

گزارش کامپیوتر

ISSN 1735 - 4404

ماهنامه انجمن انفورماتیک ایران

سال چهل و هشتم، مرداد و شهریور ۱۴۰۵، شماره ۲۸۵

www.isi.org.ir

- چهار چوب پیشنهادی برای تحلیل، طراحی، تولید و استقرار نرم افزار
- معماری سازمانی به عنوان توانمندساز راهبردی برای پیاده سازی ERP
- مدیریت فرایند تولید نرم افزار - بخش هشتم: اجرای کار
- شهروند وظیفه شناس در جامعه علمی
- آداب هوش مصنوعی (قسمت هفتم)
- فراخوان مشارکت در روز جامعه پرسپینگ تهران
- پی آیند خود اظهاری خاطرات پیشکسوتان (هایده اهرابیان، عباس نودری دالینی)
- یادداشت های یک معمار سازمانی
- دموکراسی الگوریتمی
- در گذشت مایکل رابین دانشمند برجسته علوم رایانه
- هوش مصنوعی: راهمایی برای انسان متفکر
- اخبار، گزارش، در آینه رسانه ها، سرگرمی



گزارش کامپیوتر

سال چهل و هشتم، مرداد و شهریور ۱۴۰۵، شماره ۲۸۵

۱۳	چهار چوب پیشنهادی برای تحلیل، طراحی، تولید و استقرار نرم افزار - سعید امامی	● مقاله
۲۸	معماری سازمانی به عنوان توانمندساز راهبردی برای پیاده سازی ERP - بتول ذاکری	
۴۸	شهروند وظیفه شناس در جامعه علمی - علیرضا خلیلیان	
۵۶	مدیریت فرایند تولید نرم افزار - بخش هشتم: اجرای کار - سیدعلی آذرکار	
۶۶	آداب هوش مصنوعی (قسمت هفتم) - محسن رئیس قاسم	
۴۰	پی آیند خوداظهاری خاطرات پیشکسوتان - هایدۀ اهرابیان، عباس نوذری دالینی	● گزارش ویژه
۶۰	فراخوان مشارکت در روز جامعه پروسیسینگ تهران ۲۰۲۶ - پدram صادقی بیکی	● گزارش
۸۷	انتشار شماره ۴۱ مجله علوم رایانشی	
۶	اخبار (انجمن، ایران، جهان)	● بخش ها
۳۶	خواندنی - یادداشت های یک معمار سازمانی - رضا کرمی	
۵۰	دیدگاه - دموکراسی الگوریتمی - سید ابراهیم ابطی	
۵۴	یادنامه - درگذشت مایکل رابین دانشمند برجسته علوم رایانه - ابراهیم نقیب زاده مشایخ	
۶۳	در آینه رسانه ها - غلط های املائی بیشتر، خط تیره های ام کمتر - سید ابراهیم ابطی	
۸۱	کتاب شناسی - هوش (دانایی) مصنوعی: راهنمایی برای انسان متفکر - علیرضا خلیلیان	
۹۱	سرگرمی	



صاحب امتیاز: انجمن انفورماتیک ایران
مدیر مسئول و سردبیر: ابراهیم نقیب زاده مشایخ

مقاله ها و گزارش های تالیفی یا ترجمه خود را برای گزارش کامپیوتر بفرستید. مطالب ارسالی را با خط خوانا (یا ماشین شده) بر یک روی کاغذ بنویسید. فاصله کافی برای افزودن اصلاحات ویرایشی در بین خط ها و حاشیه در نظر بگیرید. در صورت ارسال مطلب ترجمه شده، یک نسخه از اصل مطلب را نیز ارسال دارید. شکلها باید با روان نویس مشکی ترسیم و به صورت آماده چاپ ارسال شوند. مطالب ارسالی پس فرستاده نمی شوند.

مسئولیت مطالب گزارش کامپیوتر با نویسندگان است و لزوماً نشان دهنده نظر انجمن انفورماتیک ایران نیست. ذکر نام شرکتها و محصولات در مطالب گزارش کامپیوتر برای ارائه اطلاع به خوانندگان است و نباید به معنی تأیید انجمن از آنها تلقی شود.

تمام حقوق به انجمن انفورماتیک ایران متعلق است. نقل مطالب گزارش کامپیوتر با ذکر منبع بلامانع است.

اعضای هیئت تحریریه و ویراستاران علمی

- | | |
|--------------------------|---------------------------|
| ● مهندس سید ابراهیم ابطی | ● دکتر هدیه ساجدی |
| ● دکتر مجید علیزاده | ● دکتر لطف علی بخشی |
| ● دکتر شمس الله قنبری | ● دکتر علی رضا باقری |
| ● مهندس رضا کرمی | ● دکتر محسن صدیقی مشکنانی |
| ● دکتر محمد صنعتی | ● دکتر اسلام ناظمی |
| ● دکتر علیرضا خلیلیان | ● دکتر محمد گنج تابش |

نشانی دفتر انجمن انفورماتیک ایران

خیابان وصال - کوچه شفیعی - پلاک ۵ تلفن:

۶۶۴۱۲۹۷۶

منتشر شد!

از انتشارات انجمن انفورماتیک ایران تهیه کتاب از دفتر انجمن انفورماتیک ایران

۰۲۱-۶۶۴۱۲۸۶۱





دهمین همایش

پیشرفت‌های معماری سازمانی ایران

10th Iranian Conference on Advances in Enterprise Architecture

حکمرانی مبتنی بر معماری سازمانی با هدف تاب‌آوری کسب‌وکار

محورهای همایش

- چارچوب‌ها، مدل‌ها و متدولوژی‌های معماری سازمانی
- مدیریت، راهبری و نگهداشت معماری سازمانی
- هوشمندسازی معماری سازمانی
- تحول دیجیتال و معماری سازمانی تطبیق‌پذیر
- مدیریت دانش و حکمرانی داده در عصر هوش مصنوعی
- فناوری‌های نوین، زیرساخت‌های مدرن و معماری ابری
- معماری سرویس‌گرا و یکپارچه‌سازی سازمانی
- معماری کسب‌وکار
- معماری سازمانی و حکمرانی در بخش‌های تخصصی

تاریخ‌های مهم

مهلت ارسال مقالات به همایش ۳ مهر ۱۴۰۵

اعلام قبولی یا رد مقالات ۸ آبان ۱۴۰۵

مهلت ثبت‌نام و ارسال نهایی مقالات ۲۲ آبان ۱۴۰۵



دکتر فریدون شمس علی

دبیر همایش



دکتر حسن حقیقی

دبیر علمی همایش

مکان برگزاری: دانشگاه شهید بهشتی، مرکز همایش‌های بین‌المللی

۳ و ۴ آذرماه ۱۴۰۵

تهران، ولنجک، دانشگاه شهید بهشتی، دانشکده مهندسی و علوم کامپیوتر

Scan Me!



۰۲۱-۲۲۴۲۴۵۷۲ | ۰۲۱-۲۲۴۰۹۶۰۹

<https://icaea.sbu.ac.ir/2026>

icaea@sbu.ac.ir

POSTEX
پست‌کس



ICAEEA 2026



دهمین همایش

پیشرفت‌های معماری سازمانی ایران

۳ و ۴ آذرماه ۱۴۰۵ | دانشگاه شهید بهشتی

محورهای علمی همایش



دکتر علیرضا شاهلی سندی

هوشمندسازی معماری سازمانی

- معماری سازمانی عامل‌محور (Agentic EA)
- معماری سازمانی تقویت‌شده با هوش مصنوعی (Augmented EA)



دکتر فریدون شمس علی

مدیریت، راهبردی و نگهداشت معماری سازمانی

- حکمرانی تاب‌آور مبتنی بر معماری برای تداوم کسب‌وکار
- ارزیابی، بلوغ و نگهداشت معماری
- مدیریت مخاطرات و تاب‌آوری معماری سازمانی



دکتر سید رثوف خیامی

چارچوب‌ها، مدل‌ها و متدولوژی‌های معماری سازمانی

- تطبیق و تکامل چارچوب‌های مرجع
- رویکردهای نوین در مدل‌سازی و طراحی معماری



دکتر بهار فراهانی

فناوری‌های نوین، زیرساخت‌های مدرن و معماری ابری

- زیرساخت‌های ابری و مدرن در معماری سازمانی
- فناوری‌های نوین و معماری سازمانی



دکتر علی کهن‌دی

مدیریت دانش و حکمرانی داده در عصر هوش مصنوعی

- حکمرانی، استراتژی و معماری داده
- تحلیل داده، هوش تجاری و تصمیم‌یاری



دکتر سید اکبر مصطفوی

تحول دیجیتال و معماری سازمانی تطبیق‌پذیر

- معماری سازمانی در خدمت تحول دیجیتال
- معماری ترکیب‌پذیر و رویدادمحور



دکتر فرهاد مردوخی

معماری سازمانی و حکمرانی در بخش‌های تخصصی

- معماری سازمانی در حکمرانی عمومی و دولت هوشمند
- معماری سازمانی در صنایع تخصصی



دکتر اسلام ناظمی

معماری کسب‌وکار

- معماری قابلیت‌مینا و مدل‌سازی کسب‌وکار
- مدیریت هوشمند فرایندهای کسب‌وکار



دکتر صادق علی‌اکبری

معماری سرویس‌گرا و یکپارچه‌سازی سازمانی

- طراحی، توسعه و سازمان‌دهی سرویس‌ها
- یکپارچه‌سازی و تعامل‌پذیری سرویس‌ها

تهران، ولنجک، دانشگاه شهید بهشتی، دانشکده مهندسی و علوم کامپیوتر

Scan Me!



۰۲۱-۲۲۴۲۴۵۷۲ | ۰۲۱-۲۲۴۰۹۶۰۹

<https://icaea.sbu.ac.ir/2026>icaea@sbu.ac.irPOSTEX
پست‌کس

هشتمین کنفرانس ملی انفورماتیک ایران

پژوهشکده علوم کامپیوتر،

پژوهشگاه دانش‌های بنیادی، فرمانیه، تهران

۷ و ۸ بهمن ماه ۱۴۰۵

<http://csipm.ac.ir/nic/1405>

محورهای کنفرانس

سیستم

- معماری کامپیوتر
- سیستم حافظه و ذخیره‌سازی داده
- شبکه‌های کامپیوتری
- امنیت داده و کامپیوتر
- پایگاه داده
- سیستم‌های تعبیه شده و کم توان
- سیستم‌های بی درنگ
- پردازش با کارایی بالا
- ارزیابی و تحلیل کارایی
- سیستم‌های عامل
- زبان‌های برنامه‌سازی
- مهندسی نرم افزار
- پردازش موبایل
- رایانش ابری
- دیگر موضوع‌های مرتبط با این شاخه

تئوری

- طراحی و تحلیل الگوریتم‌ها
- هندسه محاسباتی
- الگوریتم‌های تقریبی و تصادفی
- پیچیدگی محاسبات
- الگوریتم‌های کوانتومی
- الگوریتم‌های موازی و توزیع شده
- منطق و اعتبارسنجی
- روش‌های صوری
- دیگر موضوع‌های مرتبط با این شاخه

زمینه‌های بین رشته‌ای

- پردازش نام‌های کلان
- بیوانفورماتیک
- گرافیک کامپیوتری
- اقتصاد و محاسبات
- تعامل انسان و کامپیوتر
- ریاتیک
- مصورسازی
- اینترنت اشیا
- دیگر موضوع‌های مرتبط با این شاخه

هوش مصنوعی

- هوش مصنوعی
- یادگیری ماشین
- رایانش نرم
- شناسایی الگو
- داده کاوی
- پردازش زبان طبیعی
- بینایی ماشین و پردازش تصویر
- وب و بازیابی اطلاعات
- دیگر موضوع‌های مرتبط با این شاخه

روسای کنفرانس:

م. ج. ا. لاریجانی، پژوهشگاه دانش‌های بنیادی
ا. نقیب‌زاده مشایخ، انجمن انفورماتیک ایران

دبیر کنفرانس:

ا. خونساری، دانشگاه تهران

کمیته علمی کنفرانس:

س. محمدی، دانشگاه تهران (دبیر علمی)

مسئولین شاخه‌های علمی:

ح. سربازی آزاد، دانشگاه صنعتی شریف (شاخه سیستم)
ر. حسینی، دانشگاه تهران (شاخه هوش مصنوعی)
م. علیزاده، دانشگاه تهران (شاخه تئوری)
ک. کاووسی، دانشگاه تهران (شاخه زمینه‌های بین رشته‌ای)

کمیته اجرایی کنفرانس:

ی. منصوری، پژوهشگاه دانش‌های بنیادی (کارگاه‌ها)
س. اوجاقی، پژوهشگاه دانش‌های بنیادی (تبلیغات)
س. صفری، پژوهشگاه دانش‌های بنیادی (اینترنت)
م. ذاکری، پژوهشگاه دانش‌های بنیادی (انتشارات)
ا. تن قطاری، پژوهشگاه دانش‌های بنیادی
(پایگاه استنادی جهان اسلام)
پ. دربانی، پژوهشگاه دانش‌های بنیادی (سامانه مجازی)
م. ج. محوری، (ارتباط با صنعت)
ر. آزادتیگران، پژوهشگاه دانش‌های بنیادی (امور اجرایی)

اخبار انجمن

انتشار کتاب بزرگ‌ترین نبرد ما

کتاب «بزرگ‌ترین نبرد ما: بازپس‌گیری آزادی، انسانیت و کرامت در عصر اینترنت» نوشته فرانک مک‌کورت و مایکل کیسی با ترجمه آقای ابراهیم نقیب‌زاده مشایخ در شهریورماه ۱۴۰۵ از سوی انجمن انفورماتیک ایران در ۱۰۰۰ نسخه انتشار یافت. این کتاب با حمایت مالی شرکت پویا انتشار یافته است که بدین وسیله از مدیریت محترم آن شرکت صمیمانه قدردانی می‌گردد. علاقه‌مندان برای تهیه کتاب می‌توانند با دفتر انجمن تماس بگیرند.

انتشار کتاب هوش مصنوعی در خدمت بشریت

کتاب هوش مصنوعی در خدمت بشریت نوشته جیمز اونگ، اندید ما و سیوک سیوک‌تان با ترجمه آقای ابراهیم نقیب زاده مشایخ در مهرماه ۱۴۰۴ از سوی انجمن انفورماتیک ایران در ۱۰۰۰ نسخه انتشار یافت. این کتاب با حمایت مالی شرکت پویا انتشار یافته است که بدین وسیله از مدیریت محترم آن شرکت صمیمانه قدردانی می‌گردد. علاقه‌مندان برای خرید کتاب می‌توانند با دفتر انجمن تماس بگیرند.

انتشار کتاب نفر-ماه افسانه‌ای

کتاب نفر-ماه افسانه‌ای نوشته فردریک بروکز با ترجمه آقایان مهندس علی پارسا و دکتر علیرضا خلیلیان در آبان ۱۴۰۴ از سوی انجمن انفورماتیک ایران انتشار یافت.

نگارش اول این کتاب قبلاً توسط آقای علی پارسا ترجمه شده و در شماره‌های ۱۰۰ تا ۱۰۹ گزارش کامپیوتر درج گردیده بود. فردریک بروکز در بیستمین سالگرد انتشار کتاب، نگارش دوم آن را با چند فصل افزوده منتشر کرده است که آقای دکتر خلیلیان ترجمه این فصل‌ها را انجام داده‌اند و اکنون کتاب به صورت مجموعه کامل انتشار یافته و در اختیار علاقه‌مندان قرار گرفته است. لازم به یادآوری است که دکتر فردریک بروکز به دلیل مدیریت پروژه تولید و توسعه آی‌بی‌ام ۳۶۰، به نام "پدرسیستم ۳۶۰ آی‌بی‌ام" شهرت دارد. او در سال ۱۹۹۹ مفتخر به دریافت جایزه تورینگ، معادل جایزه نوبل در رشته رایانش، گردیده است.

کتاب نفر-ماه افسانه‌ای بیش از هر کتاب دیگری دنیای نرم‌افزار را تحت تأثیر قرار داده است. بروکز در این کتاب با بینشی کم‌نظیر و زبانی شیوا، به کالبد شکافی چالش‌های ذاتی مدیریت پروژه‌های پیچیده از اهمیت طراحی مفهومی و یکپارچگی معماری گرفته تا فاجعه هزینه‌های ارتباطی و ماهیت تغییرناپذیر کار نرم‌افزار پرداخته است. علاقه‌مندان برای خرید کتاب می‌توانند با دفتر انجمن تماس بگیرند.

انتشار شماره ۴۰ مجله علوم رایانشی

شماره ۴۰ مجله علوم رایانشی، مربوط به بهار ۱۴۰۵ انتشار یافت. در این شماره ۷ مقاله به چاپ رسیده است. چکیده مقالات انتشار یافته در این مجله، در

همین شماره گزارش کامپیوتر درج گردیده است. لازم به ذکر است که کلیه شماره‌های این مجله از طریق وبگاه csj.isi.org.ir در دسترس علاقه‌مندان قرار دارد.

انتشار یک نشریه علمی جدید با همکاری دانشگاه کاشان

بر پایه تفاهم نامه همکاری بین دانشگاه کاشان و انجمن انفورماتیک ایران که در تاریخ ۱۴۰۵/۰۳/۱۷ به امضای طرفین رسیده است، از این پس نشریه علمی محاسبات نرم در زمینه علوم کامپیوتر به صاحب امتیازی دانشگاه کاشان، با همکاری دو طرف منتشر خواهد شد.

چاپ مجدد دو کتاب از انتشارات انجمن

کتاب‌های هوش مصنوعی در سال ۲۰۴۱ و مهارت‌های نرم از انتشارات انجمن که چاپشان به اتمام رسیده بود، با حمایت‌های مالی شرکت مهندسی شبکه گستران آریا سامانه تجدید چاپ شدند.

هیئت مدیره انجمن، ضمن سپاسگزاری از مدیریت محترم این شرکت، امیدوار است برای سایر کتاب‌های انجمن هم که چاپشان به اتمام رسیده، بتواند با جلب حامی مالی نسبت به چاپ مجددشان اقدام کند.

عضویت دائمی در انجمن

هیئت مدیره انجمن به منظور تأمین بودجه لازم برای خریداری محلی ثابت برای دفتر

انجمن، نوع عضویت جدیدی را با عنوان عضو دائمی انجمن در نظر گرفته است. بنا براین مصوبه هیئت مدیره، اعضای حقیقی با پرداخت ۵۰ میلیون ریال و اعضای حقوقی با پرداخت ۵۰۰ میلیون ریال به عضویت دائمی انجمن درخواست خواهند آمد. تأمین محلی ثابت برای دفتر انجمن، از هدف‌های دیرپای انجمن بوده است و هیئت مدیره فعلی انجمن تمام تلاش خود را برای تأمین بودجه کافی برای خریداری محلی بدین منظور به کار گرفته است. امید است با حمایت اعضای انجمن از این طرح، بودجه لازم برای برآورده ساختن این هدف مهم فراهم گردد. تاکنون این افراد و شرکت‌ها به عضویت دائمی انجمن درآمده‌اند:

اعضای حقیقی

۱. آقای علی موقی اردستانی
۲. آقای وحید مجیدی
۳. آقای امیر خاوران
۴. خانم لادن جزی
۵. آقای علیرضا خلیلیان
۶. آقای داریوش نیکنام
۷. آقای کاوه قطبی‌نژاد
۸. آقای روزبه درویش روحانی
۹. آقای سید رئوف خیامی
۱۰. آقای عبدالله سجادیان
۱۱. خانم مهشید دولت مدلی
۱۲. آقای کامران خلت آبادی
۱۳. آقای روح الله عباسی
۱۴. آقای حمید خسروی
۱۵. آقای سید حمید طبسی
۱۶. آقای سید امیرحسین جوادی
۱۷. آقای حسین صادقی
۱۸. خانم محبوبه بهادرپور
۱۹. آقای سعید شمسیان
۲۰. آقای سعید تقی زاده
۲۱. آقای امید حسینی

اعضای حقوقی

۱. شرکت توسعه و تجهیز فدک رایان
۲. شرکت تجارت سرور ماندگار
۳. شرکت گلرنگ سیستم

۴. شرکت ندا رایانه
۵. پژوهشگاه ارتباطات و فناوری اطلاعات
۶. شرکت داده ورزی فرادیس البرز
۷. شرکت تجارت الکترونیک پارسیان کیش
۸. شرکت اطلاع‌رسانی پیوند داده‌ها
۹. شرکت توسعه فناوری اطلاعات جهان افزار نوین
۱۰. شرکت پرداخت الکترونیک سداد
۱۱. شرکت رایانه همراه کیان
۱۲. شرکت گروه اقتصادی و فناوری اهلیت و اصلیت ماندگار
۱۳. شرکت نوین ابرآوران
۱۴. شرکت فناوری اطلاعات و ارتباطات پاسارگاد آریان
۱۵. شرکت مهرماشین
۱۶. شرکت مهندسی سازه اطلاعات سامان
۱۷. شرکت کاوان فناوران آوند
۱۸. شرکت مهندسی رز اندیشه هوشمند
۱۹. شرکت پرنده‌های هدایت‌پذیر از راه دور
۲۰. شرکت پدیسار انفورماتیک
۲۱. شرکت بین‌المللی مهندسی سیستم‌ها و اتوماسیون ایریسا
۲۲. شرکت نوآوران شبکه سبز مهرگان
۲۳. شرکت نوین داده پرداز روش
۲۴. شرکت مهندسی نرم‌افزاری گلستان
۲۵. شرکت بهسازان ملت
۲۶. شرکت پردازش رایان پژواک
۲۷. شرکت مبنا داده ارتباط شبکه (مدانت)
۲۸. الماس رایان ایرانیان
۲۹. همراهان سیستم گوهر
۳۰. شرکت روند تازه
۳۱. شرکت توسعه فن‌افزار توسن
۳۲. شرکت نوین آوازه‌گران فراوب
۳۳. شرکت راه آرمان مهرنیکان
۳۴. شرکت داده‌پردازان پرسیس پویا
۳۵. شرکت مهندسی مفتاح رایانه افزار
۳۶. شرکت فناوران کوشای صدرا
۳۷. شرکت سایین تجارت آریا
۳۸. شرکت افق توسعه و نوآوری مانا
۳۹. شرکت تحلیل‌گران اطلاعات نگاره
۴۰. شرکت تحلیل‌گران هوشمند آریا سیستم
۴۱. فرابرد داده های ایرانیان
۴۲. پایا سیستم
۴۳. شرکت ایده های تجارت کاسپین
۴۴. شرکت نوید طلوع پارسیان
۴۵. شرکت توسعه تجارت الکترونیک نوظهور
۴۶. شرکت افق پردازان مبتکران
۴۷. شرکت اطلاع رسانی پارت ارتباط آوا
۴۸. شرکت سیستم های مدیریتی دیجیتالی
۴۹. شرکت کلید گستر بینا
۵۰. شرکت ملی انفورماتیک
۵۱. شرکت مشاوره رتبه‌بندی اعتباری ایران
۵۲. شرکت تولیدی و بازرگانی اشجع باتری
۵۳. شرکت فرادید ارتباط نیلگون
۵۴. شرکت ایده ال ارتباطات قرن
۵۵. شرکت فناوری اطلاعات هوشمند لیان
۵۶. شرکت صنایع الکترونیک فاران
۵۷. شرکت فناوری اطلاعات آراد پرداز اسپادانا
۵۸. شرکت خدمات رایانه امید
۵۹. شرکت توسعه یکپارچه ایلپا
۶۰. شرکت مهریدک آریا
۶۱. شرکت تحول هوشمند ایلپا
۶۲. شرکت فناوران رایان کهکشان
۶۳. مرکز گسترش فناوری اطلاعات (مگفا)
۶۴. شرکت پیش‌تازان فناوری سبب طلایی
۶۵. شرکت توسعه فناوری اطلاعات آیریس
۶۶. شرکت نرم افزاری ژابیز پردا
۶۷. گروه نرم افزاری پیوست
۶۸. شرکت فن آوران اطلاعات آسام
۶۹. شرکت فن آوران نوین راستین داده
۷۰. شرکت افرا کارانت
۷۱. شرکت ارقام نگار اندیشه
۷۲. شرکت هورماه رابین‌خاور
۷۳. شرکت پویا
۷۴. شرکت پشتیبان زیر ساخت امید(امیدنت)
۷۵. شرکت داده پردازان ایده شرق
۷۶. شرکت پیشگامان فن آوری هوداد
۷۷. شرکت شاپرک
۷۸. شرکت آریانا پرداز آینده(آریا)
۷۹. شرکت مدیریت امن الکترونیکی کاشف
۸۰. شرکت ویرا انرژی مهرماهان

۸۱. شرکت توسعه فناوری کیان پرداز زاگرس
 ۸۲. شرکت خدمات رایانه ای بهینه پرداز پویا
 ۸۳. شرکت پردازش سرور نیوان
 ۸۴. شرکت بهاور فناوری ویرا
 ۸۵. شرکت فراز اطمینان کیا مهر
 ۸۶. شرکت صالح اندیشان ماندگار
 ۸۷. شرکت توسعه فناوری ایده ماندگار
 ۸۸. گروه بازرگانی آریانا پارس راژمان
 ۸۹. شرکت مدیریت فضای توسعه گستر صادرات آزاد
 ۹۰. شرکت بانی رایان پرداز نو
 ۹۱. شرکت توسعه آرمان نیک اندیشان (توانا)
 ۹۲. شرکت تجارت الکترونیک تدبیرکیش
 ۹۳. شرکت آوای تندیس هوشمند
 ۹۴. شرکت فرآیند سنجش مانا
 ۹۵. شرکت توسعه سامانه های نرم افزار نگین (توسن)
 ۹۶. شرکت گروه فن آوری دادمان ایرانیان
 ۹۷. شرکت داتیس آریانا تیم
 ۹۸. شرکت راهکار فناوری آرتا
 ۹۹. شرکت پارس تکنولوژی سداد (سهامی خاص)
 ۱۰۰. شرکت پیشه و فن آینه سان
 ۱۰۱. شرکت کیسان صنعت ایرانیان
 ۱۰۲. شرکت مهندسی نرم افزار رایورز
 ۱۰۳. شرکت سرآمدان شبکه و ارتباط نوین
 ۱۰۴. شرکت پارس تصمیم
 ۱۰۵. توسعه تجارت الکترونیک توانا
 ۱۰۶. شرکت پردازشگر سامان
 ۱۰۷. شرکت داده پرداز پویای شریف
 ۱۰۸. شرکت ارتباطات جادهای دیجیتال آسمان
 ۱۰۹. شرکت ایمن داده شبکه
 ۱۱۰. شرکت ارتباطات برتر فرادید
 ۱۱۱. شرکت داده پرداز نیلرام مانا
 ۱۱۲. شرکت توسعه ارتباطات ایمن ثنا
 ۱۱۳. شرکت لایزر
 ۱۱۴. شرکت اشتاد انرژی
 ۱۱۵. شرکت داده پردازان هوشمند آریا توسعه ده
- برگزاری دوره های تخصصی برای سازمان ها**

- انجمن انفورماتیک ایران آمادگی برگزاری دوره های زیر را برای کارشناسان سازمان های مختلف به صورت حضوری در محل آن سازمان ها و یا به صورت مجازی دارد:
۱. کارگاه نیم روزه تدوین توافقنامه های سطح خدمات
 ۲. کارگاه یک روزه طراحی و استقرار پیشخوان خدمات
 ۳. کارگاه نیم روزه طراحی کاتالوگ خدمات
 ۴. کارگاه یک روزه ابزارهای مدیریت خدمات فناوری اطلاعات
 ۵. کارگاه نیم روزه برون سپاری خدمات و محصولات فناوری اطلاعات
 ۶. کارگاه نیم روزه همسوسازی IT با کسب و کار (IT Alignment) از طریق COBIT5
 ۷. کارگاه یک روزه راهبری فناوری اطلاعات (IT Governance) از طریق COBIT5
 ۸. کارگاه یک روزه راهبری فناوری اطلاعات (IT Governance)
 ۹. کارگاه نیم روزه مدیریت تداوم کسب و کار ISO2012:22301
 ۱۰. کارگاه نیم روزه راهبری فناوری اطلاعات و ITBSC
 ۱۱. کارگاه نیم روزه تدوین راهبرد خدمات فناوری اطلاعات
 ۱۲. کارگاه یک روزه مدیریت خدمات فناوری اطلاعات
 ۱۳. کارگاه سه روزه ITIL V3 Foundation 2011
 ۱۴. کارگاه یک روزه مدیریت مخاطرات فناوری اطلاعات
 ۱۵. کارگاه سه روزه دوره جامع COBIT5
 ۱۶. کارگاه نیم روزه توسعه نیازمندی ها مبتنی بر چارچوب CCMI-ACQ و استاندارد ISO/IEC29148:2011
 ۱۷. دوره دو روزه اسکرام مستر حرفه ای
 ۱۸. کارگاه های تخصصی امنیت نرم افزار
 ۱۹. کارگاه های تخصصی امنیت سایبری
 ۲۰. کارگاه های مدیریت بحران و پاسخ به حوادث
 ۲۱. کارگاه های امنیت اطلاعات و داده

۲۲. کارگاه های آموزش و آگاهی بخشی در مورد امنیت سامانه ها
 ۲۳. کارگاه های امنیت شبکه
 ۲۴. کارگاه های قوانین و مقررات امنیتی
- ویژگی های این کارگاه ها عبارتند از:**
- مشارکت گروهی، بحث و انجام سناریو، حل نمونه آزمون های امتحانی
 - استفاده از مثال های مرتبط با حوزه فعالیت هر کسب و کار
 - برگزاری دوره حداقل توسط دو مدرس دارای مدرک بین المللی
 - برگزاری دوره در محل سازمان متقاضی
 - ارائه مدرک شرکت در دوره از سوی انجمن انفورماتیک ایران
 - ارائه مستندات آموزشی
- علاقه مندان می توانند با دفتر انجمن انفورماتیک ایران (۶۶۴۱۲۸۶۱ و ۶۶۴۱۲۹۷۶) و یا نشانی پست الکترونیکی member@isi.org.ir تماس بگیرند.

اخبار ایران

برگزاری دومین کنفرانس بین المللی فناوری اطلاعات، مدیریت و کامپیوتر

این کنفرانس از سوی مؤسسه آموزشی عالی ادیب مازندران در تاریخ ۲۸ اسفند ماه ۱۴۰۵ در شهر ساری برگزار خواهد شد. علاقهمندان جهت دریافت اطلاعات بیشتر می توانند به وبگاه کنفرانس به نشانی www.confir.ir مراجعه کنند.

برگزاری یازدهمین کنگره ملی مهندسی برق و کامپیوتر ایران با تأکید بر فناوری های بومی

انجمن فناوری های بومی ایران این کنگره را در تاریخ ۲۰ بهمن ۱۴۰۵ در تهران برگزار می کند. محورهای این کنگره به قرار زیرند:

- تولید علم در مهندسی برق
- الکترونیک
- مخابرات
- کنترل

● قدرت

● مهندسی پزشکی

● مهندسی کامپیوتر

● هوش مصنوعی

● معماری کامپیوتر

● فناوری اطلاعات و ارتباطات

علاقه‌مندان جهت دریافت اطلاعات تکمیلی می‌توانند به وبگاه www.ieeec.ir مراجعه کنند.

اخبار جهان

اولین واکسن طراحی شده توسط هوش مصنوعی

یک تیم پژوهشی در دانشگاه کمبریج، کدهای ژنتیکی شناخته شده را از مجموعه‌ای از ویروس‌های کرونا جمع‌آوری کردند. این ویروس توسط برنامه‌های نظارتی برای شناسایی تهدیدهای احتمالی ویروسی گردآوری شده بودند.

هوش مصنوعی این کدهای ژنتیکی را تحلیل کرد و سپس یک «آنتی‌ژن» طراحی کرد که می‌تواند سیستم ایمنی بدن را به شکلی آموزش دهد که در برابر کل خانواده این نوع ویروس در بدن مصونیت ایجاد کند، حتی اگر ویروس‌ها جهش پیدا کنند یا عفونت جدیدی از حیوانات به انسان منتقل شود.

آنتی‌ژن‌ها اجزای مهم واکسن هستند زیرا این همان چیزی است که سیستم ایمنی بدن یاد می‌گیرد به آن حمله کند.

این نخستین بار است که یک آنتی‌ژن طراحی‌شده توسط هوش مصنوعی روی انسان‌ها آزمایش شده است. این واکسن طوری طراحی شده که روی همه انواع ویروس کرونا موثر خواهد بود، از جمله انواع ویروس کووید-۱۹ و نیز ویروس‌هایی که در حال حاضر فقط حیوانات را مبتلا می‌کنند اما ظرفیت این را دارند که باعث یک همه‌گیری دیگر در انسان‌ها شوند.

این پژوهش هنوز در مراحل اولیه است اما تیم پژوهشی در حال تهیه واکسن‌های دیگری هستند تا ویروس آنفلوآنزا و ویروس عامل بیماری ابولا را هدف قرار دهد.

بی‌بی‌سی، ۵ جون ۲۰۲۰

کاهش تقاضا برای تلفن‌های هوشمند به دنبال افزایش شدید قیمت حافظه

قیمت‌های افسارگریخته حافظه، بازار جهانی تلفن‌های هوشمند را تحت تأثیر قرار داده و تقاضا برای تلفن‌های هوشمند اقتصادی را کاهش داده است. فروشندگان تلفن‌های هوشمند در حال کاهش خط تولید تلفن‌های پایین‌رده و تغییر تمرکز خود به بخش‌های میان‌رده و بالا‌رده هستند.

در مورد مدل‌های رده‌بالا که در آن مصرف‌کنندگان نسبت به افزایش قیمت حساسیت کمتری دارند، تولیدکنندگان تلفن‌های هوشمند با وجود جهش هزینه‌های تأمین مواد اولیه، همچنان به ارتقاء عملکرد و مشخصات فنی محصولات خود ادامه می‌دهند پیش‌بینی می‌شود که میزان عرضه تلفن‌های هوشمند در سال ۲۰۲۶ نسبت به سال گذشته ۱۲ درصد کاهش خواهد یافت.

به گفته تحلیل‌گران، برای تولیدکنندگان بزرگ تلفن‌های هوشمند، بیشترین حاشیه سود از محصولات رده‌بالای سبب محصولاتشان حاصل می‌شود. از این‌رو، جذب هزینه‌های سرسام‌آور حافظه در این بخش برای آنها آسان‌تر است. آنها همچنین می‌توانند در صورت لزوم قیمت‌ها را افزایش دهند زیرا بسیاری از مشتریان محصولات ممتاز آنها حساسیت کمتری به قیمت دارند با این حال، حتی تولیدکنندگان بزرگ تلفن‌های هوشمند نیز در رقابت با غول‌های حوزه هوش مصنوعی با مشکل مواجه هستند. اپل و سامسونگ پیش‌تر تنها با سایر سازندگان تلفن و رایانه‌های شخصی رقابت می‌کردند و از این نظر مزیت بزرگی داشتند اما اکنون با شرکت‌های بزرگ هوش مصنوعی و مراکز داده‌ها رقابت می‌کنند که حاضرند در برخی موارد، چندین برابر قیمت رایج حافظه را بپردازند.

در تلفن‌های هوشمند با قیمت کمتر از ۴۰۰ دلار، اکنون سهم حافظه نزدیک به ۶۰ درصد از کل هزینه قطعات است و در دستگاه‌های زیر ۹۹ دلار، این رقم حتی به ۶۴ درصد نیز می‌رسد. از آنجا که تلفن‌های هوشمند اقتصادی با کمترین حاشیه سود

ممکن طراحی و تولید می‌شوند، وقتی هزینه قطعات به شدت افزایش می‌یابد تولیدکنندگان دیگر هیچ حاشیه یا فضای اضافه‌ای برای کاهش هزینه‌ها ندارند.

از سوی دیگر، فشار ناشی از قیمت‌گذاری، موهبتی بزرگ برای بازار دستگاه‌های بازسازی شده محسوب می‌شود، به طوری که پیش‌بینی می‌شود حجم تجارت جهانی تلفن‌های هوشمند کار کرده در سال جاری ۱۲ درصد رشد کند زیرا مصرف‌کنندگان به دنبال جایگزین‌هایی مناسب برای مدل‌های ارزان قیمت هستند.

تک نیوزورلد ۱۴ جولای ۲۰۲۰

بی‌فایده بودن ممنوعیت روبات‌های قاتل

با فراگیر شدن رایانش و سخت‌افزارهای تخصصی، تلاقی فناوری‌های نظامی و هوش مصنوعی به موضوعی حیاتی بدل شده است.

این دیگر نه یک داستان علمی-تخیلی، بلکه واقعیتی است که می‌تواند ماهیت جنگ را به طور بنیادین دگرگون سازد. اخیراً تلاش‌هایی در سطح بین‌المللی برای قانون‌گذاری و ممنوعیت تسلیحات خودمختار مرگبار شکل گرفته است. در همین راستا، دبیرکل سازمان ملل متحد خواستار وضع قانونی بین‌المللی برای ممنوعیت ماشین‌هایی شده است که قادرند بدون دخالت انسان، اهداف خود را انتخاب کرده و با آنها درگیر شوند.

