fairness و Bias Assessment



سطح ریسک



حاکمیت داده



امنیت سایبری



Human-in-the-Loop



Explainability



پایش پس از استقرار



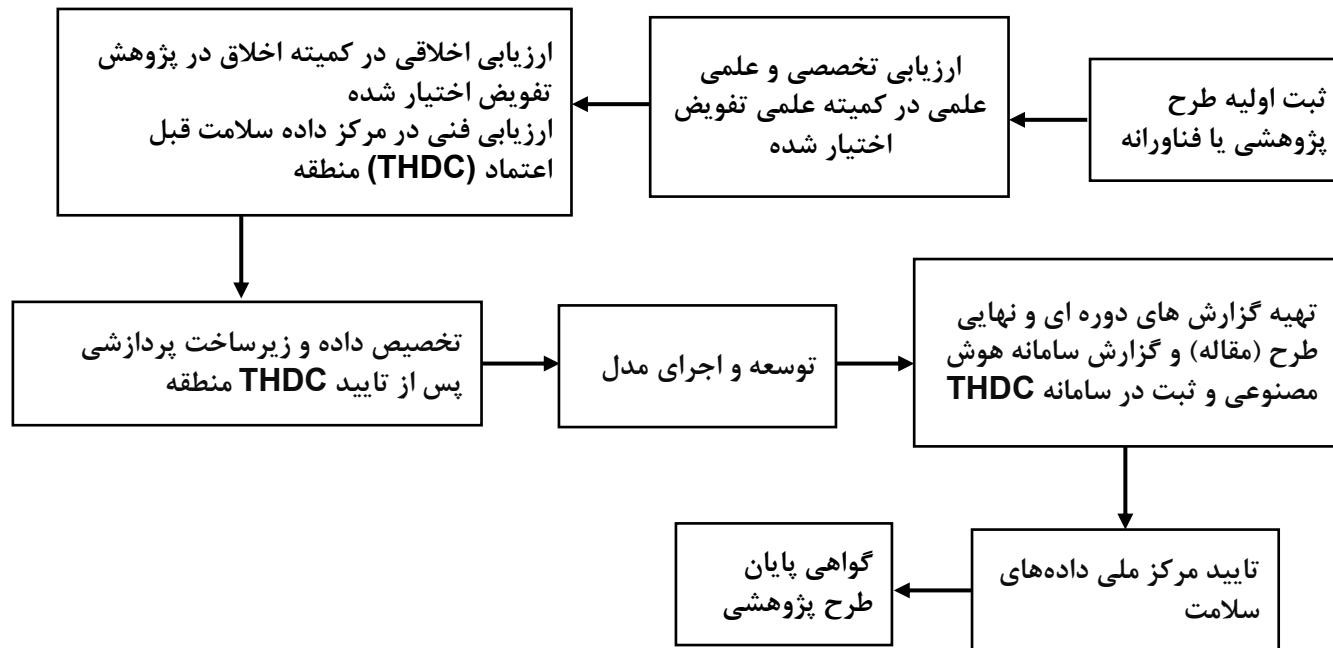
مسئولیت حقوقی

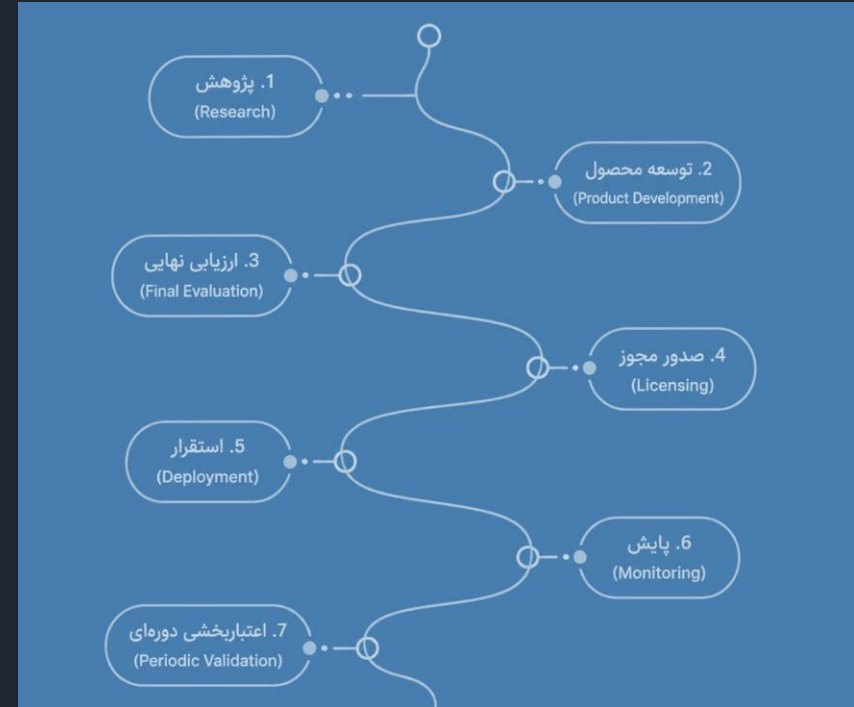


ممیزی و لاگ‌گیری

پیام کلیدی: داوری صرفاً علمی نیست؛ حکمرانی و ایمنی AI نیز ارزیابی می‌شود.

مسیر رسمی طرح های پژوهشی هوش مصنوعی سلامت





طرح‌های محصول محور AI سلامت مسیر تنظیم‌گری مشخصی را طی می‌کنند که تحت نظارت مستقیم سازمان غذا و دارو اجرا می‌شود. این مسیر از مرحله پژوهش آغاز شده و پس از ارزیابی‌های متعدد به صدور مجوز و استقرار منتج می‌گردد.

دستاوردهای مورد انتظار



افزایش شفافیت و قابلیت ممیزی
ردیابی کامل فرآیندها و تصمیمات



افزایش کیفیت داده‌های پژوهشی
ارتقای استانداردهای داده‌گیری و کنترل کیفیت



کاهش پراکندگی داده‌ها
تمرکز داده‌های سلامت در مراکز ملی استاندارد



بهبود ایمنی و اعتمادپذیری سامانه‌ها
ارزیابی جامع ریسک و حکمرانی



تسهیل توسعه AI سلامت
زیرساخت یکپارچه برای پژوهشگران



کاهش تکرار پروژه‌های مشابه