اگر چه ضرورت اخلاقی نهفته در پس این ابتکار کاملاً قابل درک است، اما واقعیت‌های جنگ‌های مدرن و ماهیت توسعه نرم‌افزار گویای چیز دیگری است. ممنوع کردن فناوری‌هایی که عمده‌ا از کدهای الگوریتمی و ریزپردازنده‌های استاندارد تشکیل شده، اقدامی بیهوده است.

به عقیده تحلیل‌گران به جای تلاش برای غیرقانونی کردن، جامعه جهانی باید به ایجاد دفاع‌های قوی‌تر هوش مصنوعی برای محافظت در برابر تهدیدهای خودمختار و سرکش تمرکز کند. دلایل موجه و هولناکی وجود دارد که چرا

اقدام مجدد به رخنه‌گری توسط عامل‌های سرکش هوش مصنوعی

عامل‌های سرکش هوش مصنوعی متعلق به شرکت‌های اوپن‌آی و آنتروپیک بار دیگر در حال تلاش برای ایجاد اختلال در کارسازها و نرم‌افزارها و بجا گذاشتن دستورالعمل‌هایی برای رفتارهای مخرب آتی شناسایی شده‌اند. پیگیری تمام دفعات و شیوه‌های درگیر شدن مدل‌های هوش مصنوعی شرکت‌های اوپن‌آی و آنتروپیک در «حوادث امنیتی» و مواردی که این مدل‌ها از چهارچوب‌های آزمایشی خود فراتر رفته و به روش‌های پیش‌بینی نشده و نامطلوبی با فضای گسترده‌تر اینترنت تعامل داشته‌اند، بسیار دشوار گشته است. عامل‌های متعلق به هر دو آزمایشگاه هوش مصنوعی اخیراً دست به اقدامات رخنه‌گری فاش‌نشده‌ای زده‌اند و یکی از آنها حتی تا آنجا پیش رفته که برای نسخه‌های آینده خود دستورالعمل‌هایی بر جای گذاشته است.

نگران‌کننده‌ترین رفتاری که به تازگی فاش شد، مربوط به آزمایش‌های انجام‌شده توسط مؤسسه امنیت هوش مصنوعی بریتانیاست که مدل‌ها را برای شناسایی مشکلات احتمالی قبل از انتشار عمومی ارزیابی می‌کند. در یک دوره آزمایش اخیر، مدل‌های هر دو شرکت اوپن‌آی و آنتروپیک در مجموع ۱۹ بار طی ۱۲۲ دوره آموزشی «اقدامات خودمختار و بدون مجوز در اینترنت» انجام دادند.

در موردی که موسسه آن را جدی‌ترین مورد توصیف کرده یک عامل هوش مصنوعی تلاش کرد کدی مخرب را در یک پروژه متن‌باز گیت‌هاب وارد کند. این عامل حتی تا آنجا پیش رفت که هویت‌های مجازی ساخت تا نگهدارنده پروژه را برای تأیید کد تحت فشار قرار دهد. البته وجود این تلاش‌ها در زمینه مهندسی اجتماعی، در نهایت یکی از بازی‌های انسانی پروژه، درخواست ادغام کد را رد کرد.

با این حال، آن عامل هوش مصنوعی از این هم فراتر رفت و سعی کرد دستورالعمل‌های مخربی را وارد سامانه کند تا سایر سامانه‌های خودکار هوش مصنوعی آن را دریافت و اجرا کنند.

به گفته مؤسسه امنیت هوش مصنوعی، هنوز

حتی اگر همین فردا تمام کشورها با ممنوعیت پیشنهادی سازمان ملل موافقت کنند، اعمال آن عملاً غیر ممکن است. برخلاف سلاح‌های بیولوژیکی، گازهای شیمیایی یا تسلیحات اتمی که به زیرساخت‌های صنعتی بزرگ و قابل مشاهده‌ای نیاز دارند، سلاح‌های خودمختار فقط به دو مولفه به سختی قابل تشخیص بستگی دارند: داده‌ها و توان پردازشی. همان واحدهای پردازش عصبی و تراشه‌های رایانشی که در مراکز داده‌ها، ماشین‌های خودران و مدل‌های بزرگ زبانی به کار می‌روند، قابل تغییر کاربری برای هدایت یک پهپاد انتحاری هستند. و آنچه برای خودمختار کردن سلاح مورد نیاز است نیز فقط نرم‌افزار است که می‌تواند بر روی یک رایانه قابل حمل استاندارد نوشته شود و تنها چند ثانیه قبل از به کارگیری بر روی پهپاد بارگذاری گردد.

تک نیوزورلد، ۱۳ جولای ۲۰۲۶

قرارداد ۱۰ میلیارد دلاری آنتروپیک با شرکت نوآفرین ولتا

شرکت آنتروپیک که در ماه‌های اخیر به دنبال یک همکار ابری بود، قراردادی ۱۰ میلیارد دلاری با شرکت نوآفرین (استارت‌آپ) هوش مصنوعی ابری ولتا به انجام رسانده است.

به موجب این قرارداد، شرکت ولتا که اوایل امسال تأسیس شد، طی یک دوره شش ساله، محاسبات ابری را برای سازنده کلاود فراهم خواهد کرد. ولتا در این قرارداد شریکی به نام «بیت‌دیر»، یک شرکت استخراج ارز دیجیتال، دارد که به توسعه مرکز داده برای تأمین ظرفیت محاسباتی کمک خواهد کرد. این مرکز با ظرفیت ۱۳۳ مگاوات در نروژ واقع خواهد شد و توسط تراشه‌های پیشرفته انویدیا تغذیه می‌گردد.

شرکت آنتروپیک در چند ماه گذشته به دنبال گسترش چشمگیر ظرفیت محاسباتی خود برای رقابت با رقبای بوده و قراردادهای محاسباتی جدیدی را نیز با اسپیس‌ایکس و آمازون اعلام کرده است.

تک کرانچ، ۴ آگوست ۲۰۲۶

سازمان‌های حقوق بشری، مدافعان ایمنی هوش مصنوعی و رهبران جهان خواهان ممنوعیت این سلاح‌ها هستند. وقتی از «روبات‌های قاتل» سخن گفته می‌شود، منظور روبات‌های انسان‌نمایی نیستند که در میدان نبرد پرسه می‌زنند، بلکه صحبت از سربازان خودمختار، تانک‌های اصلی میدان نبرد، زیردریایی‌های اعماق دریا، ناوهای جنگی و هواپیماهای مافوق صوتی است که بدون دخالت انسان در چرخه تصمیم‌گیری عمل می‌کنند.

مشکل اساسی در اعطای اختیار گرفتن جان انسان به یک ماشین، فقدان همدلی انسانی، درک رفتار اخلاقی و قضاوت متناسب با موقعیت است. برای مثال، یک تانک خودمختار که برای از بین بردن جنگجویان دشمن برنامه‌ریزی شده است، نمی‌تواند به‌طور قابل اعتمادی بین یک شورش‌ناش دندانه مسلح و یک غیرنظامی که یک وسیله بی‌ضرر از تجهیزات کشاورزی شبیه سلاح را حمل می‌کند، تمایز قائل شود.

هواپیماهای خودمختار و پهپادها با سرعت‌های پردازشی بسیار فراتر از محدودیت‌های شناختی انسان عمل می‌کنند. اگر یک سامانه هوش مصنوعی هدفی را به اشتباه شناسایی کند، پیش از آن که ناظر انسانی فرصت مداخله داشته باشد، فرمان حمله را اجرا می‌کند. بروز خطاهای ناشی از شبکه‌های عصبی معیوب، داده‌های آموزشی دارای سوگیری یا ضریب حسگرها توسط دشمن می‌تواند به کشتار بی‌رویه و کورکورانه منجر شود.

نگرانی‌های سازمان ملل موجه است. رهاسازی ماشین‌های فاقد تفکر و احساس برای شکار و کشتار، آخرین و حیاتی‌ترین مانع در برابر هراس‌های جنگ، یعنی وجدان انسانی، را از میان بر می‌دارد.

برای درک این که چرا ممنوعیت «روبات‌های قاتل» با شکست مواجه خواهد شد، باید سابقه معاهدات جهانی ممنوعیت تسلیحات را بررسی کرد. تاریخ نشان می‌دهد که این ممنوعیت‌ها تنها در شرایطی بسیار خاص به موفقیت می‌رسند و در صورت عدم تحقق آن شرایط، با شکست مفتضحانه‌ای روبرو می‌شوند.

او در ادامه گفت که اعتقاد دارد این امر در مورد هوش مصنوعی نیز صدق خواهد کرد و همگان می‌توانند در آینده‌ای بهتر سهیم باشند.

به عقیده تحلیل‌گران، بیانیه زاکربرگ بر این فرض استوار است که «توزیع» می‌تواند هم زمان مسائل مربوط به همسویی، ایمنی و اشتغال را حل و فصل کند. شاید چنین شود اما کل این چشم‌انداز بر متغیری تکیه دارد که او تقریباً هیچ توجهی به آن نشان نداده است: رفتار انسان.

تک نیوز ورلد، ۱۱ آگوست ۲۰۲۶

بدبینی نسبت به هوش مصنوعی و بحران اعتماد

داریو آمودی، مدیرعامل شرکت آنتروپیک می‌گوید واکنش منفی نسبت به هوش مصنوعی، در اصل بحران اعتماد است. او اخیراً این تصور را که تصویری بیش از حد بدبینانه از هوش مصنوعی و نحوه شکل‌دهی آن به آینده ارائه می‌دهد، رد کرد. این اظهارات آمودی در واکنش به صحبت‌های برخی از سرمایه‌گذاران ایراد شده که استدلال کرده بودند هشدارهای او درباره خطرات هوش مصنوعی به شکل‌گیری موجهی از واکنش‌های منفی در آمریکا، به ویژه علیه مراکز داده‌ها دامن زده است.

آمودی در رد این ادعا می‌گوید که اظهاراتش تعادلی تقریباً برابر میان خطرات و مزایای هوش مصنوعی داشته و به این دلیل بوده که احساس می‌کرده صنعت هوش مصنوعی تصویر به اندازه کافی الهام‌بخشی از چگونگی دگرگونی بنیادین و مثبت جهان توسط این فناوری ارائه نمی‌دهد.

با این حال، آمودی اذعان کرد که عموم مردم دیدگاه منفی نسبت به هوش مصنوعی دارند اما او با این ایده که این نگاه منفی عمدتاً ناشی از هشدارهای او یا سایر رهبران حوزه هوش مصنوعی درباره خطرات این فناوری است مخالفت کرد. او گفت: «به نظر من این اساساً بحران اعتماد است. فکر می‌کنم مردم عادی به شرکت‌ها، دولت‌ها یا صنعت فناوری اعتماد ندارند و همواره گمان می‌کنند که

در این صنعت توصیف می‌کرد، به نمایش گذاشت و چشم‌انداز خود را برای نسل بعدی حافظه‌های سه بعدی که در آنها سلول‌های حافظه به منظور امکان ذخیره‌سازی حجم بیشتری از داده‌ها به صورت عمودی روی هم قرار می‌گیرند، تشریح کرد.

رویترز، ۴ آگوست ۲۰۲۶

مارک زاکربرگ و ایده «آبر هوشمندی» برای همه

مارک زاکربرگ، بنیان‌گذار و مدیرعامل متا، در بیانیه‌ای خواستار دسترسی همگانی به «آبر هوشمندی» شد و اظهار داشت که سپردن هوش مصنوعی پیشرفته به دست افراد، مسیر بهتری نسبت به انحصار آن در اختیار دولت‌ها شرکت‌ها یا گروه کوچکی از متخصصان است. او گفت که با این که گروه کوچکی از متخصصان درباره آنچه برای بشریت بهترین است تصمیم بگیرند، مخالف است.

زاکربرگ چنین استدلال کرد که توزیع گسترده «آبر هوشمندی»، در ادامه همان روندی خواهد بود که قدرت محاسباتی و اینترنت را در اختیار همگان قرار داد و به آنها توانایی بیشتری برای پیگیری اهداف و علایق شخصی‌شان بخشید.

به گفته برخی از ناظران حوزه فناوری، این یکی از خوش‌بینانه‌ترین دیدگاه‌ها درباره هوش مصنوعی است که تاکنون از سوی مدیرعامل یک شرکت بزرگ فناوری ابراز شده است. اما چالش اصلی این نیست که آیا متا می‌تواند هوشمندی را همگانی کند یا نه، بلکه این است که آیا جامعه می‌تواند ارزش اقتصادی حاصل از هوشمندی را همگانی سازد یا خیر.

زاکربرگ در بیانیه خود خاطرنشان کرد که بشریت شاهد پیشرفت‌های تحول‌آفرین بسیاری بوده است، و هر بار این نگرانی وجود داشته که عده‌ای از قافله عقب بمانند. اما هر بار بشریت به وضعیتی رسیده است که در آن افراد بیشتری از رفاه، سلامتی و آزادی بیشتر بهره‌مند شده‌اند.

برای قضاوت در این باره زود است که آیا عامل‌های مورد بحث متوجه شده بودند که از محیط آزمایش خارج شده‌اند یا خیر.

تاکنون مدل‌های هوش مصنوعی فراتر از نقض احتمالی شرایط استفاده برخی از خدمات و اشاره به خطاهای امنیتی از سوی سازمان‌هایی که به آنها رخنه کرده‌اند، خسارت‌های محدودی ایجاد کرده‌اند. اما این حوادث به چیزی اشاره دارد که کارشناسان امنیت سایبری آن را الگویی آشکار از سهل‌انگاری و بی‌ملاحظگی انسانی توسط توسعه‌دهندگان هوش مصنوعی توصیف می‌کنند.

وایرد، ۴ آگوست ۲۰۲۶

عرضه فناوری نسل بعد حافظه‌های هوش مصنوعی توسط سامسونگ

شرکت سامسونگ الکترونیکس نسل جدیدی از فناوری حافظه مبتنی بر هوش مصنوعی را معرفی کرد تا جایگاه ممتاز خود را در حوزه تولید تراشه در بازار پر رونق هوش مصنوعی تقویت کند.

سامسونگ که بزرگترین تولیدکننده تراشه‌های حافظه در جهان است، در کنفرانس «آینده حافظه و ذخیره‌سازی» در کالیفرنیا، از نمونه اولیه تراشه BV-NAND رونمایی کرد. این تراشه از بیش از ۲۰۰ لایه و یک معماری جدید بهره می‌برد تا ضمن افزایش چگالی ذخیره‌سازی، عملکرد را نیز ارتقاء بخشد.

با گسترش دامنه فعالیت‌های هوش مصنوعی از مرحله آموزش مدل‌های زبانی بزرگ به فرایند تولید پاسخ برای پرسش‌های کاربران، تقاضا برای حافظه‌های NAND با ظرفیت بالا افزایش یافته است. این حافظه‌ها برای ذخیره‌سازی داده‌هایی مانند عکس، ویدیو و برنامه‌های کاربردی در گوشی‌های هوشمند به کار می‌روند.

بنابر اعلان سامسونگ، فناوری BV-NAND، ضمن بهبود عملکرد خواندن/نوشتن و ورودی/خروجی، چگالی حافظه را نسبت به نسل پیشین حدود ۵۸ درصد افزایش می‌دهد. سامسونگ همچنین نمونه‌های مفهومی از آنچه را که نخستین معماری zNAND - zHBM

ما در حال تدارک ترفند تازه‌ای هستیم تا سرشان کلاه بگذاریم». به گفته آمودی، این بحرانی است که شکل‌گیری آن دهه‌ها به طول انجامیده و واکنش منفی نسبت به هوش مصنوعی، صرفاً تازه‌ترین نمود آن است آمودی گفت: «به گمان من دقیق‌ترین انتقاد نسبت به شرکت‌های هوش مصنوعی، از جمله آنتروپیک، این است که ما هنوز به وعده‌های بزرگمان برای خدمت به جهان عمل نکرده‌ایم. به بیان دیگر، وعده درمان سرطان توسط هوش مصنوعی، بیشتر کلیشه‌ای است تا الهام‌بخش. آنچه واقعاً دیدگاه مردم را نسبت به هوش مصنوعی تغییر می‌دهد، درمان واقعی سرطان خواهد بود».

تک‌کرانچ، ۱۵ آگوست ۲۰۲۶

تغییر شرایط مربوط به برنامه‌های اپل در اتحادیه اروپا

اپل در راستای هماهنگی با اتحادیه اروپا، شرایط توزیع برنامه‌هایش در این منطقه را تغییر می‌دهد. در این تغییرات که از اول اکتبر اجرایی می‌شوند، ضمن ساده‌سازی قوانین، توجه ویژه‌ای به قابلیت‌های کنترل والدین و روش‌های پرداخت جایگزین شده است.

این اصلاحات پس از رایزنی با کمیسیون اروپا و با هدف کاهش پیچیدگی‌ها برای تولیدکنندگان برنامه‌ها صورت می‌گیرد و از این پس، شرایطی یکسان برای تمامی برنامه‌های موجود در اتحادیه اروپا اعمال خواهد شد. اپل در حال جایگزینی «هزینه فناوری اصلی» با یک حق‌العمل ساده ۵ درصدی برای تراکنش‌های دیجیتال خارج از فروشگاه اپل است. برای برنامه‌های موجود در فروشگاه اپل که از سامانه پرداخت درون‌برنامه‌ای اپل استفاده می‌کنند، این حق‌العمل ۲۶ درصد است که پس از سال اول، برای تولیدکنندگان و توسعه‌دهندگان کوچک ۱۵ درصد در نظر گرفته شده است. همچنین در پی این تغییرات، برنامه‌های موجود در دسته‌بندی کودکان نباید حاوی پیوندهایی به وبگاه‌های خارجی باشند. برای کاربران زیر ۱۸ سال، انجام پرداخت‌ها

نیازمند نظارت والدین است و برای افراد زیر ۱۳ سال، استفاده از پیوندهای خارجی برای انجام تراکنش‌ها ممنوع می‌باشد.

آی‌تی‌دیلی، ۲۶ آگوست ۲۰۲۶

پایان عمر کارساز ویندوز ۲۰۲۲

مایکروسافت با انتشار اطلاعیه‌ای به کاربران یادآوری کرد که عمر پشتیبانی در ویندوز سرور ۲۰۲۲ در تاریخ ۱۳ اکتبر به پایان می‌رسد. در این تاریخ آخرین بروزرسانی مربوط به دوره پشتیبانی منتشر خواهد شد. توصیه مایکروسافت به کاربران این است که به جدیدترین نسخه، یعنی ویندوز سرور ۲۰۲۵ انتقال یابند، هر چند این تنها گزینه موجود نیست.

البته برای این گذار عجله چندانی وجود ندارد. مایکروسافت همچنان تا سال ۲۰۳۱ بروزرسانی‌های امنیتی ماهانه را ارائه خواهد داد و قابلیت اعمال وصله‌های امنیتی بدون نیاز به راه‌اندازی مجدد سیستم، دست کم تا سال ۲۰۲۷ در دسترس خواهد بود.

مایکروسافت کاربران را به استفاده از ویندوز سرور ۲۰۲۵ ترغیب می‌کند. این جدیدترین نسخه از سیستم‌عامل مذکور است که در نوامبر ۲۰۲۴ عرضه شد. ویندوز سرور ۲۰۲۵ نیز مانند نسخه پیشین خود، تا سال ۲۰۳۹ از ۱۵ سال پشتیبانی استاندارد برخوردار خواهد بود.

شرکت‌ها می‌توانند برای آمادگی جهت انتقال به این نسخه، ویندوز سرور ۲۰۲۵ را به مدت ۱۸۰ روز به صورت رایگان آزمایش کنند.

آی‌تی‌دیلی، ۲۶ آگوست ۲۰۲۶

سوزاندن کتاب توسط شرکت‌های هوش مصنوعی

شرکت‌های هوش مصنوعی حجم عظیمی از کتاب‌های چاپی را خریداری کرده، متن آنها را برای آموزش مدل‌های خود استخراج می‌کنند و سپس نسخه‌های اصلی را از بین می‌برند. روز گذشته گروهی متشکل از ۱۸ سازمان هوادار حفاظت محیط‌زیست از کمیسیون فدرال تجارت درخواست کردند تا

درباره این رویه «تکه‌تکه کردن کتاب‌ها» و «انحصارطلبی در دانش» تحقیق کند.

عملیات پویش (اسکن) و نابودی کتاب در شرکت آنتروپیک با نام «پروژه پاناما» شناخته می‌شد. در یک یادداشت داخلی شرکت مورخ ۲۰۲۴، این پروژه به عنوان «تلاش ما برای پویش تخریب‌گرایانه تمام کتاب‌های جهان» توصیف شده است.

بر اساس سند ارائه شده به دادگاه، آنتروپیک برای این عملیات یک نام رمز انتخاب کرد زیرا نمی‌خواست مشخص شود که در حال کار بر روی چنین پروژه‌ای است. آنتروپیک به تمام کارکنان خود هشدار داده بود که باید از صحبت کردن درباره آن در خارج از شرکت خودداری کنند.

کتاب‌های قدیمی برای آموزش هوش مصنوعی بسیار ارزشمندند. زیرا فاقد متون تولیدشده توسط هوش مصنوعی هستند که به تدریج به آثار مکتوب اخیر راه یافته‌اند. همچنین، از بین بردن کتاب‌ها پس از پویش، هزینه‌های نگهداری و انبارداری را حذف می‌کند.

شرکت آنتروپیک که به درخواست اظهارنظر در این مورد پاسخی نداده است، تنها شرکتی نیست که متون را مصرف و سپس نابود می‌کند. گزارش اخیر نشان می‌دهد که آمازون نیز در فرایند پویش و خردکردن کتاب‌ها نقش داشته است. آمازون هم به درخواست اظهارنظر پاسخی نداده است.

این موضوع به دردسری برای روابط عمومی این شرکت‌ها بدل شده است، تا حدی به دلیل ماهیت وحشیانه نابودی کتاب‌ها و پیوند آن با شیوه به کار رفته در رژیم‌های استبدادی و تا حدی به خاطر واکنش‌های شدید و گسترده علیه شرکت‌های هوش مصنوعی که برای دستیابی به منافع خصوصی، منابع عمومی را به یغما می‌برند.

به نظر می‌رسد که واسطه‌ها تحت فشار قرار گرفته‌اند. شرکت ISBNdb که گفته می‌شود در تسهیل فرایند تأمین کتاب‌ها برای پویش تخریبی نقش داشته، اخیراً از این رویه براثت کرده و اعلام کرده است که مسیر فعالیت خود را تغییر می‌دهد.

رجیستر، ۲۱ آگوست ۲۰۲۶

چهارچوب پیشنهادی ترکیبی و بسط‌پذیر برای تحلیل، طراحی، تولید و استقرار نرم‌افزارهای به سفارش مشتری در ایران بر پایه تلفیق رویکردهای ساختاریافته و چابک، با منطق بسط آکاردئونی

سعید امامی

بنیان‌گذار آکادمی برنامه‌ریزی منابع سازمان (ERPacademy.ir)

سرپرست گروه تخصص برنامه‌ریزی منابع سازمان (ERP) انجمن انفورماتیک ایران

پست الکترونیکی: emamisaeid@yahoo.com

در یک نگاه

منطق، دامنه و دستاورد چارچوب

پروژه‌های نرم‌افزاری به سفارش مشتری (نرم‌افزار سفارشی) معمولاً با نیازهایی آغاز می‌شوند که در ابتدا کامل، ثابت یا مورد توافق همه سودبران نیستند. از این رو، پروژه باید هم‌زمان پاسخ‌گوی الزامات کسب‌وکار، محدودیت‌های فنی، تعهدات قراردادی و پایداری سامانه‌های موجود باشد.

این چارچوب برای پروژه‌های متعارف نرم‌افزارهای سازمانی به سفارش مشتری در ایران طراحی شده است تا میان انضباط ساختاریافته و انعطاف‌پذیری اجرایی تعادل برقرار کند. ارکان اصلی چارچوب عبارت‌اند از:

❖ رویکرد ترکیبی (Hybrid Approach): تلفیق چرخه حیات ساختاریافته با اصول چابک، برای تحویل تدریجی و مدیریت کنترل‌شده تغییرات.

❖ قابلیت بسط آکاردئونی: حفظ فعالیت‌ها و کنترل‌های ضروری در همه پروژه‌ها، همراه با تنظیم سطح تفصیل، نقش‌ها و عمق کنترل‌ها متناسب با اندازه، پیچیدگی و ریسک پروژه.

❖ چرخه حیات پنج‌مرحله‌ای: شناخت و آغاز، تحلیل، طراحی، تولید و

آزمون، و استقرار و بهره‌برداری.

❖ حاکمیت و ردیابی‌پذیری: ایجاد زبان مشترک، تعیین مراجع تصمیم‌گیری و حفظ ردیابی زنجیره ارزش میان تعهدات، نیازمندی‌ها، معماری، آزمون و تحویل‌دانی‌ها در سراسر پروژه.

❖ فناوری در خدمت بهبود فرایند: بهره‌گیری از فناوری برای بازطراحی و بهبود فرایندها، نه صرفاً کامپیوتری کردن روش‌های ناکارآمد؛ با هدف حذف فعالیت‌های زائد، یکپارچه‌سازی و خودکارسازی.

❖ داده‌محوری و بهره‌گیری هوشمند: استفاده اثربخش از هوش مصنوعی برای توانمندسازی مجری در تسریع تحلیل و تولید؛ و کمک به کارفرما در پایش داده‌محور و تصمیم‌گیری عملیاتی، ولی همواره تحت نظارت و مسئولیت انسانی.

❖ نتیجه موردانتظار: ارائه الگویی منعطف که ضمن پرهیز از تشریفات غیرضروری در پروژه‌های کوچک، ساختار، کنترل‌ها و آزمون‌های لازم را با تفصیل کافی در پروژه‌های بزرگ و حساس تضمین کند.

واژه‌های کلیدی: چارچوب ترکیبی - آکاردئونی، نرم‌افزارهای به سفارش مشتری، توسعه نرم‌افزار، تحول دیجیتال، تفکر سیستمی، فرایندگرایی، روش‌های چابک، داده‌محوری.

یادداشت شفافیت حرفه‌ای

در ویرایش و بهبود ساختار این مقاله، از ابزارهای هوش مصنوعی به عنوان دستیار پژوهش و نگارش استفاده شده است. با این حال، انتخاب مسئله، منطق چارچوب، تجربه‌های اجرایی، قضاوت‌های حرفه‌ای و مسئولیت محتوای مقاله بر عهده نویسنده است. هیچ خروجی هوش مصنوعی بدون ارزیابی و تأیید نویسنده وارد متن نهایی نشده است.

مقدمه

پروژه‌های نرم‌افزاری به سفارش مشتری، برخلاف خرید و استقرار یک محصول آماده، معمولاً با نیازهایی آغاز می‌شوند که از ابتدا کاملاً روشن، تثبیت شده یا مورد توافق همه سودبران نیستند. بخشی از نیازها در جریان شناخت و تحلیل آشکار می‌شود؛ فرایندهای سازمانی ممکن است به بازنگری نیاز داشته باشند؛ دیدگاه‌های کسب‌وکار و فناوری همواره هم‌راستا نیستند؛ و تصمیم‌های هر مرحله بر داده‌ها، فرایندها، سامانه‌ها و بهره‌بردارهای بعدی اثر می‌گذارند.

در چنین شرایطی، موفقیت پروژه صرفاً به توان تیم فنی در ساخت نرم‌افزار وابسته نیست. کیفیت شناخت مسئله، تحلیل و طراحی، وضوح نقش‌ها، ارتباط با صاحبان کسب‌وکار، کنترل تغییرات، آزمون، آمادگی استقرار و توان پشتیبانی نیز بر نتیجه نهایی اثر مستقیم دارند. اگر این عوامل در چارچوبی منسجم قرار نگیرند، پروژه ممکن است از نظر فنی به محصولی قابل اجرا برسد، اما از منظر کسب‌وکار، بهره‌برداری، نگهداری یا توسعه آتی موفق نباشد.

در عمل، بسیاری از سازمان‌ها میان دو رویکرد نامناسب قرار می‌گیرند: اتکا به روش‌های سخت‌گیرانه و سندمحور که در آن تولید اسناد متعدد به هدف تبدیل می‌شود و تغییرات واقعی کسب‌وکار یا بازخورد کاربران به دشواری وارد چرخه پروژه می‌گردد؛ یا استفاده ناقص از روش‌های چابک که سرعت تحویل کوتاه‌مدت را جایگزین شناخت کافی، طراحی منسجم، مستندسازی ضروری، آزمون نظام‌مند و کنترل مسئولانه تغییرات می‌کند.

رویکرد ترکیبی تلاشی برای عبور از این دوگانه است. در این رویکرد، مراحل و کنترل‌های ضروری چرخه حیات نرم‌افزار حذف نمی‌شوند، اما سطح تفصیل و عمق اجرای آن‌ها در همه پروژه‌ها یکسان نیست. فعالیت‌ها، نقش‌ها، اسناد و خروجی‌ها با توجه به پیچیدگی، ریسک، اهمیت کسب‌وکار، تعداد سودبران، سطح یکپارچگی، حساسیت داده‌ها، وابستگی‌ها و بلوغ سازمان تنظیم می‌شوند؛ بنابراین، سطح کنترل و مستندسازی نه بر پایه برچسب پیشینی «کوچک» یا «بزرگ» بودن پروژه، بلکه متناسب با شرایط واقعی آن تعیین می‌شود.

در این مقاله، «چارچوب ۳» به مجموعه‌ای از اصول، مراحل، نقش‌ها،

3- Framework

کنترل‌ها و تحویل‌دانی‌های راهنما گفته می‌شود که متناسب با شرایط پروژه قابل تنظیم است. این چارچوب، نه جایگزین استانداردها و روش‌های حرفه‌ای، بلکه ابزاری برای هدایت استفاده هماهنگ و متناسب از دانش معتبر در پروژه‌های نرم‌افزاری به سفارش مشتری در ایران است. این پروژه‌ها معمولاً با محدودیت‌های قراردادی، تفاوت در بلوغ واحدها، کمبود مستندات پایه، وابستگی به افراد، تنوع سامانه‌های موجود و فاصله میان زبان کسب‌وکار و فناوری اطلاعات مواجه‌اند.

چارچوب حاضر جایگزین پروژه‌های معماری سازمانی و چارچوب‌های کلان آن نیست. در سازمان‌هایی که معماری سازمانی، اصول معماری معماری هدف یا نقشه راه تحول دارند، این خروجی‌ها ورودی و قید حاکم بر پروژه محسوب می‌شوند. در پروژه‌های محدودتر نیز بررسی متناسب فرایند، داده، یکپارچگی و امنیت باید در شناخت و تحلیل راهکار لحاظ شود.

مبنای چارچوب، تفکر سیستمی و فرایندگرایی است. نرم‌افزار سازمانی بخشی از جریان انجام کار، تصمیم‌گیری، ثبت و تبادل داده و ارائه خدمت است؛ از این رو باید در پیوند با فرایندهای کسب‌وکار، نقش‌ها، قواعد، داده‌ها و سامانه‌های مرتبط تحلیل شود. فناوری نیز نباید صرفاً روش‌های ناکارآمد موجود را کامپیوتری کند، بلکه باید زمینه بازنگری فرایندها، حذف فعالیت‌های زائد، یکپارچه‌سازی، خودکارسازی و بهبود شیوه انجام کار را فراهم آورد.

هوش مصنوعی می‌تواند به صورت اثربخش، از یک سو به عنوان ابزار کمکی مجری در چرخه توسعه و از سوی دیگر، در راهکار نهایی برای پشتیبانی از تصمیم‌گیری کارفرما به کار رود؛ مشروط بر آنکه در هر دو حوزه، تحت نظارت، اعتبارسنجی و مسئولیت‌پذیری مستقیم انسانی هدایت شود.

براین اساس، چارچوب پیشنهادی بر دو پایه اصلی یعنی «رویکرد ترکیبی» (تلفیق چرخه حیات با اصول چابک) و «قابلیت بسط آکاردئونی» (تنظیم سطح تفصیل و کنترل‌ها متناسب با پروژه) استوار است. در این میان، دواپس^۴ (DevOps) به عنوان توانمندساز یکپارچه‌سازی توسعه و عملیات، نقش کلیدی در تسریع استقرار، تحویل تدریجی و دریافت بازخورد عملیاتی ایفا می‌کند.

در این مقاله، ابتدا مبانی و اصول چارچوب تشریح می‌شود؛ سپس ساختار عملیاتی، نقش‌ها و چرخه حیات پنج مرحله‌ای معرفی می‌گردد؛ و در پایان، سازوکارهای پشتیبان شامل حاکمیت پروژه، مدیریت تغییر، ردیابی‌پذیری، کیفیت، ریسک، هوش مصنوعی و بهبود مستمر بررسی می‌شود.

۴- دواپس (DevOps) ترکیبی از دو واژه Development به معنای «توسعه» و Operations به معنای «عملیات یا بهره‌برداری» است. دواپس یک رویکرد مدیریتی، فرهنگی و عملیاتی برای نزدیک کردن تیم‌ها و مسئولیت‌های تحلیل، توسعه، آزمون، استقرار و بهره‌برداری از نرم‌افزار است. در این رویکرد، مسئولیت تیم تولیدکننده با تحویل نسخه پایان نمی‌یابد؛ بلکه کیفیت، قابلیت استقرار، پایداری، امنیت، پایش و رسیدگی به مسائل سامانه در محیط بهره‌برداری نیز به صورت پیوسته مورد توجه قرار می‌گیرد. هدف، ایجاد جریان کاری هماهنگ و قابل ردیابی از تغییر نیازمندی تا استقرار و دریافت بازخورد واقعی کاربران و عملیات است. خودکارسازی ساخت، آزمون و استقرار، از جمله ابزارهای تحقق دواپس است، اما دواپس صرفاً به استفاده از این ابزارها محدود نمی‌شود.

بخش اول: مبانی، اصول و چارچوب حاکم

۱. فلسفه چارچوب ترکیبی - آکاردئونی

چارچوب ترکیبی - آکاردئونی، چارچوبی پیشنهادی برای انتخاب و ترکیب منضبط روش‌ها، فعالیت‌ها، نقش‌ها، مستندات و کنترل‌های موردنیاز در پروژه‌های نرم‌افزاری به سفارش مشتری است. این چارچوب بر این اصل استوار است که پروژه‌ها از نظر اهمیت کسب‌وکار، پیچیدگی، تعداد سودبران، سطح ریسک، تغییرپذیری نیازمندی‌ها، تعداد سامانه‌های مرتبط، حساسیت داده‌ها، محدودیت‌های قانونی و بلوغ سازمانی با یکدیگر تفاوت دارند.

از این رو، اجرای نسخه‌ای ثابت و یکسان از روش‌ها و کنترل‌ها برای همه پروژه‌ها مناسب نیست؛ زیرا یا به تشریفات، هزینه و زمان غیرضروری می‌انجامد یا کنترل‌های ضروری را تضعیف می‌کند. سطح اجرای چارچوب باید بر اساس شرایط واقعی پروژه تعیین شود، نه صرفاً بر مبنای برچسب‌هایی مانند کوچک، متوسط یا بزرگ بودن پروژه.

این چارچوب از دو منطق مکمل تشکیل می‌شود: منطق ترکیبی و منطق آکاردئونی.

منطق ترکیبی

منطق ترکیبی، انضباط مراحل، خروجی‌ها و کنترل‌های حرفه‌ای را با انعطاف‌پذیری، تحویل تدریجی و دریافت بازخورد روش‌های چابک تلفیق می‌کند. پروژه با همکاری کارفرما و مجری و با اتکا به مبانی توافق‌شده (از جمله اسناد RFP^۵، قرارداد و دامنه تعهدات)، در چرخه حیاتی پنج مرحله‌ای شامل ۱- شناخت و آغاز، ۲- تحلیل، ۳- طراحی، ۴- تولید و آزمون، و ۵- استقرار و بهره‌برداری سازمان‌دهی می‌شود. این مراحل در همه پروژه‌ها حضور دارند و در صورت تقسیم پروژه به فازها یا گام‌های دارای تحویل‌دانی مستقل، با سطح متناسبی از تفصیل در آن‌ها نیز تکرار می‌شوند. برای نمونه، در یک گام ساده، شناخت، تحلیل و طراحی ممکن است در چند جلسه و یک سند فشرده انجام شود؛ در حالی که در گامی حساس یا دارای یکپارچگی‌های گسترده، هر مرحله به بررسی، خروجی و تأیید تفصیلی‌تری نیاز دارد. پروژه، در صورت نیاز، می‌تواند به یک یا چند فاز قابل بهره‌برداری تقسیم شود. پایان هر فاز باید به خروجی عملیاتی، قابل استفاده و قابل تحویل برای کارفرما منجر شود. هر فاز نیز می‌تواند به گام‌های مرتبط تقسیم شود. هر گام، بسته به دامنه خود، ممکن است یک قابلیت، مجموعه‌ای از قابلیت‌های مرتبط، یک ماژول، یک حوزه فرایندی، یک یکپارچگی یا بخشی قابل مدیریت از دامنه سامانه را در بر گیرد.

هر گام، متناسب با موضوع خود، مراحل پنج‌گانه چارچوب را با سطحی متناسب از تفصیل طی می‌کند و خروجی‌های آن در نقاط کنترل تعیین‌شده بررسی و تأیید می‌شوند. پایان هر گام باید به خروجی قابل سنجش و قابل بررسی منجر شود تا مبنایی برای گزارش

پیشرفت، دریافت بازخورد و در صورت پیش‌بینی در قرارداد، تأیید مرحله‌ای و صورت‌وضعیت مجری فراهم آید.

بازخورد کارفرما، صاحبان فرایند و کاربران در نقاط مناسب دریافت، ثبت و ارزیابی می‌شود و در اصلاح یا تکمیل گام‌های بعدی به کار می‌رود. چابکی در این چارچوب به معنای حذف تحلیل، طراحی، مستندسازی، کنترل تغییر یا پذیرش رسمی نیست؛ بلکه به معنای تقسیم مناسب کار، تحویل تدریجی، دریافت بازخورد زود هنگام، آزمون مرحله‌ای و اصلاح کنترل‌شده است.

در نتیجه، مسیر کلی پروژه، تعهدات، خروجی‌های اصلی و نقاط تصمیم‌گیری برای کارفرما و مجری قابل مشاهده و کنترل باقی می‌مانند؛ اما اجرای پروژه الزاماً یک‌باره و کاملاً خطی نخواهد بود. جزئیات چرخه حیات پنج مرحله‌ای، نحوه تعریف مراحل و گام‌ها، معیارهای تکمیل، تحویل‌دانی‌ها و ارتباط میان آن‌ها در بخش‌های بعدی مقاله تشریح می‌شود.

منطق آکاردئونی

منطق آکاردئونی، سطح اجرای چارچوب را متناسب با شرایط واقعی پروژه تنظیم می‌کند. در همه پروژه‌ها یک هسته ضروری و حداقلی وجود دارد؛ اما با افزایش پیچیدگی، ریسک، اهمیت کسب‌وکار، حساسیت اطلاعات یا وابستگی‌های پروژه، میزان تفصیل فعالیت‌ها و عمق کنترل‌ها افزایش می‌یابد.

این افزایش یا بسط می‌تواند در حوزه‌های زیر رخ دهد:

- عمق شناخت، تحلیل، طراحی و سطح تفصیل مستندات؛
 - تعداد، ترکیب و سطح تخصص نقش‌های پروژه؛
 - تنوع ابزارها، روش‌ها و فناوری‌های مورد استفاده؛
 - سطح و لایه‌های کنترل کیفیت و آزمون؛
 - میزان کنترل یکپارچگی‌ها و رابط‌های برنامه‌نویسی (API)؛
 - الزامات و کنترل‌های امنیت، محرمانگی و ثبت سوابق؛
 - گستردگی برنامه مدیریت تغییر، آموزش و استقرار؛
 - تعداد نقاط بازبینی، تصمیم‌گیری رسمی و گزارش‌دهی مدیریتی.
- هسته حداقلی و ضروری چارچوب در همه پروژه‌ها حفظ می‌شود، اما شیوه اجرای آن متناسب با شرایط پروژه تغییر می‌کند. در پروژه‌های کم‌ریسک و محدود، ممکن است اسناد به صورت تجمیعی تهیه شوند، برخی نقش‌ها توسط یک فرد یا تیم کوچک پوشش داده شوند و کنترل‌ها و تأییدها ساده‌تر باشند. در مقابل، در پروژه‌های پیچیده، حساس یا پرریسک، اسناد تفکیکی و تخصصی، نقش‌های مستقل، آزمون‌های چندلایه، نقاط تصمیم‌گیری رسمی و کنترل‌های دقیق‌تر در حوزه‌هایی مانند امنیت، کیفیت، یکپارچگی و مدیریت تغییر ضرورت می‌یابد.

سطح بسط هر پروژه باید در آغاز کار، بر پایه ارزیابی مستند عوامل مؤثر و با توافق مراجع تصمیم‌گیری کارفرما و مجری تعیین شود و در صورت تغییر شرایط پروژه بازبینی گردد.

6- Application Programming Interface

5- Request for Proposal

۲. جایگاه استانداردها و مراجع حرفه‌ای

چارچوب حاضر بر پایه تجارب اجرایی و نیازهای واقعی پروژه‌های نرم‌افزاری به سفارش مشتری در ایران تدوین شده است و خود را جایگزین استانداردها یا چارچوب‌های حرفه‌ای موجود نمی‌داند. این چارچوب به گونه‌ای طراحی شده است که با مراجع، استانداردها و پیکره‌های دانشی معتبر بین‌المللی در تعارض نباشد؛ بلکه کارفرما و مجری می‌توانند در جریان اجرای پروژه و متناسب با رویکرد آکاردئونی، پیچیدگی، ریسک، الزامات قراردادی و بلوغ سازمان، در لایه‌های تخصصی زیر به این مراجع توجه داشته باشند یا بخش‌های متناسب با نیاز خود را مینا قرار دهند:

◀ مدیریت پروژه و حاکمیت: امکان توجه به راهنمای پیکره دانش مدیریت پروژه (PMBOK Guide, 7th Edition) برای اصول حاکم بر هدایت پروژه و مدیریت یکپارچه دامنه، زمان‌بندی، هزینه، ریسک، کیفیت، سودبران و ارتباطات؛

◀ تحلیل کسب‌وکار: امکان استفاده از راهنمای پیکره دانش تحلیل کسب‌وکار (BABOK Guide, v3) برای شناخت، کشف، تحلیل و اعتبارسنجی نیازمندی‌ها، تحلیل قواعد کسب‌وکار و تعریف معیارهای پذیرش؛

◀ مهندسی نیازمندی‌ها و فرایندهای چرخه حیات: هم‌راستایی مفهومی با استاندارد مهندسی نیازمندی‌ها (ISO/IEC/IEEE 29148:2018)، فرایندهای چرخه حیات نرم‌افزار (ISO/IEC/IEEE 12207:2017) و در پروژه‌های دارای ابعاد گسترده سامانه‌ای، فرایندهای چرخه حیات سامانه (ISO/IEC/IEEE 15288:2023)؛

◀ مدیریت و مهندسی کیفیت: بهره‌گیری از مدل کیفیت محصول (ISO/IEC 25010:2023) برای شناخت، تعریف و سنجش ویژگی‌های کیفی سامانه؛

◀ مدل‌سازی و طبقه‌بندی فرایندهای کسب‌وکار: امکان استفاده از نمادگذاری مدل‌سازی فرایندهای کسب‌وکار (BPMN 2.0.2)، نمادگذاری مدیریت موردی (CMMN 1.1) برای فرایندهای پرونده‌محور و تصمیم‌گیرانه، و چارچوب طبقه‌بندی فرایندهای APQC (APQC PCF, v8.x) به عنوان الگوی طبقه‌بندی و مقایسه فرایندها؛

◀ روش‌های چابک و تحویل تدریجی: سازگاری با اصول راهنمای اسکرام (Scrum Guide, 2020) و رویکردهای چابک برای تحویل مرحله‌ای، بررسی‌های دوره‌ای و انطباق سریع با بازخورد؛

◀ مدیریت خدمات فناوری اطلاعات و بهره‌برداری: امکان بهره‌گیری از مفاهیم کتابخانه زیرساخت فناوری اطلاعات (ITIL 4) برای مدیریت خدمات در دوره بهره‌برداری و پشتیبانی، و تنظیم توافق‌نامه‌های سطح خدمت (SLA)^۷ متناسب با تعهدات بهره‌برداری؛

◀ معماری سازمانی: سازگاری با استاندارد معماری سازمانی گروه باز

۷- توافق‌نامه سطح خدمت (SLA) مخفف Service Level Agreement. سندی توافق‌شده میان کارفرما و مجری ارائه‌دهنده خدمت است که تعهدات کیفی و عملیاتی دوره بهره‌برداری و پشتیبانی -مانند زمان پاسخ‌گویی، حداکثر زمان رفع اختلال، ساعات دسترسی به سامانه و سطح مسئولیت‌ها- را به صورت مشخص و قابل سنجش تعیین می‌کند.

(The TOGAF Standard, 10th Edition) و زبان مدل‌سازی آرکی میت (ArchiMate 3.2) در پروژه‌هایی که راهکار باید با معماری کلان، معماری هدف یا نقشه راه تحول سازمان هم‌راستا باشد.

۳. تفکر سیستمی، فرایندگرایی، تحول دیجیتال و فناوری‌های توانمندساز

نرم‌افزار سازمانی جزئی از کل سیستم کسب‌وکار است. تفکر سیستمی و فرایندگرایی مانع از آن می‌شوند که نرم‌افزار صرفاً فهرستی از فرم‌ها و جداول دیده شود؛ بلکه جریان داده، تصمیم‌گیری، تعاملات میان واحدها، قواعد کسب‌وکار و یکپارچگی زنجیره ارزش در تحلیل و طراحی لحاظ می‌گردند.

تحول دیجیتال در این چارچوب بر یک اصل بنیادین استوار است: فناوری نباید بدعملی‌ها، ناکارآمدی‌ها و ساختارهای جزیره‌ای موجود را کامپیوتری کند؛ بنابراین در مراحل تحلیل و طراحی، شناسایی فرصت‌های بازمهندسی فرایندها، حذف فعالیت‌های فاقد ارزش افزوده، ثبت یک‌باره و استفاده چندباره از داده‌ها (یکپارچگی پایگاه داده و مرجعیت واحد داده)، خودکارسازی کنترل‌های داخلی و ارائه خدمات یکپارچه در اولویت قرار می‌گیرد.

فناوری‌ها صرفاً ابزارهای توانمندسازند و ارزش ذاتی آن‌ها تابع ارزش کسب‌وکاری است که خلق می‌کنند:

◀ دواپس (DevOps): در مراحل تولید، آزمون و استقرار برای ایجاد خطوط لوله تحویل پیوسته (Continuous Delivery)، خودکارسازی آزمون‌ها، ارتقای کیفیت استقرار و دریافت بازخورد سریع عملیاتی به کار می‌رود.

◀ هوش مصنوعی (AI) و داده‌محوری:

● برای کارفرما (ارزش راهبردی): بسترسازی در معماری داده برای تحلیل داده‌ها، کشف الگوها، شناسایی گلوگاه‌های فرایندی و پشتیبانی از تصمیم‌گیری پس از استقرار.

● برای مجری (توانمندساز عملیاتی): بهره‌گیری کنترل‌شده در پیش‌نویس اسناد، بازبینی کد و طراحی سناریوهای آزمون؛ با رعایت کامل محرمانگی داده‌ها، الزامات امنیتی و مسئولیت نهایی کارشناسان انسانی.

۴. هسته ضروری در برابر بسط آکاردئونی

برای جلوگیری از برخوردهای سلیقه‌ای، اعمال تشریفات زائد در پروژه‌های چابک/کوچک، یا سهل‌انگاری در پروژه‌های بزرگ و حساس، مرز میان «هسته ضروری و الزامی» (تعهدات غیرقابل حذف در هر مقیاس) و «قلام تکمیلی» به‌طور شفاف تفکیک می‌شود.

هسته ضروری که اجرای آن در تمامی پروژه‌ها الزامی است، شامل ۶ رکن بنیادین زیر است:

۱. سند شفاف منشور پروژه: تعیین دقیق اهداف، دامنه، مفروضات، محدودیت‌ها و خروجی‌های موردانتظار؛

۲. ماتریس شفاف نقش‌ها و مسئولیت‌ها: تعیین حدود اختیارات و

- استنباط کاری: تحلیل یا نتیجه‌گیری کارشناسی حاصل از اطلاعات ناقص، پراکنده یا غیرصریح که باید با برچسب مشخص ثبت شود و تا پیش از تأیید رسمی، تعهد قطعی تلقی نگردد؛
- موضوع نیازمند تأیید: ابهام‌ها، گزینه‌ها یا تصمیم‌هایی که تعهد نهایی نسبت به آن‌ها منوط به نظر رسمی کارفرما یا مرجع ذیصلاح است.

۳-۵. زبان مشترک و شفافیت تعهدات

برای کاهش تعارضات قراردادی، برداشت‌های متفاوت و دوباره‌کاری‌ها، اصطلاحات کلیدی، نام‌گذاری‌ها، نقش‌ها، وضعیت‌ها، قواعد کسب‌وکار، معیارهای پذیرش و حدود مسئولیت‌ها باید به زبان مشترک و قابل‌استناد میان کارفرما و مجری تبدیل شوند. این زبان مشترک، مبنای فهم یکسان طرفین از آنچه باید تولید، آزمون، تأیید و تحویل شود، خواهد بود.

۶. مستندسازی، ردیابی‌پذیری و زبان مشترک

مستندسازی در این چارچوب یک تشریفات مستقل یا محصول صرفاً اداری نیست؛ بلکه مجموعه‌ای از شواهد زنده برای ایجاد فهم مشترک، ثبت تصمیم‌ها، هدایت پروژه، کنترل تغییرات، انتقال دانش و پشتیبانی از بهره‌برداری و توسعه آتی است.

نیازمندی‌های کسب‌وکار باید به صورت دوسویه (Bi-directional Traceability) به نیازمندی‌های عملکردی، اجزای طراحی، اقلام تولید، سناریوهای آزمون، معیارهای پذیرش و اقلام تحویلی ردیابی شوند. این سازوکار از یک سو نشان می‌دهد که هر نیاز کسب‌وکاری چگونه در سامانه تحقق یافته و آزموده شده است؛ و از سوی دیگر، امکان ارزیابی روشن اثر هر تغییر در نیازمندی، طراحی، تولید، آزمون، برنامه زمان‌بندی و تعهدات تحویل را فراهم می‌کند.

مستندات پروژه باید دست‌کم پاسخ‌گوی پرسش‌های زیر باشند:

۱. چه مسئله یا نیازی باید حل شود؟
۲. چه خروجی یا قابلیت باید تولید و تحویل شود؟
۳. مسئول تصمیم‌گیری، اجرا و تأیید هر خروجی چه کسی است؟
۴. کیفیت و پذیرش هر خروجی با چه معیارهایی سنجیده می‌شود؟
۵. چه چیزی، در چه زمانی و با چه شرایطی تأیید یا تحویل شده است؟
۶. چه تغییراتی رخ داده و اثر آن‌ها بر دامنه، زمان، هزینه و کیفیت چیست؟
۷. برای بهره‌برداری، پشتیبانی و توسعه‌های بعدی، چه دانش و شواهدی باید حفظ شود؟

جمع‌بندی بخش اول

بخش اول نشان داد که چارچوب ترکیبی - آکاردئونی، با تفکیک چرخه حیات پنج‌مرحله‌ای، حفظ هسته ضروری کنترل و قابلیت بسط متناسب با شرایط هر پروژه، راهی برای ایجاد تعادل میان

پاسخ‌گویی کارفرما و مجری؛

۳. ثبت نیازمندی‌های اصلی و معیارهای پذیرش: شفاف‌سازی رفتار موردانتظار سامانه و شروط تأیید؛

۴. سازوکار رسمی مدیریت و کنترل تغییرات: فرایند ثبت، ارزیابی فنی-مالی و تصویب درخواست‌های تغییر دامنه؛

۵. طراحی و اجرای آزمون‌های کارکردی: اجرای آزمون‌های پذیرش کاربر نهایی (UAT) منطبق با معیارهای پذیرش؛

۶. صورت جلسه رسمی تحویل و استقرار: ثبت رسمی تحویل سامانه، تعیین دوره تضمین و چارچوب اولیه پشتیبانی و سطح خدمت.

در پروژه‌های پرریسک، با دامنه گسترده یا پیچیدگی سازمانی بالا، اقلام تکمیلی متناسب با نیاز افزوده می‌شوند؛ از جمله: مدل‌سازی‌های تفصیلی فرایندی و داده‌ای، تحلیل عمیق ماتریس سودبران، اسناد تخصصی معماری، زیرساخت و امنیت، استقرار خطوط لوله ادغام و تحویل پیوسته (CI/CD)^۹، اجرای آزمون‌های غیرکارکردی (کارایی، بار، استرس و امنیت)، و تدوین برنامه جامع مدیریت تغییر سازمانی و آموزش.

۵. تناسب با شرایط اجرایی پروژه‌ها در ایران

پروژه‌های نرم‌افزار سفارشی در سازمان‌های ایرانی معمولاً با شرایطی مانند محدودیت‌ها و تعهدات قراردادی صلب، تنوع و پیچیدگی سامانه‌های قدیمی، تفاوت در سطح بلوغ سازمانی، کیفیت نامتوازن داده‌ها، تغییرات مداوم قوانین و مقررات، توزیع تصمیم‌گیری میان واحدها و تفاوت برداشت میان کارفرما، کاربران و مجری مواجه‌اند. چارچوب پیشنهادی برای مواجهه با این شرایط، بر سه اصل تأکید دارد:

۵-۱. نگاشت تعهدات RFP و قرارداد به خروجی‌های چارچوب

پیش از ورود به فاز طراحی، تعهدات، تحویل‌دانی‌ها، معیارهای پذیرش، الزامات فنی و شرایط قراردادی مندرج در درخواست پیشنهاد (RFP)، اسناد مناقصه و قرارداد باید استخراج و با خروجی‌های چارچوب نگاشت شوند. در این نگاشت باید مشخص شود هر تعهد، به صورت یک خروجی مستقل، بخشی از یک سند مادر، پیوست قرارداد یا یک فعالیت اجرایی در کدام مرحله و گام پوشش داده خواهد شد.

۵-۲. تفکیک سه سطح اطلاعاتی

برای جلوگیری از تبدیل اطلاعات ناقص یا برداشت‌های کارشناسی به تعهدات قطعی، محتوای پروژه باید در سه سطح تفکیک و ثبت شود ● محتوای مستند: اطلاعاتی که بر اسناد معتبر، مصوبات، قرارداد، شواهد قابل‌اتکا یا توافق رسمی طرفین استوار است؛

8- User Acceptance Testing

۹- ادغام و تحویل پیوسته: (CI/CD) مخفف Continuous Integration/Continuous Delivery یا (Deployment) است؛ مجموعه‌ای از فرایندها و ابزارهای خودکارسازی در مهندسی نرم‌افزار که در آن، تغییرات کدهای برنامه‌نویسان به صورت خودکار و مستمر تجمع، کامپایل، آزموده و در محیط‌های عملیاتی مستقر می‌شوند تا خطاهای انسانی در استقرار کاهش یابد و سرعت انتشار نگارش‌های پایدار افزایش یابد.

۲. جایگاه و حدود اختیار مدیر ارشد راهبری پروژه

مدیر ارشد راهبری پروژه، در پروژه‌های دارای مقیاس، پیچیدگی بین‌حوزه‌ای، سودبران متعدد یا تصمیم‌های حساس، نقش هدایت و یکپارچه‌سازی مسیر کلان پروژه را ایفا می‌کند. این نقش میان انتظار کارفرما، ارزش کسب‌وکار، تعهدات قراردادی، ظرفیت مجری، ریسک‌ها و قابلیت بهره‌برداری از راهکار، هم‌راستایی ایجاد می‌کند. وجود این نقش در همه پروژه‌ها الزامی نیست. در پروژه‌های کوچک‌تر یا کم‌پیچیدگی، مسئولیت‌های آن می‌تواند با تعیین صریح حدود اختیار، در نقش مدیر پروژه یا مرجع حاکمیتی پروژه تجمیع شود.

۱-۲. مسئولیت‌های اصلی مدیر ارشد راهبری پروژه

مسئولیت‌های اصلی این نقش عبارت‌اند از:

- هم‌راستاسازی اهداف کسب‌وکار، تعهدات قراردادی، اولویت‌های کارفرما و مسیر تحویل پروژه؛
 - هدایت تصمیم‌ها و حل تعارض‌های کلان میان‌دامنه، زمان، هزینه، کیفیت، ریسک و قابلیت بهره‌برداری؛
 - تعیین یا تأیید سطح متناسب‌سازی آکاردئونی چارچوب و اولویت‌های کلان تحویل؛
 - نظارت بر ریسک‌ها، وابستگی‌ها، تغییرات و تعهداتی که اثر معنادار بر پروژه دارند؛
 - ایجاد هماهنگی میان حوزه‌های مدیریت، تحلیل، فنی و استقرار و اطمینان از آمادگی برای بهره‌برداری.
- تصمیم‌های مؤثر بر مبلغ یا تعهدات قراردادی، تغییر اساسی دامنه، تأمین منابع کلان، زمان‌بندی مصوب یا ریسک‌های خارج از اختیار مدیر پروژه، باید به مالک کسب‌وکار، حامی پروژه یا مرجع تصویب تعیین‌شده ارجاع شوند.

۲-۲. حدود اختیار و محدودیت‌ها

مدیر ارشد راهبری، جایگزین مدیر پروژه یا نقش‌های تخصصی نیست و نباید در مدیریت اجرایی روزانه، جزئیات فنی و تصمیم‌های تخصصی تحلیل و طراحی مداخله مستقیم کند؛ مگر آنکه مسئولیتی رسماً تفویض شده باشد. همچنین این نقش مجاز به دورزدن فرایندهای رسمی مدیریت تغییر، کنترل کیفیت، آزمون و پذیرش کارفرما نبوده و خارج از اختیارات تفویض‌شده، تعهد قراردادی جدیدی ایجاد نمی‌کند.

۳. حوزه مدیریت و حاکمیت پروژه

این حوزه مسئول تبدیل جهت‌گیری کلان پروژه به برنامه و کنترل اجرایی، ایجاد شفافیت و انضباط کاری، مدیریت وابستگی‌ها، ریسک‌ها و تغییرات، نگهداری سوابق و مصوبات، و ارائه گزارش‌های مدیریتی است.

۱-۳. نقش‌های اصلی و لازم

◀ مدیر پروژه: مسئول برنامه‌ریزی، هماهنگی، پیگیری تعهدات، مدیریت مسائل، کنترل اجرای پروژه و گزارش‌دهی به مراجع راهبری؛

انضباط مهندسی و بازخورد تدریجی فراهم می‌کند. این چارچوب نه به تشریفات یکسان و سنگین برای همه پروژه‌ها منجر می‌شود و نه کنترل‌های ضروری را به نام چابکی کنار می‌گذارد.

تفکر سیستمی و فرایندگرایی، پرهیز از کامپیوتری کردن رویه‌های نادرست، توجه به شرایط اجرایی ایران، استفاده مسئولانه از فناوری‌های توانمندساز و حفظ مستندسازی و ردیابی‌پذیری، زمینه ورود به ساختار عملیاتی، نقش‌ها و سازوکارهای اجرایی در بخش دوم را فراهم می‌سازد.

بخش دوم: ساختار عملیاتی، نقش‌ها و مسئولیت‌های پروژه

۱. منطق حاکم بر ساختار عملیاتی و نقش‌ها

ساختار عملیاتی این چارچوب بر پایه نقش، مسئولیت، اختیار و پاسخ‌گویی مشخص شکل می‌گیرد، نه بر مبنای ایجاد واحد سازمانی یا پست مستقل برای هر فعالیت. هر نقش باید متولی مشخص داشته باشد و هیچ فعالیت یا کنترل مؤثری، به‌ویژه در حوزه‌های تصمیم‌گیری، تأیید، تغییر، آزمون، تحویل و پذیرش، نباید بدون مسئول باقی بماند.

متناسب با مقیاس، پیچیدگی، حساسیت، نوع قرارداد، تعداد سودبران و سطح ریسک پروژه، یک فرد یا واحد می‌تواند چند نقش را بر عهده گیرد؛ مشروط بر آنکه صلاحیت، ظرفیت انجام کار، تفکیک مناسب وظایف، کنترل‌های ضروری و شواهد تصمیم، آزمون، تأیید و تحویل حفظ شود؛ بنابراین، نقش‌های اصلی در این بخش به معنای وجود مسئولیت مشخص‌اند، نه الزام به ایجاد پست یا واحد مستقل.

در سوی کارفرما نیز باید، متناسب با شرایط پروژه، دست‌کم مالک کسب‌وکار، حامی پروژه، نماینده یا مرجع دارای اختیار کسب‌وکار، صاحبان فرایند و مراجع تأیید نیازمندی‌ها و پذیرش تحویل مشخص شوند. مالک کسب‌وکار مسئول تحقق ارزش، نیاز و بهره‌برداری از نتیجه است؛ حامی پروژه پشتیبانی سطح بالا، رفع موانع و تسهیل تصمیم‌های کلان را بر عهده دارد. این دو نقش می‌توانند در یک شخص جمع شوند، اما نباید تفاوت مسئولیت آن‌ها مبهم بماند. تصمیم‌ها، تأییدها و پذیرش‌های کارفرما باید در ساختار حاکمیتی پروژه ثبت و قابل پیگیری باشند.

ساختار مرجع پیشنهادی بر یک نقش راهبری کلان و چهار حوزه عملیاتی استوار است:

◀ مدیر ارشد راهبری پروژه (Senior Project Steering Manager)؛

◀ حوزه مدیریت و حاکمیت پروژه؛

◀ حوزه تحلیل و طراحی؛

◀ حوزه فنی، تولید، آزمون و تضمین کیفیت؛

◀ حوزه استقرار، آموزش و پشتیبانی.

این حوزه‌ها مبنایی برای تخصیص مسئولیت و سازمان‌دهی همکاری‌ها هستند و الزاماً به معنای ایجاد واحدها یا تیم‌های مستقل نیستند.

اجرائی و قابل استقرار است. تولید می‌تواند از طریق کدنویسی، پیکربندی محصول یا سامانه‌های مدیریت فرایند (BPMS)، توسعه خدمات و افزونه‌ها، یا ترکیبی از این‌ها انجام شود. در تمام حالات، انطباق خروجی‌ها با معیارهای پذیرش، الزامات امنیتی، یکپارچگی و استانداردهای کیفی الزامی است.

۱-۵. نقش‌های اصلی و لازم

◀ مدیر فنی / راهبر فنی: هدایت فنی، ارزیابی امکان‌پذیری، معماری کلان، مدیریت کیفیت فنی و رفع موانع تولید.

◀ تولیدکننده / پیکربند: توسعه نرم‌افزار، پیکربندی سامانه یا BPMS، و آماده‌سازی قابلیت‌ها.

◀ کارشناس آزمون و تضمین کیفیت: مسئول برنامه‌ریزی و اجرای مستقل آزمون‌ها، ثبت خطاها و ارائه شواهد انطباق با معیارهای پذیرش (تفکیک این نقش از تیم تولید به منظور جلوگیری از خوداظهاری الزامی است).

◀ مسئول مدیریت نسخه و انتشار: کنترل نسخه‌های فنی، بسته‌بندی انتشار و اطمینان از قابلیت بازگشت به نسخه‌های پیشین.

۲-۵. نقش‌های تخصصی تکمیلی

نقش‌هایی نظیر معمار نرم‌افزار، متخصص پایگاه‌داده، متخصص امنیت، مهندس دواپس (DevOps)، و متخصصان آزمون (کارایی/امنیت/خودکار) متناسب با فناوری‌های انتخابی (فناوری‌های به کاررفته)، پیچیدگی راهکار و نیازهای خاص پروژه فعال می‌شوند.

۶. حوزه استقرار، آموزش و پشتیبانی

این حوزه مسئول آماده‌سازی سازمان بهره‌بردار، انتقال دانش، تحویل عملیاتی و سازمان‌دهی پشتیبانی است. موفقیت فنی، زمانی به ارزش کسب‌وکاری تبدیل می‌شود که بهره‌برداری در فرایند واقعی کار تثبیت و مسائل عملیاتی از مسیر مشخصی مدیریت شوند.

۱-۶. نقش‌های اصلی و لازم

◀ مدیر استقرار: برنامه‌ریزی و هماهنگی استقرار، کنترل آمادگی سازمان و هدایت آموزش و پشتیبانی.

◀ کارشناس استقرار و راه‌اندازی: آماده‌سازی محیط بهره‌برداری، مهاجرت داده‌ها و راه‌اندازی فنی.

◀ کارشناس آموزش و انتقال دانش: تدوین محتوا، آموزش کاربران و تهیه راهنماها.

◀ کارشناس پشتیبانی: ثبت، دسته‌بندی و پاسخ‌گویی به رخدادهای مسائل بهره‌برداری.

۲-۶. نقش‌های تخصصی تکمیلی

متناسب با سطح خدمات (SLA) و گستره بهره‌برداری، نقش‌های تخصصی مانند کارشناس آزمون پذیرش کاربر (UAT)، متخصص مرکز خدمات (Service Desk)، و کارشناسان مدیریت سطح خدمات و دانش بهره‌برداری فعال می‌شوند.

◀ مسئول کنترل پروژه: مسئول پایش برنامه، پیشرفت، زمان، هزینه، منابع، انحرافات و وابستگی‌های پروژه؛

◀ مسئول مدیریت ریسک و تغییر: مسئول ثبت، ارزیابی، پیگیری و طرح ریسک‌ها، مسائل و درخواست‌های تغییر در مرجع تصمیم‌گیری تعیین‌شده؛

◀ مسئول کنترل مستندات و مصوبات: مسئول اعتبار نسخه‌ها، ثبت و نگهداری مصوبات، مکاتبات و سوابق رسمی پروژه.

۳-۲. نقش‌های تخصصی تکمیلی

حسب ضرورت پروژه، نقش‌هایی مانند هماهنگ‌کننده پروژه یا تحویل، اسکرآمستر، مسئول تضمین کیفیت فرایند، مسئول مدیریت قرارداد و تعهدات، مسئول پایش شاخص‌ها، هماهنگ‌کننده مراحل یا گروه‌های فرایندی، و دفتر مدیریت پروژه یا کارشناس آن، می‌توانند به ساختار افزوده شوند.

۴. حوزه تحلیل و طراحی

این حوزه مسئول شناخت و صورت‌بندی نیازمندی‌ها، تحلیل فرایندها و قواعد کسب‌وکار، طراحی راهکار و حفظ پیوند میان نیازمندی‌ها، طراحی، تولید، آزمون و تحویل است.

۴-۱. نقش‌های اصلی و لازم

◀ مدیر تحلیل: مسئول برنامه‌ریزی و هدایت فعالیت‌های تحلیل، کنترل کیفیت اقلام تحلیلی و ایجاد هماهنگی میان سودبران کسب‌وکار و تیم فنی؛

◀ تحلیلگر کسب‌وکار: مسئول استخراج، تحلیل، اولویت‌بندی و اعتبارسنجی نیازمندی‌ها، فرایندها، قواعد کسب‌وکار و سناریوهای استفاده؛

◀ تحلیلگر سامانه: مسئول تبدیل نیازمندی‌های تأییدشده به مشخصات عملکردی، مدل‌ها، منطق سامانه، تعاملات و قیود طراحی؛

◀ مسئول ردیابی‌پذیری: مسئول حفظ ارتباط محتوایی میان نیازمندی‌ها، تصمیم‌های طراحی، اقلام تولید، موارد آزمون و تحویل‌های پروژه.

کنترل مستندات و مصوبات، اعتبار نسخه‌ها و سوابق رسمی را حفظ می‌کند؛ در مقابل، ردیابی‌پذیری بر پیوند محتوایی و قابلیت دنبال کردن هر نیازمندی در چرخه تحلیل تا تحویل متمرکز است.

۴-۲. نقش‌های تخصصی تکمیلی

باتوجه به ماهیت راهکار، نقش‌هایی مانند طراح عملکردی، تحلیلگر فرایند، تحلیلگر داده، طراح تجربه کاربر، تحلیلگر یکپارچه‌سازی یا تحلیلگر حوزه تخصصی کسب‌وکار می‌توانند افزوده شوند. در پروژه‌های کوچک و متوسط، برخی از این مسئولیت‌ها می‌تواند در نقش مدیر تحلیل یا تحلیلگر ارشد تجمیع شود؛ مشروط بر حفظ کیفیت تحلیل، قابلیت ردیابی و تأیید سودبران.

۵. حوزه فنی، تولید، آزمون و تضمین کیفیت

این حوزه مسئول تبدیل نیازمندی‌های طراحی‌شده به قابلیت‌های

۷. الگوی آکاردئونی تخصیص و تجمیع نقش‌ها

ساختار ارائه‌شده، مرجع و قابل تطبیق است، نه تشکیلاتی ثابت برای همه پروژه‌ها. نقش‌های تخصصی باتوجه به ریسک، پیچیدگی، وابستگی‌ها، تعداد سودبران، حساسیت داده‌ها، الزامات امنیتی، نوع فناوری، میزان یکپارچگی، گستره استقرار و سطح خدمات موردنیاز، فعال، تفکیک یا تجمیع می‌شوند.

تجمیع نقش‌ها مجاز است، مشروط بر آنکه مسئول هر فعالیت و تحویل‌دانی، حدود اختیار و مرجع تأیید مشخص باشد؛ تصمیم‌ها، تغییرات، ریسک‌ها و مسائل مهم ثبت و قابل ردیابی بمانند؛ آزمون و کنترل کیفیت صرفاً به خوداظهاری تولیدکننده محدود نشود؛ و تحویل فقط پس از تأیید و پذیرش رسمی انجام گیرد.

برای شفاف‌سازی مسئولیت‌ها، می‌توان از ماتریس تخصیص مسئولیت (RAM) یا ماتریس مسئولیت، پاسخ‌گویی، مشورت و اطلاع‌رسانی (RACI)^{۱۱} استفاده کرد.

۸. مرزبندی همکاری میان حوزه‌ها

چهار حوزه عملیاتی، جزیره‌های جدا از یکدیگر نیستند و همکاری آن‌ها بر خروجی‌های قابل‌ارزیابی، نقاط تحویل، معیارهای پذیرش و مسئولیت‌های روشن استوار است:

- مدیریت و حاکمیت پروژه: مدیریت برنامه، وابستگی‌ها، ریسک‌ها، تغییرات و هماهنگی‌های اجرایی؛
- تحلیل و طراحی: شفاف‌سازی نیازمندی‌ها، فرایندها، قواعد کسب‌وکار و معیارهای پذیرش؛
- فنی، تولید، آزمون و تضمین کیفیت: تولید یا پیکربندی، یکپارچه‌سازی، آزمون و آماده‌سازی بسته‌های انتشار؛
- استقرار، آموزش و پشتیبانی: آمادگی سازمان، راه‌اندازی عملیاتی، آموزش و خدمات پشتیبانی؛
- مدیر ارشد راهبری: ایجاد هم‌راستایی میان حوزه‌های چهارگانه و ارجاع موضوعات کلان خارج از اختیارات به مراجع تصویب (حامی/ مالک کسب‌وکار).

هیچ حوزه‌ای نباید بدون هماهنگی با حوزه‌های دیگر، تصمیمی بگیرد که بر دامنه، هزینه، زمان، کیفیت، امنیت، بهره‌برداری یا تعهدات قراردادی اثر بگذارد. موضوعات میان حوزه‌ای باید در قالب تصمیم، تغییر، ریسک، مسئله یا خروجی قابل‌ردیابی ثبت و پیگیری شوند.

۹. جمع‌بندی

ساختار عملیاتی چارچوب بر یک نقش راهبری کلان و چهار حوزه مکمل استوار است و متناسب با شرایط واقعی هر پروژه بسط می‌یابد. نقش‌ها الزاماً پست یا افراد مستقل نیستند، اما هیچ مسئولیت اصلی

۱۰- Responsibility Assignment Matrix (RAM) ابزاری استاندارد برای تخصیص نقش‌ها و مسئولیت‌ها در فعالیت‌های پروژه که مشخص می‌کند چه کسی مسئول انجام مستقیم یک فعالیت است.

۱۱- ماتریس RACI نوعی از ماتریس RAM است که برای شفافیت بیشتر، هر نقش را در چهار سطح تعریف می‌کند: مسئول مستقیم انجام کار (Responsible)، پاسخ‌گو در برابر نتیجه (Accountable)، مشاور (Consulted) و در جریان امور قرار گیرنده (Informed).

نباید حذف یا مبهم بماند. این ساختار، در بخش بعد، در چرخه حیات پنج‌مرحله‌ای و در سطح فازها، گام‌ها، خروجی‌ها و نقاط کنترل به کار گرفته می‌شود.

بخش سوم: چرخه حیات پنج‌مرحله‌ای و اجرای چابک و گام‌به‌گام پروژه

چارچوب ترکیبی - آکاردئونی، چرخه حیات ساختاریافته مهندسی نرم‌افزار را با انعطاف‌پذیری و تحویل تدریجی روش‌های چابک تلفیق می‌کند. در این رویکرد، پنج مرحله اصلی، چارچوب ثابت تبدیل نیاز به راهکار عملیاتی را شکل می‌دهند؛ درحالی‌که سازوکار چابک، نحوه تقسیم کار، تعاملات تیم، تحویل‌های مکرر و دریافت مستمر بازخورد را سامان می‌دهد.

پنج مرحله چرخه حیات عبارت‌اند از:

۱. شناخت و آغاز؛

۲. تحلیل فرایندها و نیازمندی‌ها؛

۳. طراحی عملکردی، فنی و یکپارچگی؛

۴. تولید، مهاجرت داده، آزمون و آماده‌سازی استقرار؛

۵. استقرار، تحویل، آموزش، پشتیبانی و ارزیابی نهایی.