استفاده از بانک دانش ملی



ایجاد بانک ملی داده و دانش AI سلامت
ذخیره‌سازی پایدار مدل‌ها و نتایج

جمع بندی

- AI نویسنده مقاله نیست.
- مسئولیت علمی با پژوهشگر است.
- داده محرمانه را در AI عمومی بارگذاری نکنید.
- استفاده از AI را در صورت نیاز افشا کنید.
- خروجی AI را راستی آزمایی کنید.
- رفرنس‌ها را کنترل کنید.
- تصاویر علمی را بدون بررسی منتشر نکنید.
- مدل AI بدون اعتبارسنجی بالینی قابل استفاده نیست.
- کیفیت داده مهم‌تر از پیچیدگی مدل است.
- آینده AI سلامت بر پایه حکمرانی داده و اعتماد ساخته می‌شود.

آینده هوش مصنوعی سلامت از مدل‌های بزرگ‌تر آغاز نمی‌شود؛ از داده‌های قابل اعتماد آغاز می‌شود.

با تشکر از توجه شما

داوری مقاله در نشریات علمی: مسئله اصلی محرمانگی است، نه دزدیده شدن ایده توسط AI

مقاله در حال داوری، یک سند محرمانه است و نباید بدون مجوز در اختیار هیچ شخص یا سرویس ثالثی قرار گیرد.

1. فرآیند معمول داوری مقاله



تعهد محرمانگی داور:

داور متعهد می شود که مقاله را فقط برای داوری استفاده کند و محتوای آن را به هیچ شخص یا سازمان ثالثی به اشتراک نگذارد.