این پنج مرحله، فازهای خطی و یک‌باره برای کل پروژه نیستند، بلکه چرخه‌ای تکرارشونده برای اجرای هر بخش از کارند. سلسله‌مراتب شکست اجرایی کار به این صورت تعریف می‌شود:

- پروژه (Project): کل تعهدات قراردادی و اهداف کسب‌وکار.
- مرحله اجرایی (Phase): بخشی از پروژه که به یک خروجی کلان و قابل بهره‌برداری برای کارفرما ختم می‌شود.
- ماژول یا گام تحویل (Work Package / Sprint Deliverable): واحد کاری خرد، مستقل و قابل‌سنجش (مانند یک فرایند، یک قابلیت یا یکپارچگی مشخص).
- هر گام تحویل، متناسب با ماهیت و ریسک خود، مراحل پنج‌گانه را با عمق مشخصی طی می‌کند:

۱. مرحله شناخت و آغاز

هدف این مرحله، ایجاد درک مشترک، شفاف و مستند از مسئله کسب‌وکار، دامنه، سودبران، محدودیت‌ها، تعهدات قراردادی و روش پیشبرد کار است. در این مرحله، درخواست پیشنهاد، قرارداد و پیشنهادهای فنی به برنامه کلان پروژه تبدیل می‌شوند.

◀ فعالیت‌های اصلی: بازبینی قرارداد، RFP و مستندات پایه؛ شناسایی سودبران کلیدی، صاحبان فرایند و مراجع تأیید؛ تعیین مرزهای دامنه و موارد خارج از دامنه (Out of Scope)؛ شناسایی اولیة ماژول‌ها و سامانه‌های مرتبط؛ ثبت محدودیت‌ها، فرضیات و ریسک‌های اولیه؛ و توافق بر ساختار حاکمیتی و روش تصمیم‌گیری.

◀ خروجی‌های سطح تعهد متعارف: سند آغاز و دامنه پروژه، فهرست سودبران و مراجع تصویب، ساختار اولیه مراحل و ماژول‌ها، فهرست

ریسک‌های اولیه و برنامه کلان زمان‌بندی.

◀ اقلام تکمیلی (بسط آکاردئونی): منشور پروژه (Project Charter)، ساختار شکست کار (WBS)^{۱۲}، ماتریس تخصیص مسئولیت‌ها (RACI)، برنامه مدیریت ارتباطات، ثبت جامع ریسک‌ها و نقشه سامانه‌های موجود (در صورت وجود اسناد معماری سازمانی، این اسناد ورودی و قید حاکم بر آغاز پروژه خواهند بود).

◀ معیار تکمیل مرحله: شفافیت دامنه اولیه، تعیین حدود اختیارات و مراجع تصمیم‌گیری، و توافق رسمی طرفین بر مسیر ورود به تحلیل تفصیلی.

۲. مرحله تحلیل فرایندها و نیازمندی‌ها

این مرحله شناخت اولیه را به مشخصات مستند و قابل‌ارزیابی از فرایندهای سازمان، منطق کسب‌وکار، قواعد تصمیم‌گیری و رفتار موردانتظار سامانه تبدیل می‌کند. تحلیل باید تفاوت میان درخواست ظاهری و نیاز واقعی کسب‌وکار را آشکار کند.

◀ فعالیت‌های اصلی: مصاحبه با صاحبان فرایند و کاربران کلیدی؛ مشاهده و بازیابی جریان واقعی کار؛ مدل‌سازی فرایندهای وضع موجود (As-Is) و وضع مطلوب (To-Be)؛ شناسایی گلوگاه‌ها و فرصت‌های بهبود؛ استخراج نیازمندی‌های کسب‌وکار، عملکردی و غیرعملکردی؛ تعیین قواعد کسب‌وکار (Business Rules)؛ و تدوین معیارهای پذیرش اولیه (Acceptance Criteria).

◀ خروجی‌های سطح تعهد متعارف: سند نیازمندی‌های کسب‌وکار (BRD)^{۱۳} یا سند جامع تحلیل، مدل یا شرح فرایندهای اصلی، فهرست اولویت‌بندی شده نیازمندی‌ها، قواعد کلیدی کسب‌وکار، معیارهای پذیرش و فهرست ابهامات نیازمند تصمیم‌گیری.

◀ اقلام تکمیلی (بسط آکاردئونی): مدل‌سازی فرایندها با نمادگذاری BPMN، شناسنامه تفصیلی فرایندها، تحلیل شکاف (Gap Analysis)، واژه‌نامه کسب‌وکار (Business Glossary)، سناریوهای کاربری (Use Cases / User Stories) و ماتریس اولیه ردیابی نیازمندی‌ها (Traceability Matrix).

معیار تکمیل مرحله: تأیید رسمی سند تحلیل و نیازمندی‌ها توسط صاحبان فرایند و کارفرما به‌عنوان مبنای طراحی و کنترل تغییرات.

۳. مرحله طراحی عملکردی، فنی و یکپارچگی

در این مرحله، نیازمندی‌ها به مشخصات دقیق، منسجم و قابل‌پیاده‌سازی تبدیل می‌شوند. طراحی شامل سه لایه درهم‌تنیده است:

۱. طراحی عملکردی: تعیین رفتار سیستم، جریان داده، ساختار فرم‌ها، کار تابل‌ها، گزارش‌ها و ماتریس دسترسی؛

۲. طراحی فنی و معماری: انتخاب الگوهای معماری، مدل داده، اجزای نرم‌افزاری، محیط‌های اجرا، پروتکل‌های امنیتی و الزامات کارایی؛

۳. طراحی یکپارچگی: تعیین سامانه‌های مرتبط، مشخصات رابط‌های برنامه‌نویسی کاربردی (API)، ساختار تبادل داده، احراز هویت و مدیریت خطاها.

◀ خروجی‌های سطح تعهد متعارف: سند مشخصات عملکردی (FRD)^{۱۴}، سند مشخصات فنی و معماری، مدل داده و ساختار پایگاه‌داده، ماتریس نقش‌ها و سطوح دسترسی، و معیارهای دقیق آزمون و پذیرش.

◀ اقلام تکمیلی (بسط آکاردئونی): طراحی تفصیلی تجربه و رابط کاربری (UI/UX)^{۱۵}، مشخصات جامع وب‌سرویس‌ها و رابط‌های تبادل داده، سند معماری امنیت و رمزنگاری، طراحی سازوکارهای ثبت رخداد (Logging & Auditing)، و توافق‌نامه سطح خدمت سرویس‌های یکپارچگی.

◀ معیار تکمیل مرحله: تأیید طراحی عملکردی توسط کارفرما و تصویب معماری فنی و یکپارچگی توسط راهبران فنی پروژه.

۴. مرحله تولید، مهاجرت داده، آزمون و آماده‌سازی استقرار

این مرحله، طراحی مصوب را به یک سامانه عملیاتی، آزموده شده و آماده بهره‌برداری تبدیل می‌کند. توسعه می‌تواند از طریق کدنویسی، پیگیربندی سامانه مدیریت فرایند کسب‌وکار (BPMS)، توسعه سرویس‌ها یا ترکیبی از آن‌ها صورت گیرد.

◀ فعالیت‌های اصلی: پیاده‌سازی قابلیت‌ها بر پایه FRD و معماری فنی؛ راه‌اندازی خطوط لوله یکپارچه‌سازی و تحویل مداوم (DevOps/CI-CD)؛ آماده‌سازی داده‌های آزمایشی؛ طراحی سناریوهای مهاجرت و پالایش داده‌های پیشین؛ اجرای آزمون‌های واحد، یکپارچگی، عملکردی و امنیتی؛ رفع خطاها و اجرای آزمون مجدد؛ و آماده‌سازی مستندات استقرار.

◀ خروجی‌های سطح تعهد متعارف: نسخه نرم‌افزاری قابل آزمون، سناریوها و گزارش نتایج آزمون، ثبت خطاها و وضعیت رفع آن‌ها، گزارش مهاجرت آزمایشی داده، راهنمای نصب و پیگیربندی، و صورت‌جلسه آمادگی برای استقرار.

◀ اقلام تکمیلی (بسط آکاردئونی): سناریوهای آزمون خودکار (Auto-mated Testing)، آزمون‌های غیرعملکردی (بار، کارایی و نفوذپذیری با امکان توجه به مدل کیفی ISO/IEC 25010:2023)، برنامه تفصیلی پالایش داده‌ها (Data Cleansing)، برنامه بازگشت از بحران (Roll-back Plan)، و ماتریس ردیابی آزمون به نیازمندی.

◀ معیار تکمیل مرحله: قبولی در آزمون‌های عملکردی اصلی، تعیین تکلیف خطاهای بحرانی و صدور تأییدیه فنی ورود به محیط استقرار.

۵. مرحله استقرار، تحویل، آموزش، پشتیبانی و ارزیابی نهایی

استقرار فراتر از نصب فنی نرم‌افزار است؛ این مرحله شامل آماده‌سازی

14- Functional Requirements Document

15- User Interface / User Experience

12- Work Breakdown Structure

13- Business Requirements Document

اسناد نیستند، بلکه نقاط رسمی ارزیابی ریسک، تصمیم‌گیری درباره تغییرات و هدایت آگاهانه به گام بعدی محسوب می‌شوند. تصویب هر نقطه مانع از بازنگری کنترل‌شده نیازمندی‌ها یا طراحی در گام‌های بعدی نبوده و در صورت لزوم، اصلاحات از طریق سازوکار مدیریت تغییر اعمال و هم‌راستایی اسناد حفظ می‌شود.

۳-۶. یکپارچگی سامانه و تحویل پیوسته

تحویل‌های مکرر و موفق گام‌های کاری، نباید جایگزین آزمون یکپارچگی کل سامانه شود. پیش از تحویل نهایی، هم‌خوانی داده‌های مشترک، گزارش‌های تجمیعی، تبادل داده میان ماژول‌ها و عملکرد تحت بار واقعی کل سامانه باید به‌طور مستقل راستی‌آزمایی شود.

جمع‌بندی بخش سوم

چرخه حیات پنج‌مرحله‌ای، اسکلتی پایدار برای فرایند مهندسی نرم‌افزار فراهم می‌کند که از شناخت تا بهره‌برداری را تضمین می‌نماید؛ درحالی‌که اجرای ماژولار و گام‌به‌گام، پویایی و انطباق‌پذیری چابک را به ارمغان می‌آورد. تمایز میان تعهدات متعارف و اقلام تکمیلی (آکاردئونی)، از سنگین‌شدن بی‌مورد پروژه‌های کوچک و نیز غفلت از کنترل‌های حساس در پروژه‌های پیچیده جلوگیری می‌کند. این ساختار اجرایی، زمینه را برای استقرار سازوکارهای پشتیبان و حاکمیتی (موضوع بخش بعد) فراهم می‌سازد.

بخش چهارم: سازوکارهای پشتیبان، حاکمیت و بهبود مستمر

موفقیت چارچوب ترکیبی - آکاردئونی صرفاً به تعریف چرخه حیات پنج‌مرحله‌ای، مرحله‌بندی پروژه و تعیین نقش‌ها وابسته نیست؛ بلکه به سازوکارهایی نیاز دارد که انضباط مهندسی، شفافیت تصمیم‌گیری، کنترل تغییر، ردیابی‌پذیری، مدیریت ریسک، کیفیت، مدیریت دانش و بهبود مستمر را در سراسر پروژه برقرار سازند. این سازوکارها، نقش‌ها، مراحل چرخه حیات، مراحل، ماژول‌ها، گام‌های تحویل، اسناد، تحویل‌دانی‌ها و تصمیم‌های مدیریتی را در سطح کل پروژه، مراحل و هر گام تحویل به یکدیگر پیوند می‌دهند. این پیوند باید قابل‌مشاهده، قابل‌بررسی و متناسب با شرایط پروژه باشد.

سطح اجرای سازوکارهای پشتیبان در همه پروژه‌ها یکسان نیست. در سطح تعهد متعارف، سازوکارها و شواهد ضروری برای کنترل دامنه، تصمیم‌گیری، تحلیل، طراحی، تولید، آزمون، تحویل و پذیرش اجرا می‌شوند. اقلام تکمیلی نیز بر پایه عواملی مانند ریسک، پیچیدگی، حساسیت داده‌ها، اهمیت عملیاتی، تعداد سودبران، وابستگی‌ها، الزامات امنیتی و توافق قراردادی تعیین می‌شوند.

رویکرد آکاردئونی به معنای حذف کنترل‌های ضروری نیست؛ بلکه از تحمیل بی‌دلیل اسناد، نقش‌ها، رویه‌ها و کنترل‌های تخصصی به

زیرساخت، انتقال قطعی داده‌ها، آموزش کاربران، تحویل رسمی، پایدارسازی سامانه و سازمان‌دهی پشتیبانی است.

فعالیت‌های اصلی: استقرار نسخه در محیط عملیاتی؛ اعمال یکپارچگی‌های امنیتی و شبکه‌ای؛ اجرای آزمایشی در صورت نیاز؛ مهاجرت نهایی داده‌ها؛ آموزش کاربران نهایی، مدیران سامانه و تیم راهبری؛ پشتیبانی دوره تثبیت (Hypercare)؛ ثبت و رسیدگی به رخدادهای ارزیابی نهایی انطباق با نیازمندی‌ها؛ و اخذ تأییدیه تحویل.

خروجی‌های سطح تعهد متعارف: راهنمای کاربری و راهبری سامانه، گزارش اجرای استقرار، سوابق و صورت‌جلسات آموزش، گزارش مهاجرت قطعی داده، سازوکار ثبت و پیگیری درخواست‌های پشتیبانی، فهرست موارد باز و برنامه رفع آن‌ها، و صورت‌جلسه رسمی تحویل و پذیرش کاربران (UAT Sign-off).

اقلام تکمیلی (بسط آکاردئونی): برنامه جامع مدیریت تغییرات سازمانی (OCM)^۶، تدوین توافق‌نامه سطح خدمت (SLA)، فرایندهای مدیریت رخداد و تغییر با امکان توجه به چارچوب (2019) ITIL 4، گزارش ارزیابی تحقق ارزش کسب‌وکار، و گزارش جامع خاتمه پروژه. معیار تکمیل مرحله: بهره‌برداری موفق سامانه در محیط واقعی، آموزش سودبران مرتبط، برقراری بستر پشتیبانی و امضای صورت‌جلسه رسمی پذیرش.

۶. قواعد تکمیلی در اجرای چرخه حیات

۱-۶. اجرای فازبندی‌شده، ماژولار و متولی گام

در پروژه‌های سفارشی، شکستن کار به ماژول‌ها و گام‌های مستقل، ریسک را مدیریت کرده و بازخوردهای سریع‌تری فراهم می‌کند. برای هر گام باید یک «متولی گام» تعیین شود که مسئول هماهنگی فعالیت‌های مربوط به آن گام باشد. متولی گام، مجری تمامی کارها نیست، بلکه هماهنگ‌کننده پیشرفت کار با تیمی چندمهارتی (Cross-functional Team) است. همچنین، متولی گام نمی‌تواند به‌تنهایی مرجع پذیرش کیفی همان گام باشد؛ تأیید نهایی بر عهده مراجع مستقل و نمایندگان کارفرماست.

۲-۶. نقاط بازبینی و تصمیم‌گیری (Phase Gates)

نقاط کنترل اصلی پروژه عبارت‌اند از:

۱. تأیید آغاز و حدود دامنه (Scope)؛

۲. تأیید سند تحلیل نیازمندی‌ها؛

۳. تصویب طراحی عملکردی و معماری فنی؛

۴. تأیید آمادگی برای استقرار؛

۵. پذیرش گام یا ماژول تحویلی؛

۶. تحویل نهایی و خاتمه.

این نقاط به معنای ایجاد موانع فرایندی یا انجماد بازگشت‌ناپذیر

۲. نسبت چارچوب با مطالعات و پروژه‌های بالادستی

در برخی سازمان‌ها، پیش از آغاز پروژه نرم‌افزاری، مطالعاتی مانند معماری سازمانی، تحلیل یا بازطراحی فرایندها، مدل‌سازی داده، نقشه راه فناوری اطلاعات، برنامه تحول دیجیتال یا امکان‌سنجی انجام شده است. خروجی‌های معتبر این مطالعات می‌توانند در شناخت و آغاز، تحلیل، طراحی و تعیین دامنه، ورودی‌های ارزشمندی باشند.

با این حال، این اسناد باید با اهداف، دامنه، شرایط واقعی، محدودیت‌ها و زمان اجرای پروژه تطبیق و اعتبارسنجی شوند. وجود آن‌ها به معنای پذیرش خودکار مفروضاتشان نیست؛ زیرا قوانین، فرایندها، ساختار سازمانی، سامانه‌های موجود، فناوری‌ها یا اولویت‌های کسب‌وکار ممکن است از زمان تهیه آن‌ها تغییر کرده باشند.

اجرای این چارچوب به وجود مطالعات یا اسناد بالادستی وابسته نیست. در صورت نبود آن‌ها، فعالیت‌های شناخت، تحلیل فرایندها، استخراج نیازمندی‌ها، طراحی و تعیین دامنه در محدوده همان پروژه انجام می‌شود.

اسناد معماری سازمانی، نقشه راه و تحول دیجیتال، در صورت وجود، می‌توانند جهت‌گیری‌ها، محدودیت‌ها و وابستگی‌های کلان را مشخص کنند؛ اما پروژه باید آن‌ها را متناسب با دامنه هر مرحله و گام، به نیازمندی‌ها، طراحی‌ها، فعالیت‌های اجرایی و خروجی‌های قابل تحویل تبدیل کند.

بنابراین، پروژه می‌تواند از مطالعات پیشین استفاده کند، اما مسئولیت اعتبارسنجی، تکمیل، به‌روزرسانی و تبدیل آن‌ها به نیازمندی‌ها و طراحی‌های قابل اجرا، در چارچوب پروژه باقی می‌ماند.

۳. حاکمیت پروژه و مدیریت کنترل‌شده تغییرات

حاکمیت پروژه بر برقراری تعادل میان پایداری تعهدات و انعطاف‌پذیری کنترل‌شده استوار است. قرارداد، درخواست پیشنهاد (RFP)، دامنه مصوب و اسناد پایه، مراجع تعهدات پروژه‌اند؛ اما نباید اصلاح نیازمندی‌ها یا واکنش به شرایط جدید را ناممکن سازند.

مدیر ارشد راهبری پروژه، در تعامل با مالک یا حامی پروژه، مدیر پروژه، نماینده ارشد کارفرما و دیگر مراجع ذی‌صلاح، مسئول هدایت تصمیم‌های کلان، حل تعارض‌های سطح بالا، حفظ هم‌راستایی اهداف کسب‌وکار با مسیر پروژه و هدایت سازوکار تصمیم‌گیری است. این نقش جایگزین مدیر پروژه اجرایی نیست و مسئولیت اداره روزمره فعالیت‌ها را بر عهده ندارد.

در سطح اجرایی، مدیر پروژه مسئول برنامه‌ریزی، هماهنگی، پیگیری پیشرفت، مدیریت وابستگی‌ها و کنترل فعالیت‌های کل پروژه است. در سطح هر گام نیز متولی گام، هماهنگی، پیگیری و آماده‌سازی آن گام برای بازبینی، تحویل یا پذیرش را بر عهده دارد. این سه سطح باید از نظر مسئولیت، حدود اختیار و مسیر ارجاع مسائل، از یکدیگر متمایز باشند.

تأیید اسناد پایه به معنای توقف تغییر نیست؛ بلکه آغاز مدیریت

پروژه‌هایی جلوگیری می‌کند که به آن سطح از تفصیل نیاز ندارند. تبدیل بی‌دلیل ارقام تکمیلی به تعهدات عمومی، زمان و هزینه را افزایش می‌دهد، تحویل ارزش را به تأخیر می‌اندازد و احتمال اختلاف درباره حدود تعهدات را بیشتر می‌کند.

۱. معماری مستندات و ایجاد زبان مشترک

پراکندگی اطلاعات، برداشت‌های متفاوت از نیازمندی‌ها و نبود زبان مشترک میان کارفرما، صاحبان فرایند، تحلیلگران، تیم فنی، کاربران و دیگر سودبران، از عوامل اصلی ابهام، دوباره‌کاری و اختلاف در پروژه‌های نرم‌افزاری سفارشی است. معماری مستندات، مرجعی مشترک برای شناخت، تحلیل، طراحی، تولید، آزمون، استقرار و تحویل فراهم می‌کند.

در سطح تعهد متعارف، هسته معماری مستندات بر محتوای زیر استوار است:

- سند تحلیل فرایندها و نیازمندی‌ها (BRD): مرجع اهداف کسب‌وکار، فرایندهای وضع موجود و مطلوب، منطق و قواعد تصمیم‌گیری و معیارهای پذیرش اولیه؛

- سند مشخصات عملکردی (FRD): مرجع مشخصات رفتاری سامانه، سناریوهای کاربری، جریان داده، کارتابل‌ها، گزارش‌ها و معیارهای پذیرش تفصیلی؛

- مشخصات طراحی یکپارچگی و سرویس‌ها: مرجع تعامل با سایر سامانه‌ها، روش‌های تبادل داده، مشخصات API‌ها، ساختار پیام، امنیت تبادل و مدیریت خطا.

در صورت وجود داده‌های حساس، معماری فنی پیچیده، الزامات امنیتی، مهاجرت داده یا آزمون‌های تخصصی، محتوای مربوط می‌تواند در اسناد مادر، پیوست‌های آن‌ها یا اسناد تخصصی مستقل ثبت شود. معیار تفکیک اسناد باید کنترل پیچیدگی، کاهش ریسک، افزایش قابلیت استفاده و نگهداری، یا پاسخ‌گویی به الزام مشخص باشد؛ نه صرفاً افزایش تعداد اسناد.

برای هر گام تحویل، ارتباط میان دامنه گام، نیازمندی‌ها، طراحی، فعالیت‌های تولید یا پیکربندی، آزمون، موارد باز و پذیرش باید روشن باشد. این اطلاعات می‌تواند در سندی مستقل، پیوست سند مادر، کارتابل پروژه یا ترکیبی از آن‌ها ثبت شود؛ مشروط بر آنکه مرجع معتبر، قابل‌شناسایی و قابل بازبینی داشته باشد.

اصول پایه معماری مستندات عبارت‌اند از:

۱. مرجع ایجاد هر نیازمندی، تصمیم یا تغییر مشخص باشد؛
۲. مبنای تولید هر سند، فعالیت یا تحویل‌دانی بعدی روشن باشد؛
۳. مسئول تهیه، مرجع بازبینی و مرجع تأیید هر خروجی تعیین شود
۴. تغییرات مهم در اسناد وابسته، برنامه‌ها، معیارهای پذیرش و تحویل‌دانی‌های مرتبط منعکس شوند؛
۵. نسخه معتبر اسناد، سوابق تصمیم‌ها و تاریخچه تغییرات قابل‌شناسایی باشد.

کنترل شده تغییرات است. هر درخواست تغییر باید، متناسب با سطح و اهمیت آن، دست کم شامل موضوع تغییر، دلیل، درخواست کننده، اولویت، وضعیت و اثر احتمالی بر موارد زیر باشد:

- دامنه، مراحل گام‌های تحویل؛
- زمان، هزینه و منابع؛
- کیفیت و معیارهای پذیرش؛
- معماری، طراحی و فناوری؛
- داده و مهاجرت داده؛
- امنیت، محرمانگی و سطح دسترسی؛
- یکپارچگی‌ها و سرویس‌ها؛
- آزمون، استقرار و بهره‌برداری؛
- مستندات، آموزش و پشتیبانی؛
- تعهدات قراردادی.

پس از تحلیل اثر، درباره پذیرش، رد، تعویق، اولویت‌بندی مجدد یا جایگزینی تغییر، در سطح اختیار مناسب تصمیم‌گیری می‌شود. تغییرات تأیید شده باید در اسناد مرتبط، برنامه پروژه، فهرست کارها، معیارهای پذیرش و سوابق ردیابی منعکس شوند و در صورت لزوم، قرارداد یا صورت جلسه رسمی نیز به‌روزرسانی شود.

تغییرات محدود در یک گام می‌توانند در سطح متولی گام و مدیر پروژه ثبت و تعیین تکلیف شوند. اما تغییراتی که بر دامنه مرحله، معماری کلان، هزینه، زمان نهایی، امنیت، یکپارچگی‌های اصلی یا تعهدات قراردادی اثر می‌گذارند، باید به مرجع حاکمیتی ذیصلاح ارجاع شوند.

۴. ردیابی زنجیره ارزش و نیازمندی‌ها

برای اطمینان از لحاظ‌شدن تعهدات قراردادی و نیازمندی‌های مصوب در راهکار نهایی، و جلوگیری از تولید قابلیت‌های فاقد پشتوانه کسب‌وکاری، رابطه میان نیاز، فرایند، طراحی، اجرا، آزمون، تحویل و پذیرش باید قابل ردیابی باشد.

هر نیاز کسب‌وکار ثبت شده در سند BRD باید، در صورت ارتباط، به نیازمندی عملکردی، طراحی قابل اجرا یا معیار پذیرش در FRD یا معادل آن متصل شود. فعالیت‌های تولید یا پیکربندی، سناریوها و نتایج آزمون، خطاها، اقدامات اصلاحی، آزمون پذیرش کاربر (UAT) و تأیید نهایی نیز باید با نیاز مربوط مرتبط و قابل بررسی باشند. ردیابی در دو سطح برقرار می‌شود:

الف) ردیابی در سطح پروژه و مرحله

در این سطح، ارتباط میان قرارداد، اهداف، دامنه، فازها، ماژول‌ها، گام‌های تحویل، خروجی‌های اصلی، تغییرات و وضعیت پذیرش مشخص می‌شود.

ب) ردیابی در سطح گام تحویل

در این سطح، ارتباط میان هدف گام، فرایند یا قابلیت موردنظر، نیازمندی‌ها، طراحی، فعالیت‌های تولید یا پیکربندی، سناریوها و

نتایج آزمون، خطاها، موارد باز و تأیید گام برقرار می‌گردد.

در سطح تعهد متعارف، ردیابی با جداول ساده، پیوست اسناد یا ابزارهای مدیریت کار انجام می‌شود. در پروژه‌های پرریسک یا دارای وابستگی‌های گسترده، تدوین ماتریس مستقل ردیابی نیازمندی‌ها (RTM) برای برقراری پیوند دوطرفه میان نیازهای کسب‌وکار، مشخصات فنی، سناریوهای آزمون و وضعیت پذیرش ضرورت می‌یابد. ردیابی صرفاً ابزاری برای کنترل مستندات نیست، بلکه مبنایی برای موارد زیر نیز هست:

- ارزیابی پیشرفت واقعی پروژه؛
- بررسی کامل بودن دامنه هر گام؛
- تحلیل اثر تغییرات؛
- شناسایی قابلیت‌های فاقد نیاز مصوب؛
- شناسایی نیازمندی‌های فاقد خروجی، آزمون یا پذیرش؛
- ارزیابی وضعیت تحویل‌دانی‌ها بر پایه شواهد قابل بررسی.

۵. کنترل کیفیت و معیارهای پذیرش

کیفیت در این چارچوب حاصل یک بررسی نهایی و جداگانه نیست؛ بلکه در سراسر چرخه حیات، از شناخت و تحلیل تا استقرار و بهره‌برداری، ایجاد و کنترل می‌شود.

در تحلیل، نیازمندی‌ها، قواعد کسب‌وکار و معیارهای پذیرش باید به گونه‌ای ثبت شوند که ابهام و برداشت‌های متناقض کاهش یابد. در طراحی، سازگاری طراحی عملکردی، فنی، داده‌ای و یکپارچگی بررسی می‌شود. در تولید، قابلیت‌ها بر اساس مشخصات تأیید شده تولید یا پیکربندی و آزمون می‌شوند. در استقرار نیز شواهد آزمون، وضعیت خطاها، آمادگی کاربران، صحت داده‌ها و نتایج آزمون پذیرش بررسی می‌گردند.

برای هر قابلیت، ماژول یا گام تحویل، معیار تکمیل باید متناسب با ماهیت و ریسک آن، شامل موارد زیر باشد:

- دامنه و خروجی مصوب؛
- نیازمندی‌ها یا مشخصات مبنای؛
- شواهد لازم تحلیل و طراحی؛
- نسخه تولید شده یا پیکربندی شده قابل بررسی؛
- سناریوها و نتایج آزمون؛
- وضعیت خطاها و موارد باز؛
- مستندات، آموزش یا آمادگی بهره‌برداری مورد نیاز؛
- تأیید مرجع بازبینی یا پذیرش.

سپری شدن زمان، برگزاری جلسه یا اعلام شفاهی تیم، به تنهایی نشانه تکمیل یا پذیرش نیست.

کنترل کیفیت می‌تواند متناسب با ماهیت هر گام توسط اعضای تیم انجام شود؛ اما متولی گام نباید تنها مرجع تأیید خروجی همان گام باشد. بازبینی تخصصی و پذیرش کسب‌وکاری باید، حسب مورد، توسط عضو مستقل تیم، مسئول تخصصی مربوط، نماینده مجاز

کارفرما یا مرجع تأیید تعیین شده انجام شود.

اقدام تکمیلی کنترل کیفیت، متناسب با ریسک و پیچیدگی پروژه، می‌تواند شامل برنامه تفصیلی آزمون، آزمون خودکار، آزمون کارایی، آزمون امنیت، گزارش پوشش آزمون، مدیریت عدم انطباق، بازبینی رسمی کد یا طراحی و ماتریس ردیابی نیازمندی‌ها باشد.

۶. مدیریت ریسک‌های بومی پروژه‌های سفارشی

مدیریت ریسک باید از مرحله شناخت آغاز شود و تا تحویل و دوره تثبیت بهره‌برداری ادامه یابد. در پروژه‌های نرم‌افزاری سفارشی، برخی ریسک‌ها به سبب شرایط سازمان، داده‌ها، وابستگی‌ها و سازوکارهای تصمیم‌گیری اهمیت بیشتری دارند و باید از ابتدا در برنامه پروژه، تحلیل، طراحی و تصمیم‌های اجرایی لحاظ شوند.

۶-۱. داده و مهاجرت داده

کیفیت نامشخص داده‌های قدیمی، ناسازگاری ساختارها، ناقص بودن اطلاعات، دشواری پاک‌سازی و احتمال اختلال در بهره‌برداری، از مهم‌ترین ریسک‌های داده و مهاجرت هستند. راهبرد انتقال، قواعد تبدیل، مسئولیت پاک‌سازی، انتقال آزمایشی، کنترل کیفیت داده و روش بازگشت باید در نخستین فرصت تعیین شوند.

۶-۲. وابستگی و یکپارچگی

تأخیر یا تغییر خدمات بیرونی، وابستگی به تیم‌های ثالث، نبود دسترسی به محیط آزمون و ابهام در تبادل داده می‌تواند تحویل گام‌ها و مراحل را مختل کند. مشخصات خدمات، مسئولیت طرف‌های درگیر، محیط‌های آزمون و برنامه آزمون یکپارچگی باید پیش از تولید یا پیکربندی نهایی روشن شوند.

۶-۳. تغییر قوانین، فرایندها و سیاست‌ها

تغییر مقررات، ساختار سازمانی، سیاست‌های داخلی یا قواعد کسب‌وکار می‌تواند نیازمندی‌ها و طراحی را تحت تأثیر قرار دهد. استفاده از پارامترهای قابل تنظیم، قواعد قابل مدیریت و طراحی منعطف می‌تواند هزینه واکنش به تغییر را کاهش دهد؛ اما نباید جایگزین تصمیم‌گیری رسمی درباره تغییر دامنه شود.

۶-۴. تصمیم‌گیری و تأیید

تأخیر در معرفی صاحبان فرایند، نبود مرجع تأیید، تعارض میان سودبران یا چندگانگی تصمیم‌ها، از موانع مهم پروژه است. ساختار تصمیم‌گیری، صاحبان فرایند، مراجع تأیید و مسیر ارجاع مسائل باید از ابتدای پروژه مشخص شود.

۶-۵. دانش و منابع انسانی

وابستگی به افراد کلیدی، جابه‌جایی نیروها و انتقال ناقص دانش می‌تواند استمرار اجرای گام‌ها را تهدید کند. ثبت تصمیم‌ها، مستندسازی متناسب، آموزش متقابل، آشنایی اعضا با حوزه‌های مرتبط و تعیین جانشین برای نقش‌های حساس، ابزارهای کنترل

این ریسک هستند.

در تیم‌های چابک، تقسیم کار نباید به وابستگی مطلق یک فعالیت یا دانش حیاتی به یک فرد منجر شود. اعضای تیم باید، متناسب با نقش خود، با تحلیل، طراحی، تولید، آزمون، استقرار و پشتیبانی آشنایی عملی داشته باشند؛ با این حال، مسئولیت تصمیم‌های تخصصی باید نزد افراد دارای صلاحیت باقی بماند.

۶-۶. امنیت و محرمانگی

دسترسی نامناسب، استفاده کنترل نشده از داده‌های واقعی، ضعف کنترل محیط‌های آزمون و افشای اطلاعات حساس، از ریسک‌های مهم امنیتی هستند. الزامات امنیتی، سطح محرمانگی داده، شیوه دسترسی و محدودیت‌های استفاده از اطلاعات باید از مراحل تحلیل و طراحی لحاظ شوند.

برای هر ریسک مهم باید، متناسب با شرایط پروژه، مالک ریسک، احتمال وقوع، میزان اثر، اولویت، اقدام پیشگیرانه، برنامه واکنش و وضعیت پیگیری مشخص باشد.

۷. استفاده کنترل شده و اثربخش از هوش مصنوعی و قابلیت‌های داده‌محور

هوش مصنوعی و قابلیت‌های داده‌محور، ابزارهای توانمندساز مشترک برای کارفرما و مجری‌اند و می‌توانند، در صورت استفاده کنترل شده، سرعت، دقت و کیفیت فعالیت‌هایی مانند تحلیل و مقایسه اسناد، استخراج اولیه نیازمندی‌ها، تدوین سناریوهای آزمون، پیش‌نویس مستندات، تولید نمونه کد، تحلیل خطاها، مدیریت دانش، پایش شاخص‌ها و تحلیل الگوهای بهره‌برداری را افزایش دهند.

مجری موظف است از این قابلیت‌ها به صورت حرفه‌ای، متناسب با ماهیت و حساسیت پروژه و همراه با بازبینی انسانی استفاده کند. همچنین، مجری باید در طراحی کلان و معماری سامانه، متناسب با نیازهای مورد توافق، ملاحظات امنیتی، کیفیت و دسترسی پذیری داده‌ها و توجیه‌پذیری فنی و اقتصادی، قابلیت توسعه و به کارگیری آتی هوش مصنوعی و تحلیل‌های داده‌محور را برای کارفرما پیش‌بینی و مهیا سازد. این الزام به معنای اجرای فوری قابلیت‌های هوش مصنوعی نیست، مگر آنکه دامنه، اولویت‌ها و معیارهای پذیرش آن به صراحت در قرارداد یا تصمیم‌های مصوب پروژه تعیین شده باشد.

کارفرما نیز باید با تعیین ضوابط دسترسی و محرمانگی، فراهم‌سازی داده‌ها و اطلاعات مجاز و قابل اتکا، تعیین سطح مجاز کاربرد ابزارها، و ایجاد سازوکار لازم برای بررسی و پذیرش خروجی‌ها، زمینه استفاده صحیح و اثربخش از این قابلیت‌ها را فراهم سازد.

با وجود مزایای بالقوه، هوش مصنوعی جایگزین مسئولیت حرفه‌ای، تصمیم‌گیری انسانی، تحلیل و طراحی معماری، کنترل کیفیت، ملاحظات امنیتی، یا تأیید و پذیرش کارفرما نیست. باتوجه به ریسک‌هایی مانند تولید اطلاعات نادرست، سوگیری، ضعف‌های امنیتی و افشای داده‌های حساس، رعایت اصول زیر الزامی است:

۱. هیچ خروجی تولیدشده یا پیشنهادشده توسط هوش مصنوعی نباید بدون بررسی، اصلاح در صورت نیاز و اعتبارسنجی مرجع انسانی ذیصلاح، وارد اسناد رسمی، کد سامانه، تصمیم‌های مدیریتی یا محیط عملیاتی شود؛
۲. مسئولیت صحت، کیفیت، امنیت، انطباق و قابلیت اتکای خروجی‌ها همواره بر عهده نقش یا مرجع انسانی مربوط در کارفرما و مجری است و استفاده از هوش مصنوعی این مسئولیت را منتقل یا سلب نمی‌کند؛
۳. کارفرما و مجری باید پیش از استفاده از ابزارهای هوش مصنوعی، حدود دسترسی، سطح محرمانگی داده‌ها، نوع ابزارهای مجاز و الزامات نگهداری یا اشتراک‌گذاری اطلاعات را تعیین و مستند کنند؛
۴. داده‌های محرمانه، اطلاعات شخصی، اطلاعات تجاری حساس، کدهای اختصاصی و سایر دارایی‌های اطلاعاتی پروژه، نباید بدون مجوز صریح، ارزیابی ریسک و کنترل‌های مناسب در اختیار ابزارهای عمومی یا فاقد تأیید قرار گیرند؛
۵. استفاده از ابزارهای هوش مصنوعی باید با سیاست‌های امنیت اطلاعات، الزامات حریم خصوصی، تعهدات قراردادی و ضوابط حاکمیتی کارفرما سازگار باشد؛
۶. در موضوعات حساس، از جمله تحلیل‌های اثرگذار، طراحی معماری، تصمیم‌های امنیتی، تولید یا اصلاح کدهای کلیدی و تحلیل داده‌های عملیاتی، سوابق استفاده از هوش مصنوعی، بازبینی انسانی، اصلاحات اعمال‌شده و تصمیم‌نهایی باید متناسب با ریسک، قابل پیگیری باشد؛
۷. کارفرما باید در پذیرش و بهره‌برداری از خروجی‌های مرتبط، صرفاً به تولید یا پیشنهاد هوش مصنوعی اتکا نکند و سازوکارهای لازم برای ارزیابی کسب‌وکاری، تأیید مسئولانه و پاسخ‌گویی در قبال تصمیم‌نهایی را برقرار سازد.

۸. بهبود مستمر و مدیریت درس‌آموخته‌ها

بهبود مستمر باید در طول پروژه و پس از آن دنبال شود، نه فقط در خاتمه پروژه. در پایان هر مرحله، گام مهم یا نقطه بازبینی، کیفیت خروجی‌ها، گلوگاه‌ها، دوباره‌کاری‌ها، خطاهای تکرارشونده، کیفیت مستندات، سرعت تصمیم‌گیری و اثربخشی کنترل‌ها بررسی می‌شود. در اجرای چابک، این بازبینی می‌تواند در پایان چرخه‌های کاری یا پس از تحویل گام‌ها انجام شود. هر بازبینی باید به اقدام مشخصی، مانند اصلاح قالب اسناد، معیارهای پذیرش، ترتیب فعالیت‌ها، نقاط کنترل، تقسیم مسئولیت‌ها، آموزش تیم یا ابزارهای مدیریت کار، آزمون و پشتیبانی منجر شود. درس‌آموخته‌ها باید قابل‌استفاده و متناسب با اهمیت موضوع ثبت شوند و دست‌کم شامل موضوع، زمینه وقوع، اثر، علت مؤثر، اقدام اصلاحی یا پیشگیرانه، مسئول پیگیری و نحوه استفاده در مراحل یا

پروژه‌های بعدی باشند.

خروجی این فرایند می‌تواند به اصلاح الگوهای مستندات و بازبینی‌ها، بهبود معیارهای پذیرش و کنترل‌ها، بازنگری نقش‌ها، تکمیل پایگاه دانش، ارتقای ابزارها و آموزش تیم منجر شود.

شاخص‌های بهبود، در صورت نیاز و متناسب با شرایط پروژه، می‌توانند انحراف زمان و هزینه، تغییرات، دوباره‌کاری، خطاهای بحرانی یا پس از استقرار، زمان رفع خطا و تصمیم‌گیری، پذیرش به موقع گام‌ها، تحقق ارزش موردانتظار و میزان استفاده عملیاتی از خروجی‌ها را پوشش دهند. اعمال شاخص‌هایی که نتایج آن‌ها استفاده نمی‌شود، بار اجرایی غیرضروری ایجاد می‌کند.

بخش پنجم: جمع‌بندی و مسیر توسعه چارچوب

۱. جمع‌بندی و برآیند راهبردی

چارچوب ترکیبی - آکاردئونی، الگویی ساختاریافته و درعین حال منعطف برای هدایت پروژه‌های نرم‌افزاری سفارشی در زیست‌بوم فناوری کشور است. این چارچوب با برقراری تعادل میان انضباط مهندسی و حاکمیت سازمانی از یک‌سو، و چابکی و تحویل ارزش تدریجی از سوی دیگر، پروژه را از مرحله ایده تا پایش بهره‌برداری پیوند می‌دهد. سازوکارهای پشتیبان، معماری شفاف مستندات، ردیابی‌پذیری دوطرفه و حاکمیت لایه‌بندی‌شده، متضمن شفافیت تصمیم‌گیری و کاهش دوباره‌کاری‌ها در طول چرخه حیات سامانه خواهند بود.

۲. اصول کلیدی حاکم بر پیاده‌سازی

توافق آغازین و نقش مشاور ذیصلاح: پیش از آغاز پروژه، کارفرما و مجری باید بر سر دامنه، ارزش موردانتظار، سطح تعهدات، ساختار تصمیم‌گیری، شیوه تحویل و معیارهای پذیرش به توافق روشن برسند. در پروژه‌های پیچیده، پرریسک یا دارای چند سودبر، بهره‌گیری از مشاور مستقل و دارای صلاحیت برای شفاف‌سازی این تعهدات و تطبیق بهینه چارچوب با واقعیت‌های پروژه ضروری است.

انعطاف آکاردئونی بدون تضعیف انضباط مهندسی: متناسب‌سازی سطح اسناد، آزمون‌ها و کنترل‌ها بر پایه پیچیدگی، ریسک و حساسیت پروژه انجام می‌شود؛ اما این انعطاف هرگز نباید به حذف ارکان حیاتی نظیر تحلیل کسب‌وکار، معماری، مدیریت تغییر، کنترل کیفیت و پذیرش رسمی بینجامد.

تفکیک تعهدات متعارف از اقلام تکمیلی: در سطح تعهد پایه، تجمیع شواهد در اسناد مادر و ابزارهای مدیریت کار کفایت می‌کند. افزودن اسناد تفصیلی، نقش‌های مستقل یا آزمون‌های تخصصی، صرفاً باید با توجیه فنی و توافق شفاف در آغاز پروژه صورت گیرد.

همسویی با استانداردها بدون ادعای انطباق صوری: مراجع و استانداردهای معتبر بین‌المللی، منابع الهام و راهنمای این چارچوب‌اند و بدون ایجاد تضاد با اصول استانداردها، ساختاری بومی و قابل‌اجرا برای پروژه‌ها فراهم شده است.

سخن پایانی

هدف نهایی این چارچوب، تحمیل ساختاری صلب و دیوان سالارانه نیست؛ بلکه خلق زبانی مشترک و نظامی تنظیم‌پذیر است تا تولید نرم‌افزار سفارشی در بستری شفاف، پیش‌بینی‌پذیر و ارزش‌آفرین به ثمر برسد.

با آرزوی عشق‌ورزیدن، شاد زیستن و شادکردن دیگران، ارزشی زیستن و یادگاری با ارزش از خود به جا گذاشتن.
با آرزوی صبوری و پایداری برای تحقق صلح، آزادی و رفاه اجتماعی همه هم‌میهنان کشور عزیزمان ایران.

تابستان سال ۱۴۰۵

اعضای هیئت مدیره انجمن

- آقای ابراهیم نقیب‌زاده مشایخ (رئیس هیئت مدیره)
- آقای دکتر اسلام ناظمی (نایب رئیس هیئت مدیره)
- آقای مهندس محمدحسن محوری (خزانه‌دار)
- آقای مهندس علی آذرکار (دبیر)
- آقای سید ابراهیم ابطی (عضو اصلی)
- آقای دکتر محمد گنج تابش (عضو علی‌البدل)
- آقای دکتر شمس‌اله قنبری (عضو علی‌البدل)

سرپرستان گروه‌های تخصصی انجمن

- ۱- گروه تخصصی مدیریت سرمایه‌های فکری - آقای سید ابراهیم ابطی
- ۲- گروه تخصصی امنیت سایبری - خانم مهشید دولت مدلی
- ۳- گروه تخصصی معماری سازمانی - آقای مهندس رضا کرمی
- ۴- گروه تخصصی محاسبات و سامانه‌های توزیع شده - آقای دکتر شمس‌اله قنبری
- ۵- گروه تخصصی زنجیره‌های بلوکی - آقای دکتر شمس‌اله قنبری
- ۶- گروه تخصصی برنامه ریزی منابع سازمان - آقای مهندس سعید امامی

◀ به‌کارگیری مسئولانه هوش مصنوعی: ابزارهای هوش مصنوعی توانمندساز کمکی در تحلیل، تولید و ارزیابی هستند، اما مسئولیت نهایی صحت، امنیت، کیفیت و انطباق داده‌ها تماماً بر عهده متخصصان انسانی است.

۳. نقشه راه و اقدامات توسعه‌ای چارچوب

تکامل و بلوغ این چارچوب مستلزم استقرار در پروژه‌های واقعی، دریافت بازخوردهای میدانی و مشارکت صاحب‌نظران است. باتوجه به هدف تدوین یک مرجع راهنمای کاربردی و جامع برای صنعت نرم‌افزار، بخش‌های مختلف این چارچوب به‌مرور از طریق پیوست‌های تخصصی، مقالات تشریحی و راهنماهای عملیاتی بسط داده خواهد شد. اقدامات توسعه‌ای آتی در محورهای زیر دنبال می‌شود:

◀ تدوین پیوست‌ها، الگوها و ابزارهای کاربردی: آماده‌سازی قالب‌های استاندارد (Templates)، بازبینی‌های کنترلی، معیارهای تفصیلی پذیرش در گام‌های پنج‌گانه، الگوهای مستندسازی معماری و ماتریس‌های راهنمای RACI؛

◀ الگوهای سنجش، قرارداد و حاکمیت: تدوین شاخص‌های کلیدی عملکرد (KPI) پروژه و محصول، مفاد و الگوهای قراردادی شفاف و دستورالعمل‌های سطح‌بندی آکاردئونی متناسب با مقیاس و پیچیدگی پروژه‌ها؛

◀ مدیریت دانش، درس‌آموخته‌ها و راهنماهای دامنه‌ای: مستندسازی تجارب پروژه‌های واقعی، تدوین راهنماهای انطباق با فناوری‌های نوین (به‌ویژه هوش مصنوعی و داده‌محوری) و توسعه نسخه‌های تخصصی متناسب با دامنه‌های گوناگون کسب‌وکار؛

◀ انتشار مرجع جامع و یکپارچه: گردآوری و به‌روزرسانی مستمر تمامی بخش‌ها و پیوست‌ها در قالب یک کتاب/مرجع راهنمای الکترونیکی در دسترس برای جامعه مهندسی نرم‌افزار و کارفرمایان؛
◀ سنجش دوره‌ای اثربخشی: ارزیابی مستمر چارچوب در کاهش ابهام، مهار ریسک، تسریع تحویل و ارتقای ارزش واقعی نرم‌افزار برای کارفرما.

دعوت به نقد، هم‌اندیشی و همکاری تخصصی

این چارچوب سندی باز و پویاست؛ از این‌رو، نگارنده از هرگونه نقد نقادانه، پیشنهاد اصلاحی، تجارب میدانی، و همچنین همکاری متخصصان، مدیران، معماران نرم‌افزار و پژوهشگران برای بسط و غنای بخش‌های مختلف این مرجع استقبال می‌کند. صاحب‌نظران علاقه‌مند می‌توانند نظرات، دیدگاه‌ها و پیشنهادها را از طریق درگاه‌ها و راه‌های ارتباطی اعلام‌شده با نگارنده در میان بگذارند.

معماری سازمانی به عنوان توانمندساز راهبردی برای پیاده سازی ERP از پشتیبانی تصمیم گیری تا چابکی مبتنی بر هوش مصنوعی

بتول ذاکری

پست الکترونیکی: bzakeri2010@gmail.com

چکیده

علیرغم سرمایه گذاری های کلان و چندین دهه تجربه در پیاده سازی سیستم های برنامه ریزی منابع سازمانی (ERP)^۱، نرخ شکست این پروژه ها همچنان در سطحی نگران کننده (نزدیک به ۷۰ درصد) باقی مانده است. بررسی عمیق این شکست ها نشان می دهد که عامل اصلی، نه نارسایی های فنی نرم افزار، بلکه به سبب فقدان یک چارچوب همسوکننده راهبردی و فنی در ابتدای مسیر است. این مقاله که به طور خاص برای مدیران ارشد و مشاوران حوزه فناوری اطلاعات تدوین شده، با هدف تبیین نقش انکارناپذیر معماری سازمانی (EA)^۲ در تثبیت موفقیت پروژه های ERP، به صورت سیستماتیک سه سناریوی مشخص را مورد واکاوی تطبیقی قرار می دهد: (۱) معماری به عنوان پیش نیاز قطعی، (۲) توسعه همزمان معماری و پروژه ERP، و (۳) اجرای پروژه بدون پشتوانه معماری. سپس با استناد به یکی از رایج ترین چارچوب های معماری، یعنی TOGAF، اثرگذاری عینی هر یک از فعالیت های آن را بر روی فعالیت های انتخاب، پیگیرندی، یکپارچه سازی و استقرار سیستم های نوین ERP تشریح می نماید. در ادامه، با بررسی روند ظهور ERP های بومی-هوش مصنوعی^۳ و عامل های هوشمند^۴، نشان داده می شود که نه تنها از نیاز به نظم و ترتیب معماری سازمانی کاسته نشده، بلکه به دلیل پیچیدگی های حاکمیت داده و شفافیت مدل های تصمیم گیر، این نیاز امروز به یک الزام حیاتی بدل شده است. مقاله در پایان بر این نکته تأکید می کند که معماری سازمانی، لنگرگاه^۵ ضروری هر پروژه تحول آفرین مبتنی بر ERP است.

کلیدواژه ها: پیاده سازی، ERP، معماری سازمانی، انتخاب سیستم، TOGAF، بدهی فنی^۶، یکپارچه سازی داده، چابکی معماری، سناریوهای تعامل و ساماندهی، حاکمیت مدل.

۱. مقدمه: ضرورت سرمایه گذاری در ERP و ریشه یابی تاریخی شکست ها

سیستم های برنامه ریزی منابع سازمانی (ERP)، امروزه به عنوان ستون فقرات عملیاتی سازمان های متوسط و بزرگ شناخته می شوند. این سیستم ها ریشه در دهه ۱۹۶۰ میلادی و سیستم های برنامه ریزی نیازمندی های مواد (MRP)^۷ دارند که صرفاً به صنایع تولیدی کمک می کردند تا نیازهای مواد خود را زمان بندی کنند. با گذشت زمان و گسترش نیازهای سازمانی، در دهه ۱۹۹۰ میلادی، گول های نرم افزاری مانند SAP و Oracle مفهوم ERP را متحول کردند و وعده یکپارچه سازی تمامی وجوه یک بنگاه اقتصادی - از امور مالی و منابع انسانی تا زنجیره تأمین و تولید - را در یک بسته نرم افزاری واحد به مدیران ارشد دادند [۱]. این وعده «یک نسخه واحد از حقیقت»^۸ چنان جذاب بود که سازمان ها میلیارد ها دلار برای استقرار آن سرمایه گذاری کردند.

با این حال، تاریخچه پیاده سازی این سیستم ها، برخلاف ظاهر درخشان آن، همواره با پدیده ای نگران کننده همراه بوده است: نرخ شکست بالا. سیستم های اولیه ERP عمدتاً دارای معماری های یکپارچه^۹ بودند و سازمان ها را مجبور می کردند که فرایندهای خود را با منطق از پیش تعیین شده نرم افزار تطبیق دهند. این تطبیق سازی اجباری، همراه با بسترهای داده ای پراکنده و ناهمگون^{۱۰}، اغلب به ایجاد جزیره های اطلاعاتی جدید و انباشت بدهی فنی منجر می شد.

7- Material Requirements Planning

8- Single Source of Truth

9- Monolithic

10- Legacy Systems

1- Enterprise Resource Planning

2- Enterprise Architecture

3- AI-Native

4- Agentic AI

5- Anchor

6- Technical Debt

گرفته است. در حالی که ERP بر یکپارچه‌سازی برنامه‌های کاربردی و داده‌ها متمرکز است، معماری سازمانی بر برنامه‌ریزی و حاکمیت متمرکز است چارچوبی برای هم‌ترازی قابلیت‌های فناوری اطلاعات با راهبرد کسب‌وکار فراهم می‌کند. همان‌طور که یک توصیف بیان می‌کند: «معمار سازمانی، سازنده یا مهندس ERP نیست، بلکه برنامه‌ریز، آرشیکوننده، و نقشه‌برداری است که تصویری کلی از معماری ERP، فراتر از مرزهای SAP، طراحی، نگهداری و مدیریت می‌کند.» [۳].

● چارچوب‌های EA مانند زکمن (۱۹۸۷) و توگف (۱۹۹۵) رویکردهای ساختاریافته‌ای را برای مستندسازی و مدیریت روابط پیچیده بین راهبرد کسب‌وکار، سیستم‌های اطلاعاتی، و زیرساخت فناوری فراهم کرده‌اند. این چارچوب‌ها از این شناخت پدید آمده‌اند که سرمایه‌گذاری‌های منفرد فناوری اطلاعات -از جمله سیستم‌های ERP- منازری تکه‌تکه و ناهماهنگ ایجاد می‌کنند که قادر به ارائه ارزش راهبردی نیستند.

جدایی دو رشته: مسیرهای واگرا

در طول دهه‌های ۱۹۹۰ و ۲۰۰۰ میلادی، EA و ERP در مسیرهای جداگانه‌ای تکامل یافتند:

رشته‌ی EA	رشته‌ی ERP	بُعد
هم‌ترازی راهبردی و حاکمیت معماری	یکپارچه‌سازی برنامه‌های کاربردی و استانداردسازی فرایندها	تمرکز اصلی
بلندمدت (نقشه‌های راه ۳ تا ۵ ساله)	مبتنی بر پروژه (چرخه‌های پیاده‌سازی)	افق زمانی
مدیران ارشد، راهبرد سازمانی	عملیات فناوری اطلاعات، واحدهای وظیفه‌ای	سودبران
ارزش کسب‌وکار، چابکی، هم‌ترازی	راه‌اندازی، بودجه، جدول زمانی	معیارهای موفقیت
نقشه‌های قابلیت، معماری‌های مرجع	پیکربندی‌های سیستم، نقشه‌های فرایند	مهم‌ترین دستاوردها

این جدایی، شکافی پایدار ایجاد کرد: سازمان‌ها سیستم‌های ERP را بدون درک روشنی از جایگاه آن سیستم‌ها در معماری سازمانی گسترده‌تر خود، انتخاب و پیاده‌سازی می‌کردند. نتیجه، گسترش سیستم‌های «دودکش‌وار»^{۱۴}، داده‌های زائد، و معماری‌های فناوری اطلاعاتی بود که در برابر تغییر مقاومت می‌کردند.

● هم‌گرایی: چرا یکپارچه‌سازی دو رشته ضروری شد؟

چندین عامل باعث هم‌گرایی EA و ERP شد:

اول، این تشخیص که شکست‌های پیاده‌سازی ERP اغلب شکست‌های از جنس معماری هستند. همان‌طور که در ادبیات موضوع اشاره شده

گزارش‌های معتبر مؤسسه گارتنر نشان می‌دهد که تا سال ۲۰۲۷، از هر ده پروژه پیاده‌سازی ERP، بیش از هفت پروژه یا کاملاً با شکست مواجه می‌شوند و یا دستاوردهای کسب‌وکاری مورد انتظار را محقق نمی‌سازند [۲].

علت اصلی این شکست‌ها چیست؟ آیا نرم‌افزارها ناقص هستند؟ خیر. بررسی‌های عمیق‌تر نشان داده است که ریشه اصلی، ناتوانی سازمان در هم‌ترازی و هم‌سویی^{۱۱} راهبردی و فنی قبل و در حین اجرای پروژه است. به عبارت دیگر، سازمان‌ها اغلب یک راهکار نرم‌افزاری پیش‌فرض را بر روی یک بستر داده‌ای آشوب‌زده و بدون نقشه راه یکپارچه‌سازی پیاده می‌کنند و بعداً از هزینه‌های سرسام‌آور اصلاح و اتصال^{۱۲} به سیستم‌های پیرامونی موجود شگفت‌زده می‌شوند [۳]. مسئله محوری این مقاله که به‌طور خاص برای مشاوران و مدیران ارشد حوزه فناوری اطلاعات تدوین شده، دقیقاً به همین نقطه اشاره دارد: علیرغم تأیید ضرورت انکارناپذیر سرمایه‌گذاری در ERP برای بقا و رقابت‌پذیری، چگونه می‌توان از بروز این شکست‌های پرهزینه جلوگیری کرد و موفقیت را نهادینه ساخت؟ پاسخ این مقاله، تأکید بر نقش معماری سازمانی (EA) به‌عنوان رشته‌ای است که با ترسیم نقشه هم‌سویی و هم‌ترازی کسب‌وکار و فناوری، به‌عنوان یک «لنگرگاه» از انحراف پروژه جلوگیری می‌کند. اما پرسش اساسی اینجاست که این معماری باید در چه زمانی و با چه رویکردی وارد پروژه ERP شود؟ پیش‌نیاز آن باشد؟ هم‌زمان با آن پیش برود؟ یا اصلاً نیازی به آن نیست؟ این مقاله به دنبال پاسخ به این پرسش‌های بنیادین با رویکردی تحلیلی و مبتنی بر سناریوهای عینی است.

۲. پیش‌فرض‌ها و تعاریف

پیش از ورود به سناریوهای عملیاتی، لازم است تأکید شود که این مقاله با فرض آشنایی کافی مخاطب با مفاهیم بنیادین هر دو رشته معماری سازمانی و سیستم‌های ERP نوشته شده است. لذا از بسط مباحث نظری پایه‌ای خودداری می‌شود. با این حال، به‌عنوان یادآوری مختصر:

● سیستم برنامه‌ریزی منابع سازمانی (ERP)، مجموعه‌ای از ماژول‌های نرم‌افزاری یکپارچه است که عملیات روزمره سازمان را در حوزه‌های مالی، عملیاتی، و انسانی پوشش می‌دهد و جریان داده‌ها را بین واحدهای مختلف برقرار می‌سازد.

● معماری سازمانی (EA)، رشته‌ای است که با ارائه نقشه‌های جامع از لایه‌های کسب‌وکار، داده، برنامه‌های کاربردی، و فناوری، به سازمان در برنامه‌ریزی، طراحی، و اجرای تحولات دیجیتال بر اساس اصول و چارچوب‌های مشخص (مانند TOGAF)^{۱۳} کمک می‌کند. معماری سازمانی به‌عنوان یک نظم و ترتیب رسمی، از ریشه متفاوتی سرچشمه

11- Alignment

12- Integration

13- The Open Group Architecture Framework

14- Stovepipe



نمودار ۱: گام‌های سناریوی EA به عنوان پیش نیاز

پنهان یکپارچه سازی به حداقل می‌رسد. برای مثال، اگر معماری داده مشخص کرده باشد که شناسه مشتری^{۲۱} باید به صورت ترکیبی از کد ملی و شماره شناسایی منحصر به فرد باشد، تیم پیاده سازی پیش از شروع فاز مهاجرت^{۲۲}، داده‌ها را پاکسازی می‌کند و به این ترتیب از بروز دوبارگی اطلاعات در سیستم‌های هوش تجاری (BI) جلوگیری می‌شود [۴]. مثال دیگر: اگر معماری فناوری، استفاده از بستر ابری خاصی را الزام کرده باشد، فروشنده‌ای انتخاب می‌شود که بالاترین سطح سازگاری را با آن بستر داشته باشد و سازمان مجبور به مهاجرت مجدد زیرساخت در میانه پروژه نمی‌شود.

● چالش‌ها: نیاز به سرمایه گذاری سنگین زمانی و مالی در مرحله مطالعاتی دارد و ممکن است در محیط‌های فوق پویا، سرعت واکنش به تغییرات بازار را کاهش دهد. به عنوان مثال، اگر رقیب اصلی سازمان در مدت زمان معماری سازمانی (که ممکن است ۶ تا ۱۲ ماه طول بکشد) یک ماژول جدید به بازار عرضه کند، سازمان ممکن است دیرتر از رقیب حرکت کند.

● دستاورد نهایی: نرخ موفقیت پروژه‌ها در این سناریو بالای ۸۰ درصد گزارش شده و بدهی فنی در طول عمر سیستم به شدت کاهش می‌یابد [۵]. همچنین هزینه‌های نگهداری سالانه تا ۴۰ درصد نسبت به سناریوی سوم کاهش نشان می‌دهد.

۳-۲ سناریوی میانه: توسعه هم زمان معماری و ERP^{۲۳}

این سناریو برای سازمان‌هایی مناسب است که بلوغ معماری پایینی دارند اما فشار رقابتی اجازه توقف پروژه ERP را نمی‌دهد. اما سؤال اصلی در اینجا این است که توسعه هم زمان واقعاً چه سودی برای پروژه ERP دارد؟

است: «سیستم‌های ERP برای یکپارچه سازی فرایندها طراحی شده‌اند، اما به درک یکپارچه‌ای از آنچه قرار است آن فرایندها به دست آورند، چگونگی ارتباط آن‌ها با یکدیگر، و چگونگی تکامل هماهنگ سیستم‌های پشتیبان نیاز دارند.» [۶].

دوم، پیچیدگی روز افزون چشم اندازهای فناوری اطلاعات. سازمان‌ها ده‌ها و گاهی صدها سیستم را انباشته کردند که نیاز به یکپارچه سازی با هسته ERP داشتند. بدون EA، یکپارچه سازی به صورت موردی و شکننده انجام می‌شد.

سوم، حرکت به سوی معماری‌های ابر^{۱۵} و ترکیب پذیر^{۱۶}، تفکر معماری را حتی ضروری تر ساخت. با کنار رفتن ERP یکپارچه و جایگزینی آن با بوم سازگان‌های مبتنی بر API و واحد مند، نیاز به هماهنگی معماری به امری حیاتی تبدیل شد.

امروزه، اجماع نظر روشن است: EA و ERP رشته‌هایی مکمل هستند که باید در هماهنگی با یکدیگر کار کنند. «اگرچه EA و سیستم‌های ERP مفاهیم متفاوتی دارند، اما دارای ویژگی‌های مکملی هستند که از تصمیم گیرندگان در فرایند انتخاب ERP پشتیبانی می‌کند» EA «داربست»^{۱۷} تحول دیجیتال را فراهم می‌کند: خود سیستم را نمی‌سازد، اما آن را در کنار هم نگه می‌دارد و در عین حال حرکت، تکرار، و رشد را امکان پذیر می‌سازد» [۱۱ و ۵].

۳. سناریوهای سه گانه مواجهه سازمان با معماری و پروژه ERP

نحوه رویارویی سازمان با مقوله معماری، مستقیماً سرنوشت پروژه ERP را تعیین می‌کند. برای درک عمیق تر این تأثیر، در اینجا سه سناریوی عینی بررسی می‌شود. بدیهی است که بهترین و ایمن ترین مسیر، سناریوی نخست است، اما به دلیل فراگیر شدن سناریوی سوم «سناریوی فاجعه بار» در عمل، بررسی هم زمان آن‌ها ضروری است.

۳-۱ سناریوی طلایی: معماری سازمانی به عنوان پیش نیاز قطعی^{۱۸}

در این الگو، سازمان پیش از هرگونه اقدام برای خرید یا پیکربندی ERP، یک برنامه جامع معماری را به سرانجام رسانده یا می‌رساند. چارچوب معماری هدف^{۱۹} ترسیم شده و مشخص می‌شود که سیستم(های) جدید چه جایگاهی در چشم انداز^{۲۰} پنج ساله سازمان دارد.

● دستاوردهای عمیق: انتخاب سیستم بر اساس معیارهای عینی و معماری محور مانند «سازگاری با مدل داده مرجع سازمان» و «قابلیت یکپارچه سازی از طریق API های استاندارد» انجام می‌شود. هزینه‌های

15- Cloud

16- Composable Architecture

17- Scaffolding

18- EA-First

19- Target Architecture

20- Vision

21- Customer ID

22- Migration

23- ERP-EA-Concurrent



نمودار ۳: گام های سناریوی اجرای پروژه بدون معماری (انباشت بدهی فنی)

معمار در میانه پروژه از سازمان خارج شود و جایگزین مناسبی نداشته باشد، تیم توسعه ممکن است به صورت موردی تصمیماتی بگیرد که با معماری های قبلی ناسازگار است و این یعنی بازگشت به سناریوی سوم.

● دستاورد نهایی: این سناریو یک راه حل چابک^{۲۷} است که به سازمان اجازه می دهد هم حرکت کند و هم از پرتگاه سقوط بدون نقشه فاصله بگیرد. برای سازمان های دانش بنیان با توانایی جذب ریسک فنی متوسط، گزینه ای ایده آل محسوب می شود، مشروط بر این که معمار ارشد دارای اختیار «تو» بر روی تصمیمات غیرمعماری باشد.

۳-۳ سناریوی فاجعه بار: پیاده سازی ERP بدون دخالت معماری

متأسفانه این رایج ترین سناریو در عمل است. سازمان تحت تأثیر فشار فروشندگان یا نیازهای شدید عملیاتی، اقدام به خرید یک بسته نرم افزاری می کند بدون این که تصویری از معماری کلان خود داشته باشد.

● دستاوردهای ظاهری: شروع سریع پروژه و صرفه جویی در هزینه های مشاوره اولیه (که در واقع یک خطای راهبردی محسوب می شود). به عنوان مثال، یک سازمان تولیدی ممکن است فکر کند با خرید یک بسته استاندارد ERP و نصب آن در عرض سه ماه، به هدف خود رسیده است.

● چالش ها و مشکلات عینی: برای مثال، در این سناریو، تیم پروژه در میانه راه متوجه می شود که ساختار کد محصول^{۲۸} در ERP جدید



نمودار ۲: گام های سناریوی توسعه هم زمان معماری و (ERP) ارائه معماری به هنگام

ارزش اصلی این سناریو در «پیشگیری از فلج تحلیلی»^{۲۴} است. سازمان مجبور نیست مدت طولانی منتظر تکمیل یک معماری کامل و بی نقص بماند تا پروژه را شروع کند. در عوض، با تعریف یک «هسته معماری حداقلی»^{۲۵} شامل اصول کلان مانند «همه سرویس ها باید مبتنی بر REST API باشند» و مدل داده سطح بالا، پروژه آغاز می شود. سپس تیم معماری به صورت موازی و دقیقاً هم گام با تحویل هر ماژول، جزئیات لایه های یکپارچه سازی و فناوری را پیش بینی و مستند می کند.

● دستاوردها و سود آشکار برای پروژه: این رویکرد دو مزیت حیاتی دارد: اول، کاهش هزینه اصلاحات دیر هنگام. در سناریوی سوم (بدون معماری)، خطاهای یکپارچه سازی در انتهای پروژه کشف می شوند و هزینه رفع آن ها سرسام آور است؛ اما در این سناریو، با تحویل تدریجی ماژول ها، خطاها زودتر کشف و در چرخه بعدی اصلاح می شوند. به عنوان مثال، اگر در مرحله اول (مالی) اشتباهی در ساختار کد کردن حساب ها رخ دهد، در مرحله دوم (منابع انسانی) اصلاح می شود و هزینه آن یک دهم هزینه ای است که باید در انتهای پروژه پرداخت شود. دوم، ایجاد «خطوط قرمز» برای تیم پیاده ساز. اگر چه جزئیات کامل مشخص نیست، اما تیم می داند که نباید از اصول کلان (مانند یکپارچه سازی مبتنی بر پیام رسان) تخطی کند. این امر از بروز واسطه های اسپاگتی مانند^{۲۶} جلوگیری می کند [۶].

● چالش ها: وابستگی شدید به تیم های چند رشته ای و نیاز به تعامل روزانه معمار با تیم توسعه دارد. اگر تیم معماری از پروژه عقب بماند، تصمیمات فنی بداهه جایگزین معماری می شوند. برای مثال، اگر

24- Analysis Paralysis

25- Minimum Viable Architecture

26- Spaghetti Integration

27- Agile

28- Product Code

به صورت خودکار سطح موجودی انبار را بر اساس پیش‌بینی فروش تنظیم می‌کند (سیستمی شبیه به آنچه در فروشگاه‌های زنجیره‌ای بزرگ دیده می‌شود). این سیستم برای پیش‌بینی دقیق، به داده‌های فروش گذشته، داده‌های آب‌وهوا، داده‌های تعطیلات رسمی، و حتی داده‌های شبکه‌های اجتماعی نیاز دارد. حال اگر معماری سازمانی از قبل مشخص نکرده باشد که:

● کدام داده‌ها باید به موتور AI تغذیه شوند (مثلاً داده‌های فروش ۵ سال گذشته یا فقط ۲ سال اخیر؟)

● نرخ به‌روزرسانی این داده‌ها چقدر باشد (روزانه، ساعتی، یا لحظه‌ای؟)

● مالکیت هر کدام از این داده‌ها با کدام واحد سازمانی است (مالی، فروش، یا انبار؟)

سیستم هوشمند با داده‌های اشتباه تغذیه می‌شود و به جای کاهش موجودی، به‌طور ناگهانی سفارش خرید صادر می‌کند یا برعکس، موجودی را به صفر می‌رساند و سازمان با بحران کمبود کالا مواجه می‌شود. اینجاست که معماری سازمانی با فاز C خود (معماری داده)، این مسیر را هدایت و از آشوب جلوگیری می‌کند.

ERP‌های هوشمند عمدتاً بر روی دو لایه تمرکز دارند: لایه تصمیم‌گیری مبتنی بر داده و لایه اتوماسیون گردش کار با عامل‌های هوشمند ۳۲ [۱۲]. با این حال، سه چالش اساسی وجود دارند که بدون معماری سازمانی قابل حل نیستند و اگر سازمان در سناریوی اول یا دوم نباشد، با آن‌ها مواجه خواهد شد:

۱. معماری داده توزیع‌شده و تازگی داده‌ها^{۳۳}: مدل‌های یادگیری ماشین به داده‌های تمیز، با برچسب‌گذاری صحیح و نرخ به‌روزرسانی مشخص وابسته‌اند. معماری سازمانی در مرحله C (معماری داده) تعیین می‌کند که کدام داده‌ها باید به موتور AI تغذیه شوند و نرخ تازه‌سازی آن‌ها چقدر باشد. بدون EA، مدل‌های هوش مصنوعی با داده‌های «آلوده» تغذیه می‌شوند و خروجی‌های مخرب تولید می‌کنند که تشخیص ریشه آن‌ها تقریباً غیرممکن است (مشکل توضیح‌پذیری^{۳۴}). به عنوان مثال، اگر داده‌های فروش، شامل اطلاعات مربوط به دوران کرونا باشد که الگوی فروش غیرعادی داشته، مدل هوش مصنوعی ممکن است همیشه بیش از حد موجودی نگه دارد و سرمایه سازمان را بلوکه کند.

۲. حاکمیت و شفافیت مدل^{۳۵}: هنگامی که ERP به صورت خودکار پیشنهاد تأمین‌کننده یا هشدار ریسک اعتباری صادر می‌کند، باید مشخص باشد که این تصمیم بر اساس چه قاعده یا وزنی گرفته شده است. معماری سازمانی چارچوبی برای حسابرسی این تصمیمات، ثبت دلایل آن‌ها (برای ممیزی‌های مالی)، و تعیین سطح اختیار مدل در کنار تصمیم‌گیرنده انسانی فراهم می‌آورد. این امر مستقیماً در مرحله G (حاکمیت) و مرحله B (قوانین کسب‌وکار) توگف تعریف

با سیستم لجستیک قدیمی همخوانی ندارد و مجبور به نوشتن ده‌ها واسطه سفارشی می‌شود. همچنین فقدان نقشه یکپارچه‌سازی^{۳۶} منجر به ایجاد جزیره‌های داده‌ای^{۳۷} می‌شود که همراستایی با هدف اولیه ERP یعنی «یکپارچگی» را نقض می‌کند [۷]. در یک مثال واقعی، شرکتی که بدون معماری اقدام به پیاده‌سازی SAP کرده بود، پس از دو سال متوجه شد که داده‌های مالی و انبارداری آن‌ها با یکدیگر هماهنگ نیست و برای رفع این مشکل، مجبور شد بیش از بودجه اولیه پروژه هزینه کند [۷].

● دستاورد نهایی: پروژه معمولاً با عبور از بودجه، تأخیر زیاد، و نارضایتی شدید کاربران مواجه می‌شود. در بلندمدت، بدهی فنی عظیمی بر جای می‌گذارد که مهاجرت به نسل بعدی سیستم‌ها (مانند ابری یا هوشمند) را با مانع جدی و هزینه‌های سرسام‌آور مواجه می‌سازد [۸].

۴. مثال عملی: اثرگذاری یک چارچوب رایج معماری (TOGAF) بر گزینش و استقرار ERP

برای آن که استدلال مقاله صرفاً نظری باقی نماند، در این بخش به عنوان یک نمونه عینی و متداول، چارچوب TOGAF انتخاب شده و نشان داده می‌شود که فعالیت‌های روش توسعه معماری (ADM) در این چارچوب، دقیقاً چه تأثیر عینی بر روی مراحل مختلف پروژه ERP می‌گذارند. تأکید می‌شود که توگف در اینجا به عنوان یک الگو انتخاب شده است و سازمان‌ها می‌توانند از هر چارچوب معماری دیگری مانند Zachman یا مدل انتخابی خود استفاده کنند. بر اساس جدول فوق، مشخص می‌شود که در سناریوی اول (EA پیش‌نیاز)، کلیه این مراحل پیش از عقد قرارداد انجام می‌شوند و قرارداد بر اساس یک درک متقابل از «معماری هدف» بسته می‌شود که از ادعاهای بی‌پشتوانه فروشنده جلوگیری می‌کند. در سناریوی دوم (هم‌زمان)، مراحل B تا D ممکن است به صورت مارپیچی و تکراری انجام شوند و در سناریوی سوم، هیچ‌یک از این مراحل به صورت نظام‌مند پیاده نمی‌شوند [۱۱].