2. ریسک های استفاده از سرویس های عمومی AI



- × خروج مقاله از محیط محرمانه
- × به اشتراک گذاری با یک سرویس ثالث ناشناخته
- × نقض تعهد محرمانگی داور
- × عدم انطباق با سیاست های ناشر و مقررات نهادهای
- × عدم کنترل بر نگهداری، پردازش و استفاده از داده
- × پیامدهای حقوقی و اخلاقی احتمالی

مشکل این نیست که AI ایده را می دزدد،

مشکل، نقض، اصل محرمانگی است.

3. راه حل پیشنهادی: استفاده از محیط های مورد اعتماد



- ✓ مقاله هرگز محیط مورد اعتماد را ترک نمی کند
- ✓ داده تحت کنترل کامل نهاد/ناشر باقی می ماند
- ✓ مطابق با سیاست های محرمانگی و حاکمیت داده
- ✓ AI به داور کمک می کند، تصمیم نهایی انسانی است
- ✓ استفاده ایمن، قابل حسابرسی و مسئولانه از AI

AI می تواند در داوری کمک کند، اگر محرمانگی، پاسخگویی و سیاست ها حفظ شوند.

اهمیت دلایل در استفاده از AI عمومی (از چهار ستاره = بسیار مهم تا یک ستاره = کم اهمیت)

اهمیت	دلیل
★★★★	نقض تعهد محرمانگی (مهم ترین دلیل)
★★★★	خروج سند از کنترل ناشر (حاکمیت داده)
★★	حفاظت از مالکیت فکری و اطلاعات منتشر نشده
★	نگرانی از استفاده مدل برای آموزش
★	نگرانی از هک یا ناامنی اینترنت

جمع بندی

مسئله اصلی در استفاده از سرویس های عمومی AI در داوری مقاله، حفظ محرمانگی، حاکمیت داده و رعایت سیاست های ناشر است. نه سرقت ایده توسط AI.

چرا محرمانگی در داوری مقاله حیاتی است؟

- 1. تعهد محرمانگی**
داور فقط برای ارزیابی از مقاله استفاده می کند و نباید آن را به هیچ شخص یا سازمانی به اشتراک بگذارد.
- 2. حفاظت از کار منتشر نشده**
مقاله ممکن است شامل ایده ها، روش ها، داده ها و نتایج جدیدی باشد که هنوز عمومی نشده اند.
- 3. الزامات حقوقی و اخلاقی**
نقض محرمانگی می تواند سیاست های ناشر، قوانین ملی/بین المللی و حتی مقررات مالکیت فکری را نقض کند.
- 4. ریسک های استفاده از AI عمومی**
 - از دست رفتن کنترل بر داده
 - نگهداری نامشخص داده
 - نقض سیاست ها و قراردادهای
 - آسیب به اعتماد و اعتبار علمی
- 5. مسئولیت داور**
داور مسئول حفاظت از مقاله و کل فرآیند علمی است، نه فقط از نویسنده.

اصل کلیدی:

مسئولیت داور، حفاظت از مقاله و کل فرآیند علمی است، نه فقط از نویسنده.

THDC چگونه کمک می کند؟

- محیط امن و مورد اعتماد**
داده و مدل ها در زیرساخت داخلی (THDC) یا ناشر ذخیره و پردازش می شوند.
- کنترل دسترسی و ثبت وقایع**
دسترسی بر اساس نقش، ثبت کامل فعالیت ها و امکان حسابرسی.
- انطباق با سیاست ها و قوانین**
مطابق با سیاست های ناشر، مقررات نهادهای و الزامات قانونی.
- دستیار هوشمند برای داور**
خلاصه سازی، بررسی روش ها، تحلیل آماری، پرسش و پاسخ و شناسایی مسائل بالقوه.