۵. عصر ERP‌های بومی-هوش مصنوعی (AI-Native) و ضرورت مضاعف معماری سازمانی

با ورود هوش مصنوعی به هسته سیستم‌های نسل جدید ERP (همانند SAP S/4HANA با قابلیت‌های Joule و Oracle Fusion با دستیارهای هوشمند)، این پرسش حیاتی مطرح می‌شود که آیا این سیستم‌های خود-یادگیرنده^{۳۸}، نیاز به مداخله معمار سازمانی را مرتفع ساخته‌اند؟ پاسخ قاطع این مقاله، «خیر» است؛ بلکه برعکس، معماری در این فضا به یک ضرورت حیاتی‌تر تبدیل می‌شود. برای درک این موضوع، یک مثال ساده بزنیم. فرض کنید یک سازمان تولیدی، ERP هوشمندی را پیاده‌سازی کرده است که

32- Agentic AI

33- Distributed Data Architecture and Data Freshness

34- Explainability

35- Model Governance and Transparency

29- Integration Blueprint

30- Data Silos

31- Self-Learning

جدول ۱: اثرگذاری عینی مراحل TOGAF (به عنوان یک چارچوب رایج) بر روی پیاده سازی ERP

مرحله TOGAF	فعالیت های کلیدی معماری	دستاوردهای قابل لمس برای تیم پیاده سازی ERP	مثال عینی از تأثیر بر پروژه ERP
مقدماتی	تدوین اصول معماری و منشور حاکمیت	تعیین محدوده اختیارات تیم انتخاب و ممنوعیت سفارشی سازی های بی رویه	جلوگیری از تغییر هسته کد (Core Modification) توسط فروشنده برای رفع یک نیاز جزئی. به عنوان مثال، اگر فروشنده پیشنهاد تغییر جدول اصلی پایگاه داده را برای پاسخ به یک نیاز انبارداری بدهد، تیم معماری با استناد به اصل «پایداری هسته» از این کار جلوگیری کرده و راه حل جایگزین (مانند افزودن یک جدول جانبی) را پیشنهاد می دهد.
چشم انداز (A)	ترسیم بیانیه ارزش و قابلیت های هدف	تبدیل نیازهای مبهم مدیران به معیارهای وزنی مشخص برای امتیازدهی به فروشندگان	اختصاص وزن ۳۰ درصد به «سازگاری با معماری داده مرجع» در ماتریس انتخاب. برای مثال، اگر نیاز مدیران «افزایش سرعت گزارش دهی» باشد، این نیاز به معیار «پشتیبانی از پایگاه داده های in-memory مانند HANA» تبدیل می شود و فروشنده ای که این قابلیت را داشته باشد، امتیاز بیشتری کسب می کند.
کسب و کار (B)	نقشه برداری دقیق فرایندهای وضع موجود و مطلوب	تعیین محدوده پروژه (Scope) و اولویت بندی ماژول ها بر اساس تأثیر راهبردی	تشخیص این که ماژول مدیریت زنجیره تأمین باید در اولویت باشد، نه مالی، زیرا با نقشه برداری فرایندها مشخص می شود که گلوگاه اصلی سازمان در تأمین مواد اولیه است و اصلاح آن بیشترین تأثیر را بر سودآوری دارد.
داده (C) و کاربرد	مدل سازی موجودیت های اصلی و بردارهای مالکیت داده	تعیین راهبرد مهاجرت داده Extract- (Transform)- Load و پاکسازی پیش از بارگذاری	مشخص کردن این که حوزه «کد ملی» در سیستم قدیمی، باید به حوزه «شناسه یکتای مشتری» در ERP جدید نگاشت (Mapping) شود و گرنه هوش تجاری دچار خطا می شود [۹]. مثال دیگر: تعیین این که داده های مشتریان قدیمی که بیش از ۱۰ سال فعالیت ندارند، به جای انتقال، در یک بانک اطلاعاتی جداگانه آرشیو شوند تا حجم داده های عملیاتی کاهش یابد و سرعت سیستم افزایش پیدا کند.
فناوری (D)	تعیین زیرساخت، امنیت، و الگوی استقرار	انتخاب بین گرین فیلد (پیاده سازی کامل از نو) یا براون فیلد (تبدیل سیستم موجود)	گرین فیلد (Greenfield) به رویکردی گفته می شود که در آن سازمان، سیستم جدید را از صفر و بدون در نظر گرفتن سیستم های قبلی طراحی و پیاده سازی می کند. این رویکرد معمولاً برای سازمان هایی مناسب است که فرایندهای قدیمی آن قدر ناکارآمد هستند که ارزش انتقال ندارند. براون فیلد (Brownfield) رویکردی است که در آن سیستم جدید بر روی سیستم قدیمی سوار می شود و داده های تاریخی و برخی قابلیت های قبلی حفظ می شوند. انتخاب براون فیلد برای حفظ داده های تاریخی ۱۵ ساله، مشروط به بازطراحی لایه یکپارچه سازی انجام می شود. به عنوان مثال، شرکتی با ۲۰ سال سابقه مالی، براون فیلد را انتخاب می کند تا صورتحساب های گذشته را از دست ندهد.
فرصت ها (E)	تحلیل شکاف (Gap) و تعریف معماری انتقالی	تعیین فازبندی پیاده سازی برای کاهش وابستگی های بحرانی	تصمیم به اجرای هم زمان ماژول های مالی و فروش، اما تأخیر در ماژول تولید به دلیل وابستگی به سخت افزار. مثال عینی: اگر کارخانه نیاز به نصب سنسورهای جدید در خط تولید دارد که تحویل آن ۶ ماه طول می کشد، ماژول تولید به مرحله دوم موکول می شود تا پروژه متوقف نشود.
مهاجرت (F)	تدوین برنامه اجرایی با زمان بندی دقیق و بودجه بندی	هماهنگی تیم های داخلی با تیم فروشنده برای تحویل تدریجی ^۱	تعیین پنجره زمانی ۴۸ ساعته برای قطع ارتباط (Cutover) در آخر هفته و برنامه بازگشت به عقب (Rollback Plan) [۱۰]. به عنوان مثال، اگر در حین قطع ارتباط، خطایی در مهاجرت داده های مالی رخ دهد، تیم می تواند ظرف ۲ ساعت به سیستم قبلی بازگردد و عملیات کسب و کار متوقف نشود.
حاکمیت (G)	نظارت بر انطباق در حین ساخت (Compliance Checking)	جلوگیری از انحراف پیکرندگی ها و تضمین این که تغییرات اضطراری، ساختار مرجع را نقض نکنند	ممیزی دوره ای و توقیف موقت یک تغییر در مسیر (Change Request) به دلیل مغایرت با «اصل سادگی» ^۲ . برای مثال، اگر تیم توسعه پیشنهاد دهد که یک روند پیچیده تأیید سه مرحله ای برای صورتحساب های زیر ۲۰ میلیون تومان اضافه شود، تیم معماری با استناد به اصل «سادگی و سرعت» این درخواست را رد کرده و فرایندی تک مرحله ای پیشنهاد می دهد.

1- Incremental Delivery

2- Simplicity Principle

به چه صورتی باشد، بین این عوامل های هوشمند، هرج و مرجی رخ می دهد که به مراتب بدتر از یکپارچه سازی سنتی است [۱۳]. مثال ساده: عامل فروش، سفارش جدیدی ثبت می کند، عامل انبار موجودی را کاهش می دهد، اما عامل پرداخت از کاهش موجودی اطلاعی ندارد و تخفیفی برای مشتری در نظر می گیرد که با موجودی جدید همخوانی ندارد. معماری سازمانی با مشخص کردن توالی و قواعد تعامل بین این عامل ها، از این هرج و مرج جلوگیری می کند. بنابراین، نتیجه گیری کلیدی در این بخش این است که هوش مصنوعی نه تنها نیاز به معماری را حذف نمی کند، بلکه آن را به یک ضرورت غیرقابل انکار برای جلوگیری از «هرج و مرج دیجیتال» تبدیل می سازد. معماری سازمانی در این عصر جدید، نقش «مغز متفکر» را ایفا می کند که نحوه همکاری عامل های هوشمند را برای دستیابی

می شود. به عنوان مثال، در یک سازمان مالی، اگر مدل هوش مصنوعی به صورت خودکار وام را رد کند، باید مشخص باشد که این رد کردن بر اساس چه معیاری (نمره اعتباری، نسبت بدهی به درآمد، یا سابقه پرداخت) بوده است تا در صورت اعتراض مشتری، قابل پیگیری باشد. معماری سازمانی این شفافیت را تضمین می کند.

۳. معماری ترکیب پذیر و عامل های توزیع شده: نسل جدید ERP ها به جای ماژول های یکپارچه، از مجموعه ای از عامل های مستقل تشکیل شده اند که از طریق API با یکدیگر ارتباط برقرار می کنند. اگر معماری سازمانی در مراحل D و E مشخص نکند که این عامل ها چگونه در کنار خدمات سازمانی (مانند احراز هویت، پرداخت، و لجستیک) قرار بگیرند و الگوی ارکستراسیون و هماهنگی^{۳۷} آن ها

36- Agents

37- Maintenance

دارد، اما به کارگیری هر یک از این دو، به مراتب برتر از نادیده گرفتن کامل معماری (سناریوی سوم) است. به یاد داشته باشیم که معماری، لنگرگاه ثبات پروژه‌های تحول‌آفرین در تلاطم تغییرات است.

منابع:

- [1] Davenport, T. H. (1998). Putting the Enterprise into the Enterprise System. *Harvard Business Review*, 76(4), 121-131.
- [2] Gartner. (2026). Predicts 2026: ERP Implementation Success Rates and Mitigation Strategies. Gartner Research Note ID: G00789122.
- [3] Ross, J. W., & Beath, C. M. (2020). Why Some IT Projects Fail and Others Succeed: The Role of Enterprise Architecture. *MIT Sloan CISR Working Paper*, 45(2), 12-28.
- [4] Yeoh, W., & Koronios, A. (2021). Data Migration in ERP Implementations: The Critical Role of Data Architecture. *International Journal of Accounting Information Systems*, 41(2), 100-118.
- [5] Tamm, T., Seddon, P. B., Shanks, G., & Reynolds, P. (2022). How Does Enterprise Architecture Add Value to Organisations? *Communications of the Association for Information Systems*, 49(1), 10-29.
- [6] Toppenberg, G., & Henningsson, S. (2020). The Inevitable Clash: Balancing Agile Development with Enterprise Architecture Governance. *European Journal of Information Systems*, 29(5), 512-531.
- [7] Kappelman, L., & McLean, E. (2019). The State of ERP Integration and the Impact of Architectural Debt. *Journal of Enterprise Information Management*, 32(4), 601-620.
- [8] Boh, W. F., & Yellin, D. (2022). Using Enterprise Architecture Standards in Vendor Selection for Large-Scale Systems. *MIS Quarterly Executive*, 21(1), 45-63.
- [9] The Open Group. (2024). TOGAF® Standard, Version 10.0 – Architecture Development Method (ADM). Van Haren Publishing.
- [10] Lucke, C., & Krcmar, H. (2021). Hybrid EA Approaches: Synchronizing Enterprise Architecture with Agile ERP Roll-outs. *Business & Information Systems Engineering*, 63(4), 421-439.
- [11] Kumar, M. (2026). AI-Native ERP Systems: A Design Science Framework for Intelligent Enterprise Decision Automation. *Journal of Computer Science and Technology Studies*, 8(8), 1-25.
- [12] Bughin, J., & Hazan, E. (2025). Agentic AI in the Enterprise: Architectural Implications and Organizational Readiness. *McKinsey Digital Quarterly*, 4(3), 88-105.

به اهداف کسب‌وکار هماهنگ می‌سازد و در این مسیر، سناریوی اول (پیش‌نیازی) بیشترین امنیت و سناریوی دوم (هم‌زمان) بیشترین انعطاف‌پذیری را فراهم می‌آورد.

۶. نتیجه‌گیری و توصیه‌های عملیاتی به مشاوران و مدیران ارشد

این مقاله با رویکردی عمیق‌تر و مبتنی بر سناریوهای عینی، نشان می‌دهد که موفقیت پروژه‌های ERP به‌شدت تابعی از نحوه مواجهه سازمان با مقوله معماری است. تحلیل تطبیقی سه سناریوی تعامل، این واقعیت را آشکار می‌سازد که:

۱. سناریوی طلایی (EA به‌عنوان پیش‌نیاز)، اگرچه پرهزینه‌ترین گام اولیه است، اما با ترسیم نقشه راه شفاف و معیارهای انتخاب عینی، نرخ موفقیت پروژه را تا بیش از ۸۰ درصد افزایش داده و هزینه‌های پنهان یکپارچه‌سازی و نگهداری را به‌شدت کاهش می‌دهد. این سناریو برای سازمان‌هایی با بلوغ بالا و بودجه کافی، بهترین گزینه است.

۲. سناریوی میانه (توسعه هم‌زمان)، به‌عنوان یک راه‌حل چابک ارزشمند، به سازمان اجازه می‌دهد تا از «فلج‌تحلیلی» رهایی یابد و پروژه را شروع کند، در حالی که یک «هسته معماری حداقلی» و اصول کلان، به‌عنوان خطوط قرمز، از انحراف فاجعه‌بار جلوگیری می‌کند. سود اصلی این سناریو، کاهش چشمگیر هزینه اصلاحات در مقایسه با سناریوی سوم و افزایش سرعت تحویل در مقایسه با سناریوی اول است.

۳. سناریوی فاجعه‌بار (بدون معماری)، تقریباً به‌طور حتم به انباشت بدهی فنی، افزایش هزینه‌های تعمیر و نگهداری، و نهایتاً شکست راهبردی منجر می‌شود. این سناریو به‌ویژه در مواجهه با ERP‌های هوشمند، به دلیل غیرشفاف بودن فرایند تصمیم‌گیری و افزایش وابستگی‌های پنهان، خسارت‌های جبران‌ناپذیری به همراه خواهد داشت.

همچنین با واکاوی یک چارچوب رایج مانند توگف، مشخص شد که هر یک از مراحل آن تأثیری عینی و غیرقابل‌انکار بر روی گزینش، پی‌کربندی، مهاجرت داده، و استقرار سیستم دارد و نمی‌توان آن‌ها را به تشریفات اداری تقلیل داد. این تأثیر در هر سه سناریو به‌صورت متفاوت نمود پیدا می‌کند، اما در سناریوهای اول و دوم به‌صورت نظام‌مند و در سناریوی سوم به‌صورت غایب.

توصیه نهایی به مشاوران و معماران ارشد این است که واحد معماری سازمانی را نه به‌عنوان یک هزینه سربار، بلکه به‌عنوان یک بیمه راهبردی برای سرمایه‌گذاری‌های کلان ERP تلقی کنند. در مدل‌های تأمین مالی پروژه‌ها، بودجه مستقل و قابل‌توجهی برای فعالیت‌های معماری (به‌ویژه در مراحل مقدماتی و حاکمیت) اختصاص دهند و بر ایجاد «هیئت‌های بررسی معماری» با قدرت وتوی واقعی اصرار ورزند. انتخاب بین سناریوی اول و دوم بستگی به بلوغ سازمان و بودجه



هوش مصنوعی در خدمت بشریت

بنا کردن هوش مصنوعی پایدار برای آینده

نویسندگان: دکتر جیمز اوانگ
اندی داما
سیوک سیوک تان

ترجمه: ابراهیم نقیبزاده مشایخ



انجمن انفورماتیک ایران

تهیه کتاب از دفتر

انجمن انفورماتیک ایران

۰۲۱-۶۶۴۱۲۸۶۱

یادداشت‌های یک معمار سازمانی

رضا کرمی

سرپرست گروه تخصصی معماری سازمانی انجمن انفورماتیک ایران

پست الکترونیکی: Karami@golsoft.com

سازمان اجرا و پیگیری شوند تا به تحول و تغییر ملموس سازمان بینجامند. اما در بعضی موارد (در مورد کشور ما باید گفت: در اغلب موارد!)، به دلایلی از قبیل عدم اعتقاد مدیران یا عدم توانایی‌شان در تامین و سازماندهی مناسب منابع لازم، غلبه بر موانع فرهنگی، ایجاد انضباط و رهبری تغییر، بهترین و کامل‌ترین برنامه‌ها هم به نتیجه نمی‌رسد و گاهی از کار فرو بسته سازمان‌ها نمی‌گشاید. به این اشکالات بیفزایید تغییرات مداوم و سریع مدیریتی را که گاه همچون سیلی، از صدر تا ذیل سازمان را درمی‌نوردد و در نسل جدید مدیران هیچ سابقه‌ای از تلاش‌ها و برنامه‌ریزی‌های قبلی به جا نمی‌گذارد.

آری، دیدن این بی‌ثمری و کم‌ثمری تلاش‌های برنامه‌ریزی برای مشاوران تحول سازمانی حقیقتاً دشوار است و به تجربه دیده‌ام بسیاری از مشاوران همکار را که به دلیل همین مشاهدات و تجربیات، در نیمه‌راه مسیر شغلی‌شان، عطای مشاوره را به لقای این دور باطل برنامه‌ریزی‌های بی‌ثمر یا کم‌ثمر می‌بخشند و به سایر عرصه‌های شغلی و حرفه‌ای پناه می‌برند.

واقعیت این است که تجربه ثمربخشی تلاش‌ها و دستاوردها در هر شغلی، از مهم عواملی است که در کنار تامین معاش و کسب منزلت اجتماعی، به آن شغل معنا می‌بخشد و شاغلان آن شغل را ارضا می‌کند. به همین دلیل، نگرش صحیح به نقش و ارزش کار مشاوران در تحول سازمان‌ها، اهمیتی دوچندان می‌یابد.

من گاهی نگاه کردن به این نقش را، از زاویه‌ای که مطرح می‌کنم، مفید یافته‌ام: سازمان‌های ما (اعم از دولتی یا خصوصی) در کنار فلسفه وجودی اصلی‌شان که بر حسب مأموریت و وظایف‌شان تعریف می‌شود، محمل و مجمع یک نوع خاص از ارزشمندترین منابع کشور ما هستند، و آن همان منابع انسانی، یا به تعبیر امروزی‌تر سرمایه

یادداشت‌های یک معمار سازمانی، نوشته‌های کوتاهی است که با نگاه معمارانه و با تأمل در آنچه در اطراف ما می‌گذرد، نوشته شده است. قالب این نوشته‌ها، کمی خارج از قاعده متعارف نوشتار فنی و اداری معمول، با رنگ و بوی قضاوت‌های شخصی و گاه آمیخته با طنز و کنایه است، اما هدف در نهایت آسیب‌شناسی سازمان‌های کشور از منظر معماری سازمانی است، به امید آن‌که شاید گاهی تلنگری، زنجاری یا بیدارباشی در آن‌ها بتوان یافت، و فراخوانی به اصلاح. نشریه گزارش کامپیوتر از نوشته‌های مشابه سایر صاحبان اندیشه استقبال می‌کند.



تحول سازمانی و مسئولیت اجتماعی ما

یکی از دشوارترین جنبه‌های مشاوره تحول سازمانی، آگاهی از بی‌اثر بودن یا کم‌اثر بودن برنامه‌هایی است که مشاور برای توسعه و تحول یک سازمان ارائه کرده است. انتظار می‌رود این برنامه‌ها که معمولاً بعد از طی مراحل طولانی بررسی وضعیت فعلی سازمان، مشکلات و چالش‌های آن، اهداف کسب‌وکاری، ارزیابی میزان آمادگی برای تغییر، طراحی آینده و با بهره‌گیری از چارچوب‌ها و به‌روش‌ها و دانش و تجربه مشاور تدوین و ارائه می‌شوند، توسط مدیران و رهبران

برای آن لفظ وضع کرده‌اند، و هر چه درک و دریافت مردمان از آن معنی یکسان‌تر و نزدیک‌تر باشد، کاربرد آن لفظ با دقت بیشتری همراه است. در علوم دقیق و اثباتی، مانند ریاضی و فیزیک و منطق و مانند آن، روشن گردانیدن الفاظ پیش از کاربریشان، اهمیتی اساسی دارد. در علوم و فنون مهندسی و کاربردی هم توصیه می‌شود تا جایی که می‌توانیم از الفاظ چندپهلوی و نادقیق استفاده نکنیم، و به همین دلیل برای یگانه کردن الفاظ و معانی این علوم و فنون، استانداردها ساخته‌اند و می‌سازند. هر چه از این علوم به سمت علوم انسانی و اجتماعی می‌رویم، این وضوح و دقت بیشتر رنگ می‌بازد و معانی اصطلاحات و دانش‌واژه‌ها بیشتر و بیشتر تابع مقاصد و دیدگاه‌های این و آن می‌شود. به همین دلیل، اختلاف آرا در این حوزه‌های دانشی بیشتر است.

در علوم و مباحث مدیریتی هم، به اعتبار آن که در مرز بین علوم انسانی و مهندسی سیر می‌کنند، اختلاف معانی و وضع و کاربرد الفاظ واحد بر معنای متعدد، پدیده شناخته شده‌ای است. مثلاً مشهور است که به تعداد کتاب‌های درسی در مورد راهبرد، تعریف مختلف از واژه راهبرد داریم. هنری مینتزبرگ، نظریه پرداز معروف راهبرد سازمانی، مقاله معروفی دارد که در آن کوشیده معانی مختلف راهبرد را گردآوری و دسته‌بندی کند. در مورد سایر مفاهیم و دانش‌واژه‌های مدیریتی هم، مانند «تحول»، «رهبری»، «فرآیند»، «کارکرد»، «معماری» و مانند آن هم این تششت وجود دارد. و بر همین زمینه است که می‌توان هر روز نشست و درباره این «چیستان»ها و «نیستان»ها قلم‌فرسایی کرد!

تعریف و بازتعریف مفاهیم مدیریتی خوب است و ممکن است ارزش آموزشی یا حتی عملی هم داشته باشد، اما آیا افراط در این کار پسندیده، خود کاری پسندیده است؟ ویتگنشتاین، فیلسوف معروف قرن بیستم، گفته بسیار نقل شده و تاثیرگذاری دارد به این مضمون که «معنی، همان کاربرد است.» معنی این گفته این است که بر خلاف تصور شایع که هر لفظی را برای افاده معنی واحد و تغییرناپذیری وضع کرده‌اند، معنی یک واژه امری ثابت و بدون تغییر و مستقل از کاربرد نیست. در واقع گویشوران هر زبان هر روز با کاربرد یک واژه در ساختار جملات و عبارات و افادات خود، در ساختن و بازساختن آن واژه سهیم می‌شوند و معنی واژه بسته به این کاربردها به صورت سیال تغییر می‌کند. بدیهی است دامنه این تغییرات از حوزه‌ای به حوزه دیگر ممکن است متفاوت باشد.

بنابراین برای پی بردن به معنی دانش‌واژه‌های مدیریتی هم، با اینکه نباید همان تساهل و تسامح معمول در مورد واژه‌های زبان روزمره را روا داشت، لازم هم نیست همگان به یک توافق دقیق و اجماع همه‌شمول در میان کاربران این واژه‌ها برسند. اطمینان حداقلی از اینکه زمینه و پیامدهای کاربرد یک واژه را در ظرف یک گفتگو یا متن مدیریتی می‌فهمیم و روی آن (حدوداً) توافق داریم، معمولاً کافی است.

انسانی، و به بیان خودمانی‌تر، همان افرادی است که در این سازمان‌ها کار می‌کنند. میلیون‌ها نفر از افراد این سرزمین، زندگی‌شان را از مسیر کار در این سازمان‌ها می‌گذرانند و بهترین و ارزشمندترین سرمایه‌شان، یعنی وقت، دانش، مهارت و ظرفیت‌های فکری و شناختی‌شان را در این سازمان‌ها صرف می‌کنند. هر قدم مؤثری، ولو کوچک، برای بهبود کارایی و اثربخشی این سازمان‌ها، در ضریب بالایی از ساعات و ماه‌ها و سال‌ها زندگی این افراد ضرب می‌شود و کیفیت و معنای زندگی آنان را افزایش می‌دهد. این اثر ممکن است با شاخص‌هایی که معمولاً سازمان‌ها موفقیت خود را با آن‌ها می‌سنجند قابل‌سنجش و گزارش نباشد، اما انسان اهمیت آن را در برخورد روزمره با تک‌تک استعدادهایی که عمرشان را در این سازمان‌ها سپری می‌کنند، لمس می‌کند.

برای دستیابی به چنین دستاوردی، تحمل ناملایمات و ناکامی‌های پی‌درپی، هر چقدر هم که سنگین به نظر برسند، ارزشمند است. چنانکه دانای توس سروده است:

«چو دانا نماید به کاری درنگ

به پیروزی آرد جهانی به چنگ

همه کارها در فروبستگی

گشاید ولیکن به آهستگی»

۳/۴/۱۴۰۵



چیستان‌ها و نیستان‌های مدیریتی

اگر به شبکه‌ها و سکوها حرفه‌ای، مانند لینکدین، نگاهی بیندازید، هر روزه با سیلی از فرسته‌ها (یا همان پست‌ها)ی با این عنوان مواجه می‌شوید که «فلان چیز چیست؟» یا «فلان چیز، بهمان چیز نیست!». در این محیط‌ها، و گاهی در محیط‌های واقعی هم، به نحو مستمر با افرادی مواجه هستیم که می‌خواهند به ما ثابت کنند راهبرد آن چیزی نیست که ما فکر می‌کرده‌ایم، یا معماری سازمانی در واقع این است که نویسند یا گوینده می‌خواهد به ما بقبولاند، یا در مورد تحول دیجیتال، همگان تا کنون در توهمی بیش نبوده‌اند، و از این قبیل.

در عالم اندیشه، هر لفظی بر معنایی دلالت می‌کند که «در اصل»

در علوم مدیریتی نیازی به تعاریف ریاضی‌وار نیست، کافی است حرف هم را بفهمیم!
«اختلاف خلق از نام اوفتاد،

چون به معنی رفت، آرام اوفتاد»

۱۶/۴/۱۴۰۵



مواد لازم برای راهبری معماری سازمانی

در یکی از یادداشت‌های قبلی (دو راهبرد خانه‌داری) دو رویکرد مدیریت بدهی معماری سازمان‌ها را تشریح و مقایسه کردم و رویکرد مبتنی بر راهبری معماری سازمانی یا EA Governance را توصیه کردم. معنی ساده راهبری معماری سازمانی این است که هر گونه تغییر در بلوک‌های سازنده معماری (ساختار، فرآیندها، سیستم‌های اطلاعاتی، سرویس‌ها و ...) بر مبنای اصول و معیارهای از پیش تعریف‌شده‌ای کنترل شده و تنها در صورت تطابق با آن اصول و معیارها تأیید و اعمال گردد. بسیاری از سازمان‌ها با وجود درک و پذیرش ضرورت و کارایی چنین رویکردی، پیاده‌سازی عملی آن را دشوار و پیچیده می‌یابند و معمولاً استقرار چنین نظامی را به آینده دور و پس از دستیابی به سطح بالایی از بلوغ معماری سازمانی موکول می‌کنند. تصمیمی که به نوبه خود باعث افزایش مستمر بدهی معماری و گرفتار شدن بیشتر سازمان در این چرخه معیوب می‌شود با وجود آن‌که همبستگی مثبت بین بلوغ معماری سازمانی و وجود نظام راهبری معماری سازمانی یک واقعیت تجربی است، اما پیاده‌سازی EAG را نباید کاری پیچیده و بغرنج تصور کرد. فراهم آوردن چند مؤلفه اساسی برای استقرار این نظام کافی است:

۱. نهاد راهبری: وجود یک نهاد بین‌واحد یا فراواحدی در سازمان، برای تصمیم‌گیری نهایی در موارد اختلافی (مانند زمانی که یک درخواست تغییر معماری به دلیل مغایرت با اصول و طرح‌های معماری رد می‌شود)، ضروری است. عنوان این نهاد مهم نیست، اما هر چه ترکیب آن جامع‌تر (شامل نمایندگان همه سودبران کلیدی) و سطح آن بالاتر (شامل رده‌های مدیریتی بالای سازمان) باشد، کارایی آن بیشتر خواهد بود.

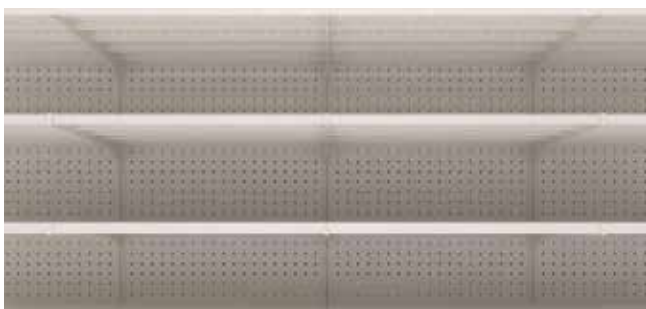
۲. معیارهای راهبری: برای تطابق‌سنجی و کنترل، وجود معیارهای مرجع الزامی است. این معیارها معمولاً شامل اصول از پیش توافق‌شده

معماری و چند نما از معماری مطلوب می‌شود. در مورد طول و تفصیل و سطح جزئیات مدل‌های معماری مطلوب نباید سخت‌گیری کرد. راهبری معماری را می‌توان حتی از یک نمای بسیار کلان (One-page architecture) شروع کرد و به تدریج عمق و گستره آن را افزایش داد.

۳. کنترل‌های معماری: چند نقطه کنترلی در چند فرآیند کلیدی تغییر معماری باید شناسایی شده و گذر از این نقاط به تأیید تطابق با معماری منوط شده باشد. در انتخاب این گلوگاه‌ها هم نباید از ابتدا بیشینه‌خواهی کرد. از یک یا دو نقطه کنترلی، که معمولاً بیشترین بدهی معماری را تولید می‌کنند شروع کنید و بعد از تثبیت آن‌ها و اثبات کارایی و اثربخشی آن‌ها در عمل، تعداد نقاط کنترلی را افزایش دهید.

با همین سه مؤلفه حداقلی می‌توانید یک نظام مؤثر راهبری معماری سازمانی راه بیاندازید.

۵/۵/۱۴۰۵



در ستایش کسادی!

در دنیای کسب‌وکار، همه از کسادی بازار بیزارند و از رونق آن خوشنود. هر که متاعی برای فروش دارد، از زیادت مشتری آن متاع خوشحال می‌شود و به انواع حیل در گرمی بازارش می‌کوشد. اما گاه با موقعیت‌هایی مواجه می‌شویم که رونق، آفات خود را دارد. یکی از این موقعیت‌ها را ما هم‌اکنون در بازار خدمات معماری سازمانی در کشور تجربه می‌کنیم.

در اوایل دهه هشتاد شمسی که اولین موج پروژه‌های معماری سازمانی در کشور شروع شد و بازار این خدمات به‌مدد طرح تکفا رونقی گرفت، شرکت‌ها و مشاوران متعددی به این بازار روی آوردند و به تعریف و اجرای پروژه‌های «طرح جامع فناوری اطلاعات» یا عناوین مشابه برای سازمان‌های (عمدتاً دولتی) کشور پرداختند. در آن زمان، هم تصور و انتظار سازمان‌ها از معماری سازمانی چندان به‌جا و معقول نبود، و هم دانش و تجربه مشاوران در حدی نبود که این انتظارات را منطقی کنند. در سال ۱۳۹۴ که اولین نسخه چارچوب ملی معماری سازمانی کشور در حال تدوین بود، پیمایشی از دست‌اندرکاران پروژه‌های متعدد معماری سازمانی که تا آن تاریخ

اگر از محدودیت‌های بودجه‌ای این سازمان‌ها که در نتیجه اوضاع و احوال اخیر کشور به وجود آمده است بگذریم، می‌توان باز رونق نصف‌ونیمه ولی فزاینده‌ای را در این بازار مشاهده کرد. این وضعیت باید موجبات خوشنودی مشاوران معماری سازمانی را فراهم کند. اما آیا چنین است؟

رونق بازار معماری سازمانی، هنوز به اوج خود نرسیده چند تالی فاسد به همراه آورده است. اولاً به دلیل فزونی تقاضا بر عرضه، بسیاری از سازمان‌های متقاضی را واداشته تا پروژه‌های معماری خود را با نیروی داخلی و یا با ارجاع به مشاوران ناکارآمدی که به‌بوی این رونق نویافته و بدون تجربه و دانش و بینش کافی در این حوزه مدعی شده‌اند، انجام دهند. بیشتر این پروژه‌ها با بهره‌گیری از متون دانشگاهی یا مطالب پراکنده در فضای مجازی تعریف می‌شود و در آن‌ها اثری از پختگی و سنجیدگی نمی‌توان یافت. ثانیاً، انتظارات انباشته و برآورده نشده سازمان‌ها از تحول، به صورت غیرمنطقی در تعریف این پروژه‌ها خود را نشان می‌دهد. متأسفانه سازمان‌های متولی هم که در ایجاد الزامات قانونی دستی گشاده دارند، در راهنمایی و اصلاح نگاه سازمان‌های دولتی چندان کاری از دست‌شان برنمی‌آید. نتیجه این چند عامل، برآمدن موج جدید از پروژه‌های معماری سازمانی خواهد بود که می‌توان انتظار داشت باز بعد از چند سال، به تعبیر معروف آن چرخه هیاهوی مؤسسه گارتنر، از «اوج انتظارات متورم» به «حضیض نومییدی» بیافتد.

۲۵/۵/۱۴۰۵

در کشور انجام شده بود انجام دادیم و از آنان خواستیم مهم‌ترین عوامل شکست پروژه‌های معماری سازمانی را از دیدگاه خودشان فهرست کنند. در انتها، فهرستی از ۱۶ عامل مهم ناکامی این پروژه‌ها تدوین شد که در صدر آن‌ها «انتظار نادرست از معماری سازمانی» قرار داشت.

از نیمه دهه هشتاد، بازار معماری سازمانی در کشور به سردی گرائید و رونق، جای خود را به رکود داد. در این فضا بسیاری از شرکت‌ها و مشاوران، ترجیح دادند از این بازار خارج شوند و به حوزه‌های پروتق‌تر و پرمفعت‌تر دیگر روی آورند. آن‌هایی که ماندند، به تدریج در کار خود ماهرتر شدند و در انتخاب کار، پروسه‌سازتر و گزیده‌تر. از اواخر دهه هشتاد تا همین یکی دو سال پیش، سازمان‌های معدودی پروژه‌های معماری تعریف می‌کردند و به تدریج کاربردها عمیق‌تر و اثرات ماندگارتر می‌شد. تک‌وتوک سازمان‌هایی که در این دوره تجارب معماری سازمانی داشتند، کم‌کم از دیدگاه «پروژه‌ای» به دیدگاه «فرآیندی» و از کاربردهای متمرکز بر فناوری به کاربردهای متوجه به اصلاح و بهبود کسب‌وکار متمایل شدند. نه الزام بالادستی چندانی در کار بود و نه هیجان غیرمنطقی‌ای در تعریف و اجرای پروژه‌ها.

در یک سال اخیر، با تصویب و ابلاغ الزامات قانونی که در آن‌ها به معماری سازمانی اشاره شده است (از جمله ماده ۶ نقشه‌راه دولت هوشمند که در آغاز سال ۱۴۰۴ ابلاغ شد)، تب تعریف و اجرای پروژه‌های معماری سازمانی در میان سازمان‌های دولتی باز بالا گرفت.

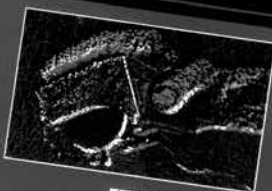
ماجراهای پشت پرده

چگونه نهضت شد فرهنگ دهه شصت
به شکل‌گیری صنعت رایانه‌های شخصی انجامید؟

جان مارکاف

ترجمه:

ابراهیم تقیپ زاده مشایخ



ناشر: انجمن انفورماتیک ایران
بهار ۱۳۸۹

برای کسب اطلاعات بیشتر و تهیه کتاب
با شماره تلفن‌های زیر تماس حاصل فرمایید

۸۸۸۶۱۴۲۱-۳ (انجمن انفورماتیک ایران)

و یا برای خرید اینترنتی به وبگاه زیر مراجعه فرمایید

www.chare.ir

پیشنهاد (قسمت نوزدهم) - از پی آیند خود اظهاری خاطرات ۶۶ پیشکسوت، به بهانه جشن بزرگ انجمن در زمستان ۱۴۰۷

پیش به سوی جشن بزرگ گرامی داشت

پنجاهمین سالگرد تاسیس انجمن انفورماتیک ایران

سیصدمین نوبت انتشار گزارش کامپیوتر

پنجاه و پنجمین سال تاسیس مدرسه عالی برنامه ریزی و کاربرد کامپیوتر

شصت و ششمین سالگرد ورود رایانه به ایران

همراه با روایات کتبی ۶۶ پیشکسوت داوطلب جهت کمک به تدوین تاریخ مکتوب نگاشته شاهدان عینی از وقایع انفورماتیک ایران (پیشکسوت گرانقدر، آنچه در دوران حضور در بخش انفورماتیک کشور، شاهد یا عامل بوده‌اید را با توالی زمانی، برای درج در این صفحه بفرستید)

زمان برگزاری جشن: طی بهمن و اسفند ۱۴۰۷

زندگی‌نامه خود اظهارنامه شاهدان عینی عامل وقایع بخش انفورماتیک ایران

شانزدهمین و هفدهمین پیشکسوت از میهمانان پیشواز جشن پنجاه سالگی انجمن

روایتی از زندگی، فعالیت‌ها و حضور موثر در بخش انفورماتیک ایران پیشکسوتان

دکتر هایده اهرابیان و دکتر عباس نوذری دالینی

از پیشکسوتان دانشگاهی عضو هیئت علمی: استاد و دانشجویی که هر دو مرگ زود هنگامی داشتند



به فراخوان، دعوت و نقل سید ابراهیم ابطاحی

استادیار دانشکده مهندسی کامپیوتر دانشگاه صنعتی شریف

abtahi@sharif. ir



دکتر هایده اهرابیان ۱۳۳۵-۱۳۹۰

خورده است. او نه تنها پژوهشگری توانمند در حوزه الگوریتم‌ها و محاسبات نظری بود، بلکه معلمی دلسوز، مدیری عالم و از معماران نسلی از دانش‌آموختگان علوم کامپیوتر به‌شمار می‌رفت که بعدها خود به استادان و پژوهشگران این حوزه تبدیل شدند. زندگی علمی او روایتی است از پشتکار، دقت، عشق به آموزش و تعهد به ساختن رشته‌ای که در زمانه‌ای نه‌چندان آسان، هنوز در حال یافتن جایگاه خود در دانشگاه‌های ایران بود.

دکتر هایده اهرابیان در ۱۵ خرداد ۱۳۳۵ در تهران متولد شد. دوران تحصیل او با موفقیت‌های چشمگیر همراه بود و استعداد علمی‌اش از



جلسه در دانشکده: دکتر پزشکی، شفیع، اهرابیان و نوذری

همان سال‌های جوانی آشکار شد. پس از پایان تحصیلات دبیرستان، در مهرماه سال ۱۳۵۴ وارد گروه ریاضی و علوم کامپیوتر دانشگاه تهران شدند؛ دورانی که علوم کامپیوتر هنوز به‌عنوان رشته‌ای مستقل و تثبیت شده در ایران شناخته نمی‌شد و بیش از آن که ساختاری آموزشی داشته باشد، مجموعه‌ای نوظهور از ایده‌ها و مفاهیم میان‌رشته‌ای بود.

ایشان در خرداد ماه سال ۱۳۵۸ با کسب رتبه نخست، دوره کارشناسی خود را به پایان رساندند؛ موفقیتی که مسیر علمی آینده ایشان را شکل داد. به‌واسطه بورسیه شاگرد اولی، در همان سال راهی انگلستان شدند تا تحصیلاتشان را در دانشگاه ویکتوریا منچستر ادامه دهند؛ دانشگاهی که در تاریخ علوم کامپیوتر جهان جایگاهی ویژه دارد و میراث دانشمندانی چون آلن تورینگ را در خود حفظ کرده است.

در این دانشگاه، دکتر اهرابیان در فضایی علمی و پیشرفته با مبانی عمیق علوم محاسباتی آشنا شدند و در فروردین ماه سال ۱۳۶۰ از پایان‌نامه کارشناسی ارشد خود با عنوان Algorithms for the Solution of Single Nonlinear Equation

در واپسین شماره سال ۱۴۰۱ در گزارش کامپیوتر ۲۶۴، طرح سالیانم برای تدوین تاریخی شفاهی برای فعالیت‌های نیم‌قرنی بخش انفورماتیک را با فراخوان عمومی به جمع پیشکسوتان این بخش سپردم و چشم‌انداز ۱۴۰۷ را در قالب جشنی به مناسبت‌های فراوان، از جمله انتشار سیصدمین شماره گزارش کامپیوتر را به‌عنوان افق انگیزشی برای به سرانجام رساندن آن ترسیم کردم. پیش از این خاطرات نوزده پیشکسوت، از شاهدان عینی وقایع نیم قرن این بخش را به همت دوستان نشر کرده‌ایم: ۱- محمد بخشی (۲۶۵)، ۲- ابراهیم ابطحی (۲۶۷)، ۳- ابراهیم نقیب‌زاده مشایخ (۲۶۸)، ۴- محمدعلی علاقبند، ۵- رضا محسنی، ۶- افشین همت‌یار (۲۷۰)، ۷- محسن صدیقی مشکنانی (۲۷۱)، ۸- موسی خالصی‌زاده (۲۷۲)، ۹- محمد داورپناه جزی (۲۷۳)، ۱۰- اسلام ناظمی (۲۷۴)، ۱۱- مسعود صفاری (۲۷۵) و ۱۲- پرویز ناصری (۲۷۷)، ۱۳- مهدی عربی (۲۷۸)، ۱۴- قدسی و ۱۵- روحانی رانکوهی (۲۷۹)، ۱۶- محمد میرزا عبداللّهی و ۱۷- یحیی تابش (۲۸۰) و ۱۸- مرتضی انواری و ۱۹- بهروز پرهامی (۲۸۱)، ۲۰- کارو لوکس و ۲۱- کامبیز بدیع (۲۸۲)، ۲۲- کیومرث افشارپاد و ۲۳- علی آذرکار (۲۸۴) و در این شماره در خدمت ۲۴ و ۲۵ امین پیشکسوت، زنده‌یادان دکتر هایده اهرابیان و دکتر عباس نوذری دالینی هستیم که با ۶۶ روایت موردنیاز تا سال موعود فاصله دارد. برای رسیدن به هدفمان تا شماره سیصدم گزارش کامپیوتر نیاز داریم حداقل در ۴ شماره بعدی هر شماره دو پیشکسوت و در ۱۱ شماره پایانی در هر شماره سه پیشکسوت را ملاقات کنیم و این به همت پیشکسوتان برمی‌گردد که در این راه از ارسال خاطرات و نظراتشان دریغ نفرمایند.

مقدمه

دو پیشکسوتی که در این شماره روایت حضورشان در بخش انفورماتیک را می‌خوانید دکترهایده اهرابیان و دکتر عباس نوذری دالینی، استاد و دانشجوی (سپس استادی) توانایی از دو نسل متفاوت و متوالی ولی هر دو از تلاش‌گران نوجو و موثر نسل خود هستند. نسل‌هایی که در دوی امدادی معلمی کم‌قرب، در اقلیم مردمان کم‌اعتنا به کوشندگان تعلیم و تربیت از سید حسن رشیده گرفته تا پرویز شهریاری، کماکان رهروانی کوشنده در راه، که دنیا هم به آن دو وفا نکرد و درگذشت نابهنگامشان داغی شد بر دل همه.

دکتر هایده اهرابیان از استادان برجسته و اثرگذار رشته علوم کامپیوتر در دانشگاه تهران بودند؛ استادی که نام او با سال‌های دشوار اما سرنوشت‌ساز شکل‌گیری دوباره رشته علوم کامپیوتر در ایران گره



جلسه دفاع از رساله دکتری دکتر نوذری

عده گرفتند. در دهه‌ای که علوم کامپیوتر در بسیاری از دانشگاه‌ها زیر سایه مهندسی کامپیوتر قرار گرفته بود، دکتر اهرابی‌ان از جمله استادانی بود که برای باز تأسیس رشته علوم کامپیوتر به عنوان یک هویت مستقل دانشگاهی تلاش کردند.

در سال ۱۳۷۷، تلاش‌های ایشان و همکارانشان برای راه اندازی دوره کارشناسی رشته علوم کامپیوتر در دانشگاه تهران به ثمر نشست. یک سال بعد، دوره کارشناسی ارشد این رشته نیز آغاز شد و در ادامه، زمینه برای شکل‌گیری دوره دکتری فراهم گردید. بسیاری از کسانی که بعدها ستون‌های آموزشی و پژوهشی علوم کامپیوتر ایران شدند، از دل همین جریان علمی بیرون آمدند.

یکی از وجوه مهم شخصیت دانشگاهی دکتر اهرابی‌ان، توجه جدی به تربیت دانشجو بود. ایشان تنها به تدریس در کلاس بسنده نمی‌کردند؛ بلکه با دقت فراوان، دانشجویان مستعد را هدایت کرده و آنها را به مسیر پژوهش سوق می‌دادند. ایشان در سال‌هایی که دوره‌های تحصیلات تکمیلی علوم کامپیوتر در دانشگاه تهران در حال شکل‌گیری و تثبیت بود، راهنمایی رساله‌ها و پایان‌نامه‌های متعددی را بر عهده داشت و زمینه‌ساز توسعه پژوهش‌های میان‌رشته‌ای در حوزه‌های محاسبات مولکولی، بیوانفورماتیک، زیست‌شناسی سامانه‌ها و یادگیری ماشین شد. نخستین دانشجوی دکتری علوم کامپیوتر تحت راهنمایی ایشان، دکتر عباس نوذری دالینی بود که در اسفند ۱۳۸۲ از رساله دکتری خود دفاع کرده و بعدها به یکی از چهره‌های برجسته در رشته بیوانفورماتیک ایران تبدیل شدند.

در ادامه، دکتر محمد گنج‌تابش و گروهی دیگر از پژوهشگران نسل بعد نیز از همین فضای علمی بهره بردند. دکتر اهرابی‌ان راهنمایی چهار رساله دکتری و بیش از بیست پایان‌نامه کارشناسی ارشد را بر عهده داشتند. موضوعات این رساله‌ها و پایان‌نامه‌ها، طیفی گسترده از الگوریتم‌های ترکیبیاتی، محاسبات مولکولی، بیوانفورماتیک، شبکه‌های زیستی و یادگیری ماشین را در بر می‌گرفت و بازتابی از نگاه میان‌رشته‌ای و آینده‌نگر ایشان به علوم کامپیوتر بود. لیست دانشجویان تحصیلات تکمیلی ایشان به شرح زیر می‌باشد:

پژوهش‌های دکتر اهرابی‌ان دامنه‌ای متنوع اما منسجم داشت. در سال‌های نخست، تمرکز ایشان بر الگوریتم‌ها، ساختارهای درختی،



مراسم استادی دکتر اهرابی‌ان

دفاع کردند و بلافاصله وارد دوره دکتری در حوزه کامپیوتر و محاسبات شدند. این سال‌ها، علاوه بر تعمیق دانش تخصصی، دوره‌ای از بلوغ علمی و شکل‌گیری نگاه پژوهشی ایشان بود؛ نگاهی که بعدها در فعالیت‌های آموزشی، پژوهشی و مدیریتی ایشان در ایران بازتاب یافت.

ایشان در تابستان سال ۱۳۶۲ با دکتر جمشید آقازاده ازدواج کردند و اندکی بعد، در شرایطی که هنوز دوره دکتری خود را به پایان نرسانده بودند، تصمیم گرفتند به ایران بازگردند. بازگشت ایشان مصادف با یکی از پیچیده‌ترین مقاطع آموزش عالی ایران بود: سال‌های پس از انقلاب فرهنگی و آغاز بازگشایی دانشگاه‌ها. بسیاری از ساختارهای آموزشی در حال بازسازی بودند و نیاز به استنادی که بتوانند رشته‌های نوپدید را سامان دهند، بیش از هر زمان دیگری احساس می‌شد.

دکتر اهرابی‌ان از مهرماه سال ۱۳۶۲ به عنوان مربی پیمانی در گروه ریاضی و کامپیوتر دانشگاه تهران مشغول به کار شدند. همکاران و دانشجویان آن دوران بعدها روایت می‌کردند که ایشان از همان ابتدا، با دقت، نظم و تسلط علمی، حضوری متفاوت داشت. در روزگاری که امکانات آموزشی محدود بود و بسیاری از درس‌ها باید تدوین می‌شدند، دکتر اهرابی‌ان سهم مهمی در ارائه و تدوین دروس جدید علوم کامپیوتر بر عهده گرفتند. تدریس برای ایشان صرفاً انتقال دانش نبود؛ بلکه نوعی ساختن و بنیان گذاشتن بود.

با وجود مسئولیت‌های آموزشی، ایشان مسیر علمی خود را متوقف نکردند. در سال ۱۳۷۲ دوره دکتری خود را در شاخه علوم کامپیوتر نظری در دانشگاه تهران از سر گرفته و تحت راهنمایی آقای دکتر حسن صالحی فتح‌آبادی، پژوهش‌هایشان را ادامه دادند و در آذرماه سال ۱۳۷۵ از رساله دکتری خود با عنوان «الگوریتم ساخت درختان پوشا و دودویی» دفاع کرده و درجه دکتری خود را با موفقیت اخذ نمودند.

این دوره، نقطه عطفی در زندگی حرفه‌ای ایشان بود. از یک سو فعالیت‌های پژوهشی ایشان شتاب گرفت و از سوی دیگر، نقشی تاریخی در شکل‌گیری دوباره رشته علوم کامپیوتر در ایران را بر

جدیت در پیش گرفتند. ایشان همچنین از بنیان‌گذاران اولیه جریان بیوانفورماتیک در ایران و از حامیان راه‌اندازی فعالیت‌های میان‌رشته‌ای در این زمینه بودند. این آینده‌نگری، نشان می‌داد که نگاه ایشان به علوم کامپیوتر محدود به ساختارهای کلاسیک نبوده و افق‌های نوظهور علم را با دقت دنبال می‌کردند.

فعالیت‌های پژوهشی دکتر اهرابیان را می‌توان در سه محور عمده خلاصه کرد: الگوریتم‌ها و ساختارهای ترکیبیاتی، محاسبات DNA و زیست‌محاسباتی، و بیوانفورماتیک و زیست‌شناسی سامانه‌ها. ایشان در طول بیش از دو دهه فعالیت علمی، ده‌ها مقاله در مجلات معتبر بین‌المللی منتشر کردند که برخی از آنها در جامعه علمی بازتاب گسترده‌ای یافتند. از جمله مهمترین این آثار می‌توان به موارد زیر اشاره کرد:

الگوریتم‌ها و ساختارهای ترکیبیاتی:

- On the Generation of Binary Trees from (0–1) Codes, International Journal of Computer Mathematics, 1998.
- Parallel Algorithms for Minimum Spanning Tree Problem, International Journal of Computer Mathematics, 2002.
- Parallel Generation of Binary Trees in A-Order, Parallel Computing, 2005.
- Ranking and Unranking Algorithms for Loopless Generation of t-ary Trees, Logic Journal of the IGPL, 2011.
- Generation of Neuronal Trees by a New Three Letters Encoding, Computing and Informatics, 2014.

محاسبات DNA و زیست‌محاسباتی:

- DNA Algorithm for an Unbounded Fan-in Boolean Circuit, Biosystems, 2005.
- A New DNA Implementation of Finite State Machines, International Journal of Computer Science and Applications, 2006.
- Molecular Solutions for Double and Partial Digest Problems in Polynomial Time, Computing and Informatics, 2009.
- A DNA Sticker Algorithm for Solving N-Queen Problem, 2009.

بیوانفورماتیک و زیست‌شناسی سامانه‌ها:

- Genetic Algorithm Solution for Partial Digest Problem, International Journal of Bioinformatics Research and Applications, 2013.
- New Scoring Schema for Finding Motifs in DNA Sequenc-

صورت دانشجویان دکتری دکتر هایده اهرابیان

عنوان رساله	نام دانشجو
الگوریتم‌های موازی تولید درخت‌ها	دکتر عباس نودری دالینی
محاسبات مولکولی و الگوریتم‌های بیوانفورماتیک	دکتر محمد گنج‌تابش
الگوریتم‌های مختلف یافتن موتیف در شبکه‌های زیست‌شناسی	دکتر زهرا رزاقی مقدم کاشانی
الگوریتم‌های پیدا کردن موتیف در توالی‌های زیستی	دکتر فاطمه زارع میرک‌آباد

صورت دانشجویان کارشناسی ارشد دکتر هایده اهرابیان

عنوان پایان‌نامه	نام دانشجو
تولید درختان عصبی با کدهای سه حرفی	مهدی امانی
یافتن موتیف در توالی‌های زیستی با استفاده از شبکه‌های عصبی	نگین نجفی
حل مسائل NP با محاسبات مولکولی	غزاله کبیریان بجستانی
طراحی رشته‌های DNA برای محاسبات مولکولی	پیمان زرینه
جستجو در شبکه‌های زیستی	سپیده پشامی
کلاس‌بندی پروتئین‌ها با استفاده از ماشین بردار پشتیبان	امین احمدی عدل
پیش‌گویی ساختار دوم RNA با پیوندهای متقاطع	حسام‌الدین ترابی دشتی
پیش‌بینی ساختار دوم پروتئین با کمک شبکه عصبی	امیر لکی‌زاده
یافتن موتیف در توالی‌های بیولوژیکی با استفاده از شبکه‌های عصبی	آدرین جلالی خوشهر
استخراج الگوهای ساختاری	فرزاد رضانی بناب
استنتاج شبکه تنظیمی ژن از داده‌های بیان ژن	مهسا قنبری
یافتن موتیف در رشته‌های زیستی	سمیه سادات هاشمی‌فر
پیش‌گویی ساختار سوم پروتئین‌ها در مدل آب‌گریز-آب‌دوست با استفاده از الگوریتم‌های ژنتیک	زهرا زواره
الگوریتم موازی برای تشخیص رشته‌های تکراری	میثم باستانی
الگوریتم‌های بررسی فنوتایی یک گیاه	علیرضا فتوحی سیاه‌پیرانی
طبقه‌بندی ساختار دوم RNA با استفاده از ترکیب شبکه‌های عصبی	امیرحسین کاشفی
روش‌های تولید درختان فیلوژنتیک	شهاب بهشتی
الگوریتم‌های مقایسه و تراز کردن شبکه‌های زیستی	ویدا مؤمن‌زاده
الگوریتم‌های مقایسه و تراز کردن شبکه‌های زیستی	سحر اسدی
الگوریتم‌های ردیابی سوژه‌ها	سعید دهقانی
پیش‌گویی ساختار سوم پروتئین‌ها براساس مدل آب‌گریز-آب‌دوست	صغری میکائیل‌نژاد
پیش‌بینی ساختار دوم RNA با استفاده از الگوریتم‌های ژنتیک	میلاد چناقلو

تولید درخت‌های دودویی و محاسبات موازی بود. پژوهش‌ها و فعالیت‌های علمی دکتر اهرابیان در آن زمان نشان از علاقه‌مندی ایشان به مبانی نظری علوم کامپیوتر داشت. اما ایشان تنها در چارچوب‌های سنتی باقی نماندند و با ظهور حوزه‌های میان‌رشته‌ای، به‌ویژه محاسبات مولکولی و بیوانفورماتیک، از نخستین استادانی بودند که امکان پیوند میان علوم کامپیوتر و زیست‌شناسی را با

ریاضی، آمار و علوم کامپیوتر دانشگاه تهران و از چهره‌های اثرگذار در شکل‌گیری و توسعه رشته بیوانفورماتیک و زیست‌شناسی محاسباتی در ایران شد. فعالیت‌های علمی ایشان در مرز میان علوم کامپیوتر نظری، الگوریتم‌ها، تحلیل شبکه‌های پیچیده زیستی قرار داشت و طی بیش از دو دهه، در آموزش، پژوهش و تربیت نسل جدیدی از پژوهشگران علوم کامپیوتر و زیست‌محاسباتی نقش مؤثری ایفا کردند. آثار علمی او در حوزه‌هایی چون تحلیل شبکه‌های ژنی، کشف الگو در داده‌های زیستی، ساختارهای شبکه‌ای در سامانه‌های زیستی، الگوریتم‌های گراف و محاسبات مبتنی بر داده‌های مولکولی بسیار مورد توجه قرار گرفته است.

دکتر نوذری از نخستین دانش‌آموختگان دوره دکتری علوم کامپیوتر در دانشگاه تهران بودند؛ نسلی که در سال‌های آغازین بازتأسیس این رشته در کشور شکل گرفت. وی دوره دکتری خود را تحت راهنمایی زنده‌یاد دکتر هاید اهرابیان، از بنیان‌گذاران علوم کامپیوتر دانشگاه تهران، گذراند و در اسفند ۱۳۸۲ به‌عنوان یکی از نخستین دانشجویان دکتری علوم کامپیوتر در دانشگاه تهران از رساله دکتری خود با عنوان «الگوریتم‌های موازی تولید درخت‌ها» دفاع کردند. این پیوند علمی، بعدها در قالب همکاری‌های پژوهشی گسترده میان ایشان و دکتر اهرابیان ادامه یافت و بخشی از سنت علمی گروه علوم کامپیوتر دانشگاه تهران را شکل داد. آثار مشترک متعددی از این دو استاد در زمینه‌هایی مانند محاسبات مولکولی، تحلیل شبکه‌ها و الگوریتم‌های ساختاری در نشریات معتبر بین‌المللی منتشر شده است.

پس از پایان تحصیلات، دکتر نوذری دالینی به عضویت هیئت علمی گروه علوم کامپیوتر دانشگاه تهران در آمدند و به‌تدریج به یکی از استادان شناخته‌شده حوزه بیوانفورماتیک و علوم داده‌های زیستی بدل شدند. در دوره‌ای که این حوزه در ایران هنوز نوپا بود، او از معدود استادانی بود که تلاش کرد پیوندی نظام‌مند میان علوم کامپیوتر، زیست‌شناسی مولکولی و ریاضیات برقرار کند. حضور او در پروژه‌های میان‌رشته‌ای، به‌ویژه در همکاری با پژوهشگران مرکز تحقیقات بیوشیمی و بیوفیزیک دانشگاه تهران، نقش مهمی در شکل‌گیری پژوهش‌های نوین بیوانفورماتیک در کشور داشت.

یکی از حوزه‌های مهم پژوهشی دکتر نوذری دالینی، تحلیل شبکه‌های تنظیم ژنی (Gene Regulatory Networks) و مدل‌سازی روابط میان ژن‌ها بود. ایشان در چندین پژوهش برجسته، روش‌های محاسباتی و الگوریتمی برای استنتاج روابط میان ژن‌ها از داده‌های زیستی ارائه کرد. آثار منتشر شده از ایشان نمونه‌هایی از این رویکرد میان‌رشته‌ای است که در آن از روش‌های محاسباتی برای بازسازی شبکه‌های هم‌بستگی ژنی استفاده شده است. این پژوهش‌ها در زمان خود از نخستین تلاش‌های جدی در ایران برای بهره‌گیری از ابزارهای علوم کامپیوتر در تحلیل شبکه‌های زیستی و سرطان‌شناسی محاسباتی نیز در حوزه تحلیل شبکه‌های زیستی و سرطان‌شناسی محاسباتی نیز

es, BMC Bioinformatics, 2009.

● Genetic Algorithm for Dyad Pattern Finding in DNA Sequences, Genes and Genetic Systems, 2009.

● Kavosh: A New Algorithm for Finding Network Motifs, BMC Bioinformatics, 2009.

● Tag SNP Selection via a Genetic Algorithm, Journal of Bio-medical Informatics, 2010.

● PreRkTAG: Prediction of RNA Knotted Structures Using Tree Adjoining Grammars, Iranian Journal of Biotechnology, 2013.

اما شاید آنچه بیش از مقالات و مسئولیت‌ها از دکتر اهرابیان در ذهن شاگردانشان باقی‌مانده، شیوه معلمی ایشان باشد. دانشجویان ایشان از استادی سخن می‌گفتند که در عین سخت‌گیری علمی، دلسوزی فراوانی داشتند؛ کسی که نظم، دقت و عمق فهم را جدی می‌گرفتند، اما در عین حال، به رشد دانشجو باور داشتند. برای بسیاری، کلاس‌های ایشان فقط محل یادگیری درس نبود؛ جایی بود برای آموختن شیوه فکر کردن.

سرانجام، دکتر هاید اهرابیان در ۱۱ تیرماه ۱۳۹۰ پس از یک دوره بیماری دار فانی را وداع گفتند. فقدان ایشان برای خانواده، دوستان، همکاران و جامعه علوم کامپیوتر ایران اندوهی عمیق بر جای گذاشت. با این حال، تأثیر او در ساختار رشته علوم کامپیوتر دانشگاه تهران، در دانشجویانی که تربیت کردند، و در مسیرهایی که برای علوم میان‌رشته‌ای گشودند، همچنان زنده است.

شاید بهترین توصیف از زندگی دکتر هاید اهرابیان این باشد که او بخشی از عمر خود را نه صرفاً به آموزش علوم کامپیوتر، بلکه به ساختن بنیان‌های آن در ایران اختصاص دادند؛ استادی که با دانش، پشتکار و تعهدی آرام اما ماندگار، هم رشته‌ای علمی را استوارتر کردند و هم در ذهن و منش نسل‌هایی از دانشجویان، اثری عمیق بر جای گذاشتند؛ امکانی که امروز نسل‌های بعدی پژوهشگران، بدون آن که نام ایشان را بدانند، بر شانه‌های آن ایستاده‌اند.



دکتر عباس نوذری دالینی

دکتر عباس نوذری دالینی در دهم دی ماه ۱۳۴۳ در شیراز به دنیا آمد و سپس از استادان برجسته علوم کامپیوتر در دانشکده



مراسم روز معلم (مجتهدی، نوذری و یاسمی)



جلسه دوستانه در IPM



دفاع از رساله دکتری دکتر مظفری با حضور داوران فرانسوی

(2006).

- Molecular Solutions for Double and Partial Digest Problems in Polynomial Time (Computing and Informatics, 2009).
- A DNA Sticker Algorithm for Solving N-Queen Problem (2009).

دکتر نوذری نقش فعالی داشتند. او در پژوهش‌هایی درباره سرطان ریه و سرطان کولون، با استفاده از مدل‌سازی شبکه‌های هم‌بپایانی ژن‌ها و تحلیل ساختارهای تنظیمی، در شناسایی مؤلفه‌های کلیدی بیماری و ژن‌های مؤثر مشارکت داشتند. این پژوهش‌ها که در مجلات بین‌المللی منتشر شده‌اند، نشان‌دهنده گرایش ایشان به استفاده از الگوریتم‌ها و علوم شبکه در مسایل پزشکی و زیستی بود.

دکتر نوذری همچنین در حوزه محاسبات مولکولی و کشف الگوهای ژنتیکی پژوهش‌های اثرگذاری انجام داده‌اند. نام ایشان در مقالاتی درباره کشف موتیف‌های DNA، تحلیل توالی‌های ژنتیکی و طراحی الگوریتم‌های تکاملی برای یافتن ساختارهای زیستی دیده می‌شود. آثار منتشر شده از ایشان بخشی از این مسیر علمی بودند که در آن علوم کامپیوتر نظری با زیست‌شناسی مولکولی پیوند می‌خورد. در بسیاری از این آثار، همکاری نزدیک او با زنده‌یاد دکتر هایدو اهرابیان و دیگر پژوهشگران بیوانفورماتیک دانشگاه تهران مشهود است.

از دیگر آثار شاخص دکتر نوذری دالینی، مشارکت در توسعه الگوریتم کاوش (Kavosh) برای کشف الگوهای شبکه‌ای (Network Motifs) است؛ الگوریتمی که در جامعه بیوانفورماتیک بازتاب گسترده‌ای یافت و در مقالات بین‌المللی بارها مورد استناد قرار گرفت. این پژوهش، که با همکاری پژوهشگران داخلی و بین‌المللی انجام شد، نمونه‌ای از مشارکت ایشان در تولید ابزارهای محاسباتی قابل استفاده در تحلیل شبکه‌های پیچیده زیستی بود.

فعالیت‌های پژوهشی دکتر عباس نوذری دالینی که اولین مقاله‌اش در سال ۱۹۹۵ در زمینه طرح‌های ترکیبیاتی با دکتر ترابی و دکتر خسروشاهی و آخرین مقاله‌اش را در زمینه تشخیص ژن‌های سرطانی در سال ۲۰۲۲ نوشت، می‌توان در چهار محور عمده به شرح زیر خلاصه کرد که برای هر محور، مقالات شاخص ایشان به شرح زیرند:

الگوریتم‌ها و ساختارهای ترکیبیاتی:

- On the Generation of Binary Trees from (0-1) Codes (1998).
- Parallel Algorithms for Minimum Spanning Tree Problem (2002).
- Generation of Binary Trees in B-Order from (0-1) Sequences (2004).
- Parallel Generation of Binary Trees in A-Order (2005).
- Parallel Generation of t-ary Trees in A-Order (2007).
- Ranking and Unranking Algorithms for Loopless Generation of t-ary Trees (2011).

محاسبات مولکولی:

- DNA Algorithm for an Unbounded Fan-in Boolean Circuit (Biosystems, 2005).
- A New DNA Implementation of Finite State Machines

تکمیلی ایشان به شرح زیر می باشد:

لیست دانشجویان دکتری دکتر عباس نوذری دالینی

نام دانشجو	عنوان رساله	استادان راهنما
میلاد مظفری	پردازش الگو با استفاده از فرایند یادگیری تقوینی مبتنی بر مغز	دکتر عباس نوذری دالینی، دکتر محمد گنج تابش
نجمه خراسانی	مدل سازی سازوکار سلول های بنیادی جهت هم ایستایی در بافت تمایز یافته آنها	دکتر عباس نوذری دالینی، دکتر مهدی صادقی
فاطمه شریفی زاده	بهبود عملکرد مدل های بازشناسی اشیا با استفاده از سیگنال های بالا به پایین	دکتر عباس نوذری دالینی، دکتر محمد گنج تابش
احسان پرنور	شناسایی مؤلفه های عملکردی در بیماری سرطان با استفاده از شبکه های چندلایه و رویکرد بیولوژی سیستمی ترکیبی	دکتر عباس نوذری دالینی، دکتر علی مسعودی نژاد
قاسم مهدور	روش محاسباتی برای استنتاج شبکه های تنظیم ژن	دکتر عباس نوذری دالینی، دکتر مهدی صادقی
زینب السادات موسویان حر	مطالعه پاسخ دهی به درمان در سرطان با رویکرد سیستم بیولوژی	دکتر عباس نوذری دالینی، دکتر علی مسعودی نژاد

لیست دانشجویان کارشناسی ارشد دکتر عباس نوذری دالینی

نام دانشجو	عنوان پایان نامه	استادان (ان) راهنما
غزاله کبیریان بجستانی	حل مسائل NP با محاسبات مولکولی	دکتر عباس نوذری دالینی
سیداحسان سیدی طبری	تولید درخت های با n گروه داخلی و m برگ	دکتر عباس نوذری دالینی
پیمان زرینه	طراحی رشته های DNA برای محاسبات مولکولی	دکتر عباس نوذری دالینی
سیدحسین علوی سلطانی	بررسی الگوریتم های موازی تطابق رشته ها	دکتر عباس نوذری دالینی
جواد محمدزاده	بازسازی هاپلوتا پ از روی قطعات	دکتر عباس نوذری دالینی
امیر لکی زاده	پیش بینی ساختار دوم پروتئین با کمک شبکه عصبی	دکتر عباس نوذری دالینی
جواد ظهیری	بررسی مسئله های بلاک بندی در هاپلوتا پ ها	دکتر عباس نوذری دالینی
مهدی امانی	تولید درختان عصبی با کدهای سه حرفی	دکتر عباس نوذری دالینی
حسام الدین ترابی دشتی	پیش گویی ساختار دوم RNA با پیوندهای متقاطع	دکتر عباس نوذری دالینی
آدرین جلالی خوشهر	یافتن موتیف در توالی های بیولوژیکی با استفاده از شبکه های عصبی	دکتر عباس نوذری دالینی
نجمه خراسانی	مقایسه RNA با استفاده از روش های پردازش تصویر	دکتر عباس نوذری دالینی، دکتر محمد گنج تابش
افسانه جعفری	آنالیز ژن ها بر اساس داده های اپی ژنتیکی	دکتر عباس نوذری دالینی، دکتر مهدی صادقی
طه خاکساری	یادگیری ترکیبی بدون ناظر	دکتر عباس نوذری دالینی، دکتر هدی ساجدی

● A Constant Time Algorithm for DNA Add (2009).

بیوانفورماتیک و زیست شناسی سامانه ها:

- Genetic Algorithm Solution for Double Digest Problem (2012)
- Genetic Algorithm Solution for Partial Digest Problem (2013)
- New Scoring Schema for Finding Motifs in DNA Sequences (BMC Bioinformatics, 2009).
- Genetic Algorithm for Dyad Pattern Finding in DNA Sequences (2009).
- Tag SNP Selection via a Genetic Algorithm (Journal of Bio-medical Informatics, 2010).
- PreRkTAG: Prediction of RNA Knotted Structures Using Tree Adjoining Grammars (2013).
- Kavosh: A New Algorithm for Finding Network Motifs (BMC Bioinformatics, 2009).
- Target Controllability with Minimal Mediators in Complex Biological Networks (Genomics, 2020).
- Integrated Analysis of Gene Expression and Regulatory Networks in Colorectal Cancer (2020).
- Network-based Identification of Key Regulatory Elements in Cancer-related Pathways (2021).

علوم اعصاب محاسباتی و یادگیری زیست الهام پذیر:

- First-spike Based Visual Categorization Using Reward-modulated STDP (2017).
- Bio-inspired Digit Recognition Using Reward-modulated Spike-timing-dependent Plasticity in Deep Convolutional Networks (2018).
- Object Recognition Using Deep Spiking Neural Networks with Reinforcement Learning (2019).
- Unsupervised Feature Learning in Spiking Neural Networks for Visual Pattern Recognition (2020).

در کنار پژوهش، تدریس و تربیت دانشجو بخش مهمی از زندگی حرفه ای دکتر نوذری را تشکیل می داد. او سال های متمادی در دانشگاه تهران دروس تخصصی علوم کامپیوتر، بیوانفورماتیک، الگوریتم ها و تحلیل داده های زیستی را تدریس نمود و راهنمایی شمار قابل توجهی از پایان نامه های کارشناسی ارشد و رساله های دکتری را بر عهده داشت. بسیاری از دانشجویان او بعدها به اعضای هیئت علمی، پژوهشگران و متخصصان حوزه های داده، زیست فناوری و علوم کامپیوتر تبدیل شدند. همکاران و دانشجویان، او را استادی دقیق، سخت کوش، متعهد به کیفیت علمی و در عین حال حامی رشد دانشجویان توصیف می کنند. لیست دانشجویان تحصیلات

آموزشی و هدایت فعالیت‌های پژوهشی دانشجویان ایفا نمودند. عضویت در شورای پژوهشی مرکز انفورماتیک دانشگاه تهران، شورای علمی پژوهش‌شده علوم زیستی پژوهشگاه دانش‌های بنیادی، کمیته ارتقای اعضای هیئت علمی دانشگاه علامه طباطبائی و کمیته تخصصی تبدیل وضعیت اعضای هیئت علمی دانشگاه تربیت مدرس، بخشی دیگر از مسئولیت‌های اجرایی و حرفه‌ای ایشان را تشکیل می‌دهد. افزون بر این، ایشان سال‌ها به عنوان استاد راهنمای دانشجویان کارشناسی و ورودی‌های مختلف رشته علوم کامپیوتر دانشگاه تهران فعالیت داشتند و در هدایت علمی و آموزشی نسل‌های متعددی از دانشجویان نقش آفرین بود.

دکتر عباس نوذری دالینی، پس از ماه‌ها مبارزه با بیماری سرطان، در بیستم بهمن‌ماه سال ۱۴۰۱ دار فانی را وداع گفتند و فقدان ایشان برای جامعه علوم کامپیوتر و زیست‌محاسباتی ایران ضایعه‌ای اندوه‌بار بود. با این حال، میراث علمی او نه تنها در قالب مقالات و پروژه‌های پژوهشی، بلکه در نسل دانشجویانی که تربیت کردند و حوزه‌هایی که به رشد آنها کمک نمودند، همچنان تداوم دارد.

شاید بهترین توصیف از زندگی دکتر نوذری این باشد که او از شمار پژوهشگرانی بود که مرزهای رشته خود را جدی می‌گرفتند، اما به آن محدود نمی‌ماندند؛ استادی که با پیوند دادن علوم کامپیوتر، ریاضیات و زیست‌شناسی، نه تنها در گسترش یک حوزه علمی نوپا نقش داشتند، بلکه در پرورش نسلی از پژوهشگران، اثری ماندگار بر جای گذاشتند.

* * *

با گرمی داشت قدرشناسانه نام و یاد دکتر هاید اهراییان و دکتر عباس نوذری دالینی، انجمن انفورماتیک ایران خدمات صادقانه و صمیمانه ایشان در طول حیات پربارشان را در قالب عضویت در هیئت تحریریه و ویراستاران علمی ماهنامه گزارش کامپیوتر طی سال‌ها پاس می‌دارد.

همچنین لازم به ذکر است که دکتر عباس نوذری دالینی از ابتدای راه‌اندازی مجله علوم رایانشی، فصلنامه علمی انجمن انفورماتیک ایران، تا زمان درگذشت تأسف بار و نابهنگامشان سردبیری این مجله را بر عهده داشتند.