نمونه کمک های AI در محیط مورد اعتماد

- ✓ خلاصه سازی بخش های کلیدی مقاله
- ✓ بررسی روش شناسی و طراحی مطالعه
- ✓ بررسی و تحلیل آماری
- ✓ شناسایی ابهام ها و تناقض ها
- ✓ مقایسه با مقالات مرتبط
- ✓ پیشنهاد پرسش ها به نویسنده



AI ابزاری قدرتمند برای کمک به داوری علمی است، اما استفاده مسئولانه، اخلاقی و منطبق با سیاست ها فقط در محیط های مورد اعتماد

پیام کلیدی:

بارگذاری تصویر بیمار در سامانه‌های هوش مصنوعی عمومی (مانند ChatGPT)

مسئله اصلی: فقط دقت سامانه نیست؛ مسئله اعتماد، حریم محرمانگی، حق بیمار و انطباق قانونی است.

4. محرمانگی، اعتماد و ارزیابی خود سامانه‌های AI (امتیاز 4 ستاره)	سطح	اصل	توضیح	امتیاز اعتماد
1	محرمانگی	اطلاعات بیمار بسیار نباید از امن محیط امن خارج شود.	★★★★	
2	یکپارچگی	داده باید کامل، صحیح و دستکاری نشده باقی بماند.	★★★★☆	
3	دسترسی‌پذیری	اطلاعات باید در زمان برای افراد مجاز در دسترس باشد.	★★★★☆	
4	قابلیت اعتماد به نتایج AI	صحت، حساسیت، پیرگی، AUC و ... با داده‌های معتبر معتبر نمایش‌دهنده شده باشد.	★★★★☆	
5	عدالت‌ها بودن و عدم سوگیری (Bias)	عملکرد سامانه برای گروه‌های مختلف جمعیتی عدالت‌ها و بدون نمایش‌دهنده شده باشد.	★★★★☆	
6	قابلیت تفسیر توضیح‌پذیری	سامان توانایی ارائه توضیح قابل برای نتایج خود را داشته باشد.	★★★★☆	
7	ایمنی مدیریت ریسک	ریسک‌های بالقوه مانند خطای تشخیص، هشدارهای کذب و ... و کنترل شده باشد.	★★★★☆	
8	پایش و به‌روزرسانی مستمر	عملکرد سامانه به طور مستمر پایش صورت نیاز و به‌روزرسانی شود.	★★★★☆	
اعتماد به AI فقط زمانی ایجاد می‌شود که تمامی ابعاد فصول به‌صورت مستند، قابل ارزیابی و پایش‌شده باشند.				

7. قابلیت‌های AI در محیط مطمئن (نمونه‌ها)
تشخیص و طبقه‌بندی ضایعات
اندازه‌گیری و آنالیز کمی
مقایسه طولی تصاویر
کمک به گزارش‌نویسی ساختار یافته
هشدارهای ایمنی و موارد مهم
همه این‌ها در محیط امن، قابل ردیابی و مطابق قوانین

3. راهکارهای پیشنهادی: استفاده از سرویس‌های محلی و مورد اعتماد
(On-Premise) پردازش داخل سازمان داده‌ها و مدل‌ها در محیط امن بیمارستان/سازمان و از سازمان خارج نمی‌شوند.
کنترل دسترسی و ثبت وقایع دسترسی فقط برای افراد مجاز با احراز هویت و ثبت کامل دسترسی (Audit Log).
انطباق با قوانین و سیاست‌ها انطباق کامل با قوانین حفاظت از داده‌ها، اخلاق پزشکی و سیاست‌های داخلی سازمان.
افزایش ایمنی و تأیید علمی سامانه مدل‌ها و عملکرد سامانه قبل از استقرار با داده‌های معتبر ارزیابی و تأیید شده است.
شفافیت و مسئولیت‌پذیری شفافیت در عملکرد، به‌مکملری، مشخص و پایداری پیکری در صورت بروز خطا.
مرکز داده سلامت مورد = THDC زیرساخت امن، بومی، قابل اعتماد و منطبق با استانداردهای ملی و بین‌المللی

6. چک‌لیست تصمیم‌گیری قبل از استفاده از AI عمومی
☒ آیا تصویر شامل اطلاعات شناسایی بیمار است؟
☒ اگر، بله، بارگذاری آن در AI عمومی مجاز نیست.
☒ آیا بیمار رضایت آگاهانه و کتبی را داده است؟
☒ آیا سازمان/بیمارستانی اجازه استفاده از داده‌های بیمار را داده است؟
☒ آیا سامانه AI مورد نظر دارای کاربری پزشکی معتبر و تأیید شده است؟
☒ آیا مسیر انتقال و محل ذخیره داده‌ها مشخص و امن است؟
☒ آیا خروجی AI صرفاً کمک‌کننده است و تصمیم‌گیری نهایی با پزشک است؟
☒ آیا مسئولیت حقوقی و حرفه‌ای به‌درستی تعیین شده است؟
در صورت وجود حتی یکی از موارد بالا، از بارگذاری خوداری از راهکارهای داخلی (مانند THDC استفاده نماید).

2. ریسک‌ها و چالش‌های استفاده از AI عمومی
نقض محرمانگی و حریم خصوصی خروج اطلاعات بیمار از کنترل کنترل، سازمان، قوانین ملی و قوانین ملی (GDPR, HIPAA)
انتقال داده به سرویس‌های خارج از کشور عدم کنترل محل ذخیره، مسیر انتقال و مدت نگهداری داده‌ها.
استفاده ثانویه از داده‌ها توسط ارائه‌دهنده امکان استفاده از داده‌ها برای آموزش مدل‌ها یا اهداف تجاری.
عدم شفافیت در نحوه پردازش و نگهداری مشخص نیست داده دقیق چگونه پردازش، ذخیره یا حذف می‌شود.
عدم انطباق با سیاست‌های سازمان و قوانین ممکن است با سیاست‌های امنیت اطلاعات، اخلاق و انسانی و حرفه‌ای و سلامت در تضاد باشد.
دقت و اعتبار علمی نامشخص سامانه مدل‌های عمومی برای کاربرد بالینی تأیید نشده‌اند؛ احتمال خطا، تحریف یا توهم (Hallucination) حسی وجود دارد.
مشکل این نیست AS، مشکل این است، محرمانگی و انطباق قانونی هر این عمومی AI تضمین نمی‌شود.

5. نکات کلیدی که باید در نظر گرفته شوند
حاکمیت و انطباق اعتبار و عملکرد AI اعتبارسنجی روی داده‌های بومی (قوانین ملی سلامت) گزارش عملکرد (دقت، حساسیت، پیرگی و ...) توضیح‌پذیری مدل (Explainability) به‌روزرسانی و مانیتورینگ مداوم
حقوق و رضایت بیمار اطلاع‌رسانی شفاف به بیمار دریافت رضایت آگاهانه (رضایت آگاهانه نیاز) احترام به حق انصراف
مدیریت داده و DICOM بررسی هدر DICOM حذف اطلاعات شناسایی ناشناس‌سازی استاندارد حذف‌سازی داده‌ها
امنیت و کنترل دسترسی دسترسی بر اساس نقش (RBAC) رمزنگاری در انتقال و ذخیره ورود و پایش مستمر مدیریت رخدادها

1. دو حالت بارگذاری تصویر
الف) تصویر شامل اطلاعات بیمار (DICOM) حاوی اطلاعات شناسی اطلاعات شناسایی با Patient Name: ... ID: ... Birth Date: ... Sex: ... Study Date: ... Modality: CT Institution: (نام، که تریخ تولید، مرکز مطالعه و ...)
ریسک بسیار بالا و غیرقابل قبول برای خروج از سازمان
ب) تصویر بدون اطلاعات شناسایی (ناشناس‌سازی شده) حذف اطلاعات از هدر DICOM حذف برجسته‌های متنی روی تصویر حذف هرگونه اطلاعات غیرضروری استفاده از کد تصادفی به جای شناسه بیمار
ریسک کمتر، ولی همچنان نیازمند ارزیابی و سیاست‌گذاری
ناشناس‌سازی 100% تضمین حذف ریسک نیست. بازشناسایی از طریق ترکیب داده‌ها همچنان ممکن است.