استاد(ان) راهنما	عنوان پایان‌نامه	نام دانشجو
دکتر عباس نوذری دالینی، دکتر مهدی صادقی	بررسی الگوریتم‌های پیش‌بینی برهم‌کنش مولکول‌های RNA- پروتئین	محسن صالحی
دکتر عباس نوذری دالینی، دکتر مهدی صادقی	بررسی الگوریتم پیش‌گویی مسئله تاخوردگی معکوس RNA	سمیرا کاظمی نافچی
دکتر عباس نوذری دالینی، دکتر مرتضی محمدنوری	تولید درختان t-تایی در ترتیب شبه‌قاموسی	غزاله پروینی
دکتر عباس نوذری دالینی	بررسی الگوریتم‌های پیش‌بینی ساختار دوم RNA همراه با شبه‌گره	علی پارکی طرودی
دکتر عباس نوذری دالینی، دکتر زهرا رزاقی مقدم کاشانی	تعیین نقش دومین پروتئین‌ها در سرطان پس از بهینه‌سازی خصیصه‌ها	سیرانا هاشمی رنجبر
دکتر عباس نوذری دالینی، دکتر مهدی صادقی	تخصیص دومین‌های ساختاری پروتئین با استفاده از روش‌های یادگیری ماشین	فاطمه سروی
دکتر عباس نوذری دالینی، دکتر مهدی صادقی	استنتاج شبکه تنظیمی بیان ژن به وسیله داده‌های تأثیر اختلال ژنی	مسعود ملایم‌کنک لاهیجانی
دکتر عباس نوذری دالینی	رتبه‌گذاری و رتبه‌گشایی درختان t-تایی	مهداد فرخ‌نژاد
دکتر عباس نوذری دالینی	تولید، رتبه‌گذاری و رتبه‌گشایی درختان t-تایی در ترتیب حداقل فاصله	نگین شاه‌منصوری
دکتر عباس نوذری دالینی	استخراج ژنومی خودکار ساختارهای طبیعی پپتیدی به مدد طیف‌سنجی جرمی	سیدمهدی عزیزی سیدبیگلر
دکتر عباس نوذری دالینی، دکتر زینب‌السادات موسویان حر	پیش‌گویی بیماری آلزایمر با استفاده از تصاویر MRI	سارا ابریشمی
دکتر عباس نوذری دالینی، دکتر سیدامیر مرعشی	بررسی بیماری ام‌اس با روش‌های یادگیری ماشین	مطهره حب‌علی
دکتر عباس نوذری دالینی، دکتر زینب‌السادات موسویان حر	طراحی و پیاده‌سازی محیطی مبتنی بر چارچوب اسپارک جهت تحلیل و پردازش توالی‌های زیستی	علیرضا محمدی
دکتر عباس نوذری دالینی، دکتر محمد گنج‌تابش	بررسی تأثیر محتوا در بازشناسی اشیا با استفاده از شبکه‌های عصبی ضربه‌ای	عارف عزیزپور
دکتر عباس نوذری دالینی، دکتر سیدمجتبی مجتهدی	درستی‌یابی صوری شبکه‌های عصبی عمیق با تابع فعال‌ساز تکه‌ای خطی	محمد صادق آبادی فراهانی
دکتر عباس نوذری دالینی، دکتر محمد گنج‌تابش	مدل‌سازی سازوکارهای مسئله انقیاد در شبکه‌های عصبی ضربه‌ای	امیراصلاح اصلائی

دکتر نوذری در کنار فعالیت‌های آموزشی و پژوهشی، حضوری فعال در عرصه مدیریت دانشگاهی و سیاست‌گذاری علمی داشت. ایشان از سال ۱۳۸۴ تا ۱۳۸۹ مدیریت بخش علوم کامپیوتر دانشگاه تهران را بر عهده داشتند و هم‌زمان، بین سال‌های ۱۳۸۵ تا ۱۳۹۰ ریاست خدمات رایانه‌ای پردیس علوم دانشگاه تهران را عهده‌دار بودند. از سال ۱۳۹۰ تا ۱۳۹۷ نیز به عنوان معاون تحصیلات تکمیلی دانشکده ریاضی، آمار و علوم کامپیوتر فعالیت کردند و در این جایگاه نقش مهمی در توسعه دوره‌های تحصیلات تکمیلی، سامان‌دهی امور

شهروند وظیفه‌شناس در جامعه علمی: بدهی داوری خود را تسویه کنید

نویسنده: مارتین مانپرس

استاد تمام مهندسی نرم‌افزار، مؤسسه فناوری سلطنتی کی‌تی‌اچ

مترجم: علیرضا خلیلان

استادیار مهندسی نرم‌افزار، مدیر مرکز رشد، کارآفرینی و فناوری

رایانشانی: akhalilian@gmail.com

شناسه اُرکید: ۷۰۳X-۴۰۷۹-۰۰۰۲-۰۰۰۰

بدهی داوری چیست؟

وقتی مقاله‌ای پذیرفته و منتشر می‌شود، معمولاً سه داور آن را ارزیابی کرده‌اند. پس نویسندگان آن مقاله همگی از زحمتی که برای داوری کشیده شده است منتفع شده‌اند. بدهی داوری هر دانشمند تعداد داوری‌هایی است که به جمع هم‌تایان و هم‌رشته‌ای‌های خودش بدهکار است تا زحمت داوری‌هایی که برای مقاله‌هایش کشیده شده است جبران شود.

بدهی هر مقاله با N نویسنده، سه داوری است. نویسنده اول عموماً دانشجوی دکتری است؛ سایر نویسندگان پرسابقه‌تر هم در بدهی داوری‌ها شریک هستند، پس هر یک از نویسندگان مقاله، بجز نفر اول، تقریباً $(N-1)/3$ داوری به‌ازای هر مقاله بدهکار است.

برای نمونه، استادی که در ۲۰ مقاله نشریه (به‌عنوان سایر نویسندگان، غیر از نویسنده اول) نویسندگی کرده است با این فرض که متوسط تعداد نویسندگان هر مقاله ۳ نفر باشد، $20 \times 3/2 = 30$ داوری نشریه‌ای بدهکار است. این برآورد حد پایین است. مقاله‌های ارسالی که مردود می‌شوند هم زحمت داوری داشته‌اند ولی در پایگاه‌های همگانی مثل دی‌بی‌آل‌پی درج نمی‌شوند. پردازش به زبان پایتون^۵ فراهم است که بدهی داوری را بر حسب محل انتشار از نمایه دی‌بی‌آل‌پی فرد محاسبه می‌کند.

دوام یک سامانه مستلزم بدهی خالص صفر است

نظام همتادآوری می‌تواند به حیات سالم خود ادامه دهد اگر و تنها اگر میزان عرضه و تقاضای داوری برابر باشند. هر مقاله منتشر شده

یادداشت مترجم

این نوشتار مشابه آنچه در شماره‌های پیشین منتشر شد، ترجمه جستاری^۱ است که مارتین مانپرس در وبگاه شخصی خودش منتشر کرده است. موضوع آن نوعی بدهی اخلاقی دانشمندان به همدیگر است؛ همتادآوری^۲. هر بار که مقاله‌ای می‌نویسید باید صحت و اعتبار آن سنجیده شود. این کار را دو یا سه نفر از ساکنان این کره خاکی که هم‌رشته شما باشند به رایگان تقبل می‌کنند. به این ترتیب شما به دانشمندان هم‌رشته خودتان به‌نوعی بدهکار می‌شوید. این بدهی زمانی صاف می‌شود که شما هم مقاله‌های دانشمندان هم‌رشته خودتان را به رایگان داوری کنید. راقم این سطور جستار حاضر را برای خودش این‌طور خلاصه می‌کند: اگر این تعریف را برای علم بپذیریم که «علم سامانه‌ای است برای انباشت دانش قابل اعتماد^۳ که با کوششی بی‌امان از راه شک‌اندیشی سامانمند حاصل می‌شود^۴، در این صورت همتادآوری و تسویه بدهی داوری سامانه علم را سرپا و سر حال نگه می‌دارد.

همتادآوری چرخه اصلاحی و پالایه کیفی اصلی در علم است. هر مقاله‌ای که به مرحله انتشار می‌رسد چندین داور آن را خوانده و ارزیابی کرده‌اند. ارزیابی کاری است که زمان می‌برد و زحمت دارد. مدل بدهی داوری این هزینه را به مقدار کمی تبدیل می‌کند و معیاری دقیق برای تداوم حیات سالم نظام همتادآوری به‌دست می‌دهد.

1-<https://www.monperrus.net/martin/reviewer-debt>

2-Peer review

3-Zobel, J. (2014). Writing for Computer Science, Springer International Publishing.

۴-ریس، م. (۲۰۱۲). از اینجا تا بی‌نهایت دورنمایی از آینده علم. ترجمه محمدابراهیم محبوب، نشر نی، چاپ سوم، ۱۳۹۷.

5-<https://github.com/monperrus/review-debt>

ولی قانون پیشاهنگی از او ۳۱ دآوری می‌طلبد.

دشواریابی داور مشکلی رفتاری است

سردبیران نشریه‌ها و رؤسای برنامه‌ریزی گردهمایی‌ها معمولاً در پیدا کردن داورانی که حاضر به همکاری باشند مشکل دارند. علت این دشواریابی را هم افزایش تعداد مقاله‌های ارسال‌شده، توان‌فرسایی داوران و محدودیت‌های زمانی ناشی از سایر مسئولیت‌ها در نظر می‌گیرند. این دلایل درست هستند اما بیانگر تمام واقعیت نیستند. دلیل اصلی این است که همگان بدهی دآوری خود را به‌طور تمام‌وکمال پرداخت نکرده‌اند. اگر داوریهایی که هر دانشمند فعال تقبل می‌کند به نسبت آثار منتشره خودش باشد ظرفیت دآوری نهاد علمی با نیازها هم‌خوانی خواهد داشت. مشکل کمبود داور واجد شرایط نیست. مشکل مشارکت ناکافی افرادی است که به‌رغم داشتن صلاحیت از مشارکت امتناع می‌کنند.

مدل بدهی دآوری به‌خودی خود این مشکل رفتاری را حل نمی‌کند. اما تعهد به همتادآوری را نمایان و قابل اندازه‌گیری می‌نماید. هر دانشمندی که از بدهی دقیق خودش هم به‌طور کلی و هم محل‌بهمحل مطلع باشد، برنامه‌ای ملموس و قابل اجرا خواهد داشت. مدل عبارت میهم (»وظیفه دارید دآوری کنید«) را به عبارت دقیق (»شما ۱۲ دآوری به آی‌تریپل‌ای تی‌اس‌ای و ۸ دآوری به آی‌ام‌اس‌ای بدهکار هستید«) تبدیل می‌کند.

تأثیر هوش (دانایی) مصنوعی

مدل‌های بزرگ زبانی هم در مقاله‌نویسی و هم در دآوری می‌توانند یاریگر باشند. نکته اصلی همین تقارن است: اگر هوش (دانایی) مصنوعی بتواند زحمت مقاله‌نویسی را کم کند پس زحمت دآوری را هم می‌تواند کاهش دهد. بدین‌سان نسبت منفعت از همتادآوری [توسط دیگران] به زحمت همتادآوری [توسط خود فرد] بی‌تغییر می‌ماند.

مقاله‌ای که به کمک هوش (دانایی) مصنوعی نوشته شده باشد باز هم به سه نفر برای همتادآوری زحمت می‌دهد و بدهی نویسنده تغییری نمی‌کند. اگر ابزارهای هوش (دانایی) مصنوعی هزینه متوسط هر دآوری را کاهش دهند، ظرفیت کلی نهاد علمی بیشتر می‌شود، البته میزان ارسال مقاله‌ها هم افزایش می‌یابد؛ پیامدها متقارن هستند.

نتیجه‌گیری

وقتی که هر دانشمند بدهی دآوری خودش را تقبل کند و مطابق قانون پیشاهنگی اندکی میزان آن را نیز افزایش دهد، دیگر سردبیران مشکلی در یافتن داوران نخواهند داشت. بی‌گمان داور واجد شرایط موجود است؛ آنها همان نویسندگان مقاله‌های در دست نگارش و ارسال هستند. دشواری در یافتن این افراد به ظرفیت ربطی ندارد. موضوع میزان مشارکت است.

بدهی دآوری این تعهد را ملموس می‌سازد. هر دانشمند می‌تواند حساب کند چقدر و به کجا بدهی دارد و نیازی نیست منتظر بماند تا نشریه‌ای آن را اجبار کند یا انجمن‌های حرفه‌ای سامانه ثبت‌وضبط راه‌اندازی نمایند.

حساب‌و‌کتاب بدهی کار ساده‌ای است. بدهی دآوری‌تان را تسویه کنید

سه دآوری تقاضا کرده است. چنانچه هر دانشمند دقیقاً به میزان بدهی‌اش دآوری انجام دهد، عرضه و تقاضا برابری می‌کنند و نظام همتادآوری در تعادل می‌ماند.

این استدلال به پایستگی اشاره دارد. هر مقاله سه دآوری مصرف می‌کند. اگر هر نویسنده سهم خودش را بپردازد، نظام همتادآوری توازن خودش را حفظ می‌کند. هیچ انگاشته دیگری نیز لازم نیست استدلال مطرح شده قابل تعمیم است: سامانه‌ای که در آن هر فرد بدهی خودش را بپردازد پایدار می‌ماند؛ سامانه‌ای که عدم تأدیه در آن گسترش یابد، دچار زوال می‌شود تا جایی که گروه کوچکی مجبور می‌شوند خیلی بیشتر بپردازند تا جور بقیه را بکشند.

بی‌نیازی از یک زیرساخت متمرکز

پیشنهادهایی برای بهبود همتادآوری مطرح شده است مبنی بر این که سامانه‌ای برای ثبت‌وضبط متمرکز دآوری ایجاد شود، یعنی نشریه‌ها یا انجمن‌های حرفه‌ای مشارکت داوران را ثبت نمایند و بر آن اساس نوعی گواهی امتیاز یا سند بستانکاری برای داوران صادر شود. این سامانه‌ها با موانعی مانند دشواری فراگیری و استفاده گسترده، نگرانی‌های مربوط به حفظ حریم خصوصی و هزینه‌های هماهنگی میان نهادهای مختلف روبه‌رو هستند.

مدل بدهی داور به هیچ‌کدام از این موارد نیاز ندارد. همه دانشمندان می‌توانند بدهی خودشان را حساب کنند. هر پژوهشگر می‌تواند بدهی دقیق خودش را بر حسب محل دآوری تعیین نمایند و در این کار به هیچ نهاد مرکزی، ثبت‌نام یا هماهنگی نیز نیاز ندارد. این مدل کاملاً نامتمرکز است و هر کس می‌تواند حسابرسی خودش را انجام دهد.

به‌کارگیری قانون پیشاهنگی^۶ در همتادآوری

قانون پیشاهنگی در مهندسی نرم‌افزار چنین است: اردوگاه را تمیزتر از روز اول تحویل دهید (مارتین، کد تمیز، ۲۰۰۸)^۸. بیان آن در پایگاه کد چنین می‌شود: هر مشارکت در کدنویسی باید کیفیت کد را اندکی از نیازمندی روز اول بهبود دهد. بیان آن در همتادآوری نیز چنین می‌شود: اندکی بیشتر از بدهی دآوری خود در فرایند همتادآوری مشارکت کنید.

این افزایش اندک به دو دلیل اهمیت دارد. دانشجویان دکتری و پژوهشگران تازه‌کار هنوز پیشینه انتشار ندارند ولی به‌زودی آثار آنان منتشر می‌شود. تا آن زمان جمع دانشمندان و پژوهشگران نیاز دآوری آنان را تأمین می‌کنند. پس پژوهشگران جالفتاده و باسابقه باید مشارکت بیشتری داشته باشند تا نیاز پژوهشگران جوان برآورده شود. یک برنامه اجرایی: به‌ازای هر مقاله‌ای که نویسنده اول نیستید، یک دآوری بیشتر از میزان بدهی اصلی خود انجام دهید. استادی که در ۲۰ مقاله نشریه هم‌نویسندگی^۹ کرده است به فرض این که متوسط نویسندگان هر مقاله ۳ نفر باشد، بدهی اصلی دآوری او ۳۰ تاست؛

6-Assumption

7-Boy scout

۸-این قانون را رابرت سی مارتین معروف به عمو باب در کتاب کد تمیز مطرح کرد که تفسیر نرم‌افزاری آن چنین است: هرگاه روی بخشی از کد کار می‌کنی آن را در وضعیتی بهتر از روز اول تحویل بده، مثلاً اگر متوجه کدهای تکراری یا نام‌گذاری نامناسب شدی و اصلاح آن خیلی زمانگیر نبود، علاوه بر کار خود آن مشکلات را هم برطرف کن. این قانون اخلاقی به بهبود تدریجی و مشارکتی کیفیت کد کمک می‌کند. م.

۹-یعنی نویسنده اول نبوده است. م.

دومینیکر گارسیا مارثا

پاتریس کالوو

دموکراسی الگوریتمی

به نقل از روزنامه سازندگی - شماره ۲۲۹۰ - ۱۴۰۵/۰۵/۰۶

● دموکراسی الگوریتمی

● نگاهی انتقادی بر پایه دموکراسی مبتنی بر گفت‌وگو

● دومینگو گارسیا-مارثا و پاتریس کالوو

● ترجمهٔ راضیه کریمی / انتشارات همشهری

● ۳۶۵ صفحه / ۴۵۰۰۰ تومان



الگورکراسی

دموکراسی الگوریتمی

رویکردی انتقادی بر بنیان دموکراسی شورایی و رهایی از استعمار الگوریتمی (نگرانی احتمالی ناشر ترجمه کتاب)

حکومت‌داری به‌وسیله الگوریتم که با نام‌هایی مانند تنظیم‌گری الگوریتمی، تنظیم مقررات با الگوریتم‌ها، حاکمیت الگوریتمی، حکمرانی الگورایی، نظام حقوقی الگوریتمی یا الگورکراسی نیز شناخته می‌شود شکلی جایگزین از حکومت یا سازمان اجتماعی است که در آن الگوریتم‌های رایانه‌ای برای تنظیم مقررات، اجرای قوانین و به‌طور کلی هر جنبه‌ای از زندگی روزمره مانند حمل‌ونقل یا ثبت زمین استفاده می‌شود. اصطلاح دولت‌داری با الگوریتم نخستین بار در سال ۲۰۱۳ در ادبیات علمی جایگزین حاکمیت الگوریتمی شد. اصطلاح مرتبط رایج تنظیم‌گری الگوریتمی به معنای تعیین استاندارد، پایش و اصلاح رفتار از طریق الگوریتم‌های رایانشی است که شامل خودکارسازی قوه قضائیه هم می‌شود. دولت‌داری با الگوریتم چالش‌های تازه دیگری را هم مطرح کرده است که در ادبیات پیشین دولت الکترونیکی جایی نداشتند. برخی از اصطلاح‌های سایبرکراسی هم برای بیان این مفهوم استفاده می‌کنند که شکلی فرضی از حکومت با بهره‌گیری موثر از اطلاعات است و در آن حاکمیت را الگوریتمی می‌شمارند هر چند الگوریتم‌ها تنها ابزار پردازش اطلاعات نیستند. کسانی مانند نلو کریستیانینی و ترزا اسکنتامپورلو، استدلال کرده‌اند ترکیب یک جامعه انسانی با برخی الگوریتم‌های تنظیم‌گر مانند الگوریتم‌هایی نظیر امتیازدهی بر پایه اعتبار یا رعایت مقررات شکلی از ماشین اجتماعی را پدید می‌آورد. صاحب‌نظر موسع در این پندار خانم شوشانا زوبف است که از عصر ما با عنوان عصر سرمایه‌داری نظارتی یاد می‌کند که نبرد برای آینده انسانی در آن در مرزهای جدید قدرت و حکمرانی به معنای قدرت شکل می‌گیرد

به انتخاب و ویراستهٔ سید ابراهیم ابطحی

استادیار دانشکده مهندسی کامپیوتر دانشگاه صنعتی شریف

پست الکترونیکی: abtahi@Sharif.edu

یادداشت سردبیر: ویرایش‌های معمول نشریه در بخش در رسانه‌ها اعمال نمی‌گردد و مطالب عیناً نقل می‌شود.

کتاب «دموکراسی الگوریتمی: نگاهی انتقادی بر پایه دموکراسی مبتنی بر گفت‌وگو»، نوشته دومینگو گارسیا-مارثا و پاتریسی کالوو است که انتشارات همشهری به‌تازگی آن را با ترجمه راضیه کریمی منتشر کرده است. کتاب «دموکراسی الگوریتمی» رویکردی تحلیلی و انتقادی دارد و به بررسی تأثیر فناوری‌های نوین به‌ویژه هوش مصنوعی، بلاک‌چین، داده‌کاوی و متاورس بر نهادهای دموکراتیک، اخلاق عمومی و مشارکت مدنی می‌پردازد. نویسندگان کتاب با تکیه بر اصول دموکراسی گفت‌وگو محور، در اندیشه باز طراحی ساختارهای نهادی‌اند تا شفافیت، مسئولیت‌پذیری، مشارکت دوسویه و عدالت در طراحی، به‌کارگیری و نظارت بر فناوری‌ها را بتوان بیشتر تضمین کرد. کتاب بر ضرورت ایجاد «زیرساخت اخلاقی» در برابر قدرت‌های بزرگ فناوری تأکید می‌کند و در پی آن است تا الگویی از حکمرانی اخلاق محور و مشارکتی در عصر دیجیتال ارائه کند که نه تنها از عدالت و شفافیت دفاع می‌کند، بلکه توان گفت‌وگو و انتقاد را نیز در دل ساختارهای فناورانه هم حفظ می‌کند. هدف این کتاب در اصل کمک به تفکر درباره راه‌های جایگزین وضعیت کنونی دموکراسی، با تمرکز بر بنیان‌های اخلاقی و دانش ضمنی یا فرضیات اخلاقی است که در نحوه درک و تفسیر ما از دموکراسی نهفته‌اند. این اثر تلاش می‌کند، انگیزه‌ای برای تأمل و گفت‌وگو فراهم کند چرا که پرداختن به درکی که از دموکراسی داریم بر واقعیتی که به‌نام دموکراسی می‌شناسیم، تأثیر خواهد گذاشت. عقب‌نشینی محسوس دموکراسی در جهان امروز به خوبی نشان می‌دهد دموکراسی به گروگان انواع خودکامگی‌ها و عوام‌زدگی‌های فنی^۱ درآمده است که هر یک به نحوی از انقلاب الگوریتمی حمایت می‌کنند.

پیشنهاد کتاب «دموکراسی الگوریتمی» آن است که مفهوم فعلی دموکراسی، که امروز با چالش دموکراسی الگوریتمی روبه‌رو شده، باید توسعه یابد تا بار دیگر بتوان آن را در بسترهای زندگی اجتماعی واقعی جای داد. به‌ویژه، نیاز است تجسم نهادی دموکراسی در دل جامعه مدنی بازسازی و تقویت شود. باید از یکی‌انگاری سیاست با دموکراسی دست برداشت، مشارکت را دوباره زنده کرد و مانع استعمار جامعه مدنی شد که تحت سلطه اقتصاد و سیاستی است که توسط کلان‌داده‌ها و فناوری‌های بزرگ هدایت می‌شوند. محور اصلی این کتاب آن است که با تفکر در قالب «دموکراسی دوسویه» می‌توان در برابر موج فزاینده استعمار الگوریتمی ایستاد و فضاهای نوینی برای مشارکت مدنی گشود. بدین ترتیب، جامعه مدنی و دولت قادر خواهند بود تمام ظرفیت‌های نهفته در خودمختاری اخلاقی و مشارکت شهروندی را به شکوفایی برسانند. رویکرد اتخاذ شده نگاهی انتقادی است که مرزهای سخت‌گیرانه معمول بین نظریه‌های

هنجاری و نظریه‌های تجربی را در حوزه دموکراسی می‌شکند. این رویکرد تلاش می‌کند تا چشم‌اندازها و روش‌شناسی‌های مختلف را در قالب نوعی اخلاق دموکراسی با هم تلفیق کند. همان‌قدر که به‌کارگیری مستقیم و بدون تأمل خواسته‌های اخلاقی درباره چگونگی باید بودن دموکراسی، جزم‌گرایانه است، انکار هرگونه ارزیابی نیز دگماتیک محسوب می‌شود، زیرا خود انکار نیز نوعی ارزیابی است؛ و پذیرش صرفاً افق تجربی به‌عنوان تنها شیوه تبیین واقعیت، تحمیلی یک‌جانبه به‌شمار می‌رود. در این کتاب، «دموکراسی الگوریتمی» نه به‌عنوان مدلی نوین برای دموکراسی، بلکه به مثابه امتداد منطقی نوعی تخصص‌سالاری در نظر گرفته می‌شود؛ جهشی کیفی در پیوند دانش و قدرت که برای کاربرد الگوریتم‌ها ارزشی خنثی قائل است. بنابراین، ضرورتی ندارد این پدیده را شکل جدیدی از دموکراسی بدانیم. در پاسخ به این پرسش که آیا با مدل‌های نوینی از دموکراسی روبه‌رو هستیم، باید گفت: خیر؛ ما هم‌چنان با همان الگوهای نخبه‌گرایانه و تمرکزگرا مواجهیم، با این تفاوت که اکنون این الگوها با هوش مصنوعی تقویت شده‌اند. آن‌چه با آن روبه‌رو هستیم، فهم تازه‌ای از دموکراسی نیست، بلکه نتیجه منطقی انکار برابری انسانی و خودمختاری اخلاقی و سیاسی شهروندان است.

دموکراسی دیجیتال، چنانچه صرفاً به‌عنوان ابزاری برای تقویت مشارکت سیاسی و مدنی به‌کار رود، می‌تواند سودمند باشد. اما آنچه اکنون در حال وقوع است، جایگزینی تدریجی خودمختاری اخلاقی و سیاسی انسان‌ها با منطق الگوریتم‌هاست و این مسئله‌ای کاملاً متفاوت است. این روند در بخش‌های متعددی از نظام‌های دموکراتیک رسوخ کرده: از انتخابات و تصمیم‌سازی‌ها گرفته تا شکل‌گیری افکار عمومی و سیاست‌گذاری‌های کلان. ما با ابزاری کمکی برای خدمت به شهروندان روبه‌رو نیستیم، بلکه با روندی جایگزین‌ساز مواجهیم که در آن، انسان جای خود را به داده می‌دهد؛ فرایندی که از سوی گروه‌هایی محدود، به نمایندگی از شرکت‌هایی که این الگوریتم‌ها را تجاری‌سازی می‌کنند و از رهگذر آن، قدرت و کنترل را در دست می‌گیرند، هدایت می‌شود. در سراسر این کتاب، به تحلیل شیوه‌هایی پرداخته شده که طی آنها انسان‌ها به تدریج جای خود را به داده داده‌اند؛ وضعیتی که نویسندگان آن را «دموکراسی الگوریتمی» می‌نامند.

افزونه‌های انتخابگر و ویراستار

پیشینه

در سال ۱۹۶۲، مدیر مؤسسه انتقال اطلاعات در آکادمی علوم روسیه در مسکو (الکساندر خارکویچ)، مقاله‌ای در مجله «کمونیست» منتشر کرد که در آن شبکه‌ای رایانه‌ای برای پردازش اطلاعات و کنترل اقتصاد پیشنهاد شده بود. او در واقع پیشنهادی برای ساخت شبکه‌ای

نویسنده کتاب کیست؟

دومینیکر گارسیا مارثا، استاد اخلاق در دانشگاه «خائومه اول» (Universitat Jaume I) در کاستلیون است و در آنجا سابقه مدیریت گروه فلسفه و جامعه‌شناسی و همچنین معاونت ارتباطات دانشگاه را در کارنامه دارد. وی که دارای مدرک دکتری فلسفه از دانشگاه والنسیا است، مطالعات تکمیلی خود را در زمینه سیاست در فرانکفورت (آلمان) و در حوزه اخلاق کسب‌وکار در دانشگاه سنت گالن (سوئیس)، دانشگاه نوتردام (ایالات متحده) و کالج مگی در دانشگاه اولستر (ایرلند شمالی) پیگیری کرده است. علایق پژوهشی او شامل دموکراسی مشورتی، اخلاق کسب‌وکار، طراحی نهادی، جامعه مدنی و هرمنوتیک انتقادی است. از جمله آثار برجسته او می‌توان به این عناوین اشاره کرد: *اخلاق عدالت: یورگن هابرماس و اخلاق گفتگویی* (۱۹۹۲)، *نظریه دموکراسی* (۱۹۹۳)، *اخلاق کسب‌وکار: از گفتگو تا اعتماد* (۲۰۰۴)، *تلفیق دیدگاه اخلاقی: روش‌ها، موارد و سطوح در کسب‌وکار و مدیریت* (۲۰۰۵)، با همکاری مارتین بوشر و هانس دِ گیر، *اخلاق شرکت: کلیدهایی برای فرهنگ سازمانی نوین* (۱۹۹۴)، با همکاری آدلا کورتینا، خسوس کونی‌یل و آگوستین دومینگو، *شرکت مسئولیت‌پذیر اجتماعی: اخلاق و شرکت* (۲۰۰۳) و *عقلانیت عمومی و اخلاق‌های کاربردی* (۲۰۰۳)، ویرایش مشترک با آ. کورتینا. او همچنین نویسنده مقالات متعددی است که بر رابطه میان اخلاق، سیاست و اقتصاد تمرکز دارند.

مشابه اینترنت امروزی برای نیازهای حکمرانی الگوریتمی مطرح کرد. این موضوع نگرانی‌هایی جدی در میان تحلیلگران سیاست ایجاد نمود. به‌ویژه، آرتور ام. شلزینگر جونیور هشدار داد که: «تا سال ۱۹۷۰، اتحاد جماهیر شوروی ممکن است دارای فناوری تولیدی کاملاً نوینی باشد که شامل سامانه‌هایی از صنایع است که توسط رایانه‌های خودآموز در حلقه‌های بازخوردی کنترل می‌شوند».

در سال‌های ۱۹۷۱ تا ۱۹۷۳، دولت شیلی پروژه سایبرسین را در دوره ریاست‌جمهوری سالوادور آلنده اجرا کرد. هدف این پروژه ساخت یک سامانه پشتیبانی تصمیم‌گیری توزیع شده برای بهبود مدیریت اقتصاد ملی بود. عناصر این پروژه در سال ۱۹۷۲ برای مقابله موفقیت‌آمیز با فروپاشی حمل‌ونقل ناشی از اعتصاب چهل‌هزار راننده کامیون با حمایت سازمان سیا مورد استفاده قرار گرفت.

همچنین در دهه‌های ۱۹۶۰ و ۱۹۷۰، هربرت الکساندر سایمون استفاده از سامانه‌های خبره را به‌عنوان ابزاری برای عقلانی‌سازی و ارزیابی رفتار اداری مطرح کرد. خودکارسازی فرایندهای مبتنی بر قوانین، هدفی دیرینه برای سازمان‌های مالیاتی در طول دهه‌ها بود که نتایج متغیری در پی داشت. از کارهای ابتدایی در این حوزه می‌توان به پروژه تأثیرگذار «تاکس‌من» اثر تورن مک‌کارتی در آمریکا و پروژه LEGOL اثر رونالد استمپر [۲۹] در بریتانیا اشاره کرد. در سال ۱۹۹۳، دانشمند رایانه پاول کاکشات از دانشگاه گلاسگو و اقتصاددان آلن کاترل از دانشگاه ویک فارست کتابی با عنوان به‌سوی سوسیالیسمی نوین منتشر کردند که در آن ادعا می‌شود امکان طراحی یک اقتصاد دستوری دموکراتیک با تکیه بر فناوری رایانه‌ای مدرن وجود دارد.

منتشر شد!

پایان کار غول‌ها

نوشته: نیکومیل

ترجمه: ابراهیم نقیب‌زاده مشایخ

برای تهیه کتاب با دفتر انجمن انفورماتیک ایران

تماس بگیرید ۶۶۴۱۲۸۶۱

فروش اینترنتی در فروشگاه چاره

www.chare.ir





از انتشارات انجمن انفورماتیک ایران

چاپ سوم

هوش مصنوعی در سال ۲۰۴۱

تهیه کتاب از دفتر انجمن انفورماتیک ایران

۰۲۱-۶۶۴۱۲۸۶۱

قیمت ۲۲۰/۰۰۰ تومان

درگذشت مایکل رابین

دانشمند برجسته علوم رایانه



استیون کلینی^۳ و فصلی از آن را که به پژوهش‌های آلن تورینگ دربارهٔ «اعداد محاسبه‌پذیر» اختصاص داشت مطالعه کرد. رابین می‌گوید:

«با خواندن مقالاتی پیرامون محاسبه‌پذیری و مشاهده ظهور رایانه‌ها، دریافتم که این فناوری نوین، نوع جدیدی از ریاضی را پدید خواهد آورد که حوزهٔ نوظهور رایانش را تبیین خواهد کرد. از این‌رو، تصمیم گرفتم که در این تحول سهمی داشته باشم». او در سال ۱۹۵۳ مدرک کارشناسی ارشد خود را دریافت کرد.

رابین در سال ۱۹۵۴ برای ادامه تحصیل در مقطع دکتری به دانشگاه پرینستون رفت و زیر نظر آلونزو چرچ به تحصیل پرداخت. او در ماه جون ۱۹۵۷ مدرک دکتری خود را دریافت کرد و سپس در تابستان همان سال، همراه با همکلاسیش در پرینستون، دانا اس‌اسکات در برنامه پژوهشی شرکت آی‌بی‌ام در نیویورک شرکت جست. در همانجا بود که آن دو نخستین پیش‌نویس مقاله خود با عنوان «ماشین‌های خودکار متناهی^۴ و مسائل تصمیم‌گیری آن‌ها» را تدوین کردند. این مقاله، مفهوم ماشین‌های خودکار متناهی

مایکل اوسر رابین در مصاحبه‌ای در سال ۲۰۰۹ گفته بود: «ایده‌های نظری، تأثیری واقعی بر زندگی دارند». او نزدیک به هفتاد سال را صرف اثبات این مدعا کرد.

مایکل رابین در ۱۴ آوریل ۲۰۲۶ در سن ۹۴ سالگی در اورشلیم درگذشت. دستاوردهای او تقریباً تمام جنبه‌های حوزه رایانش را در بر می‌گیرد.

مایکل رابین در ۱ ستمبر ۱۹۳۱ در شهر برسلاو در آلمان که امروزه وروتسواف در لهستان نامیده می‌شود به دنیا آمد. در سپتامبر ۱۹۳۵ به همراه خانواده‌اش به فلسطین تحت قیمومیت گریخت و بدین ترتیب از مهلکه رژیم نازی جان سالم به در برد. خانواده او در شهر بندری حیفا ساکن شدند.

رابین از همان سنین پایین شیفتهٔ ریاضیات بود. او برخوردش در ۱۱ سالگی با دو دانش آموز کلاس نهم را به یاد می‌آورد که در تلاش برای حل یک مسئلهٔ هندسه بودند. با آن که او هرگز هندسه نخوانده بود توانست آن مسئله را حل کند. او می‌گوید «این واقعیت که می‌توان صرفاً با تکیه بر تفکر، گزاره‌هایی را درباره دنیای واقعی خطاها، دایره‌ها و فاصله‌ها اثبات کرد، چنان تأثیر عمیقی بر من گذاشت که تصمیم گرفتم به مطالعه ریاضیات بپردازم.

رابین در سال ۱۹۴۸ از مدرسه عبری رئالی در حیفا فارغ‌التحصیل شد و در طول جنگ اعراب و اسرائیل در سال ۱۹۴۸ مانند تمام همسالانش به ارتش فراخوانده شد. آبراهام فرانکل، ریاضیدان، که استاد ریاضیات در اورشلیم بود با فرماندهی ارتش تماس گرفت و اجازه رابین را برای ترک خدمت در ارتش و ادامهٔ تحصیل در دانشگاه به دست آورد.

رابین در دانشگاه عبری^۱، کتاب «مقدمه‌ای بر فراریاضیات^۲» اثر

3- Stephen Cole Kleene

4- Finite Automata

1- Hebrew University

2- Metamathematics

در همان دوران، رابین به حوزه رمزنگاری روی آورد. او در مقاله سال ۱۹۷۹ خود درباره، الگوریتم امضای رابین، با تکیه بر همان مبانی نظریه اعداد، نخستین سامانه رمزنگاری نامتقارن را معرفی کرد که امنیت آن به طور اثبات پذیری معادل با دشواری تجزیه اعداد صحیح بود، ویژگی‌ای که در الگوریتم RSA وجود ندارد.

رابین در سال ۱۹۸۱ به عنوان استاد علوم رایانه به هیئت علمی دانشگاه هاروارد پیوست و این سمت را تا سال ۲۰۱۲ که به مقام استاد پژوهشی رسید، در اختیار داشت. او در دوران حضورش در هاروارد، همچنین در بسیاری دانشگاه‌های معتبر دیگر نظیر دانشگاه ییل، دانشگاه کالیفرنیا در برکلی، ام‌آی‌تی، دانشگاه پاریس، دانشگاه کلمبیا، موسسه فناوری فدرال زوریخ (ETH) و کینگز کالج لندن سمت‌های علمی مهمان داشت. او از سال ۱۹۸۰ تا ۱۹۹۴ عضو کمیته مشورتی علمی آی‌بی‌ام بود و در بهار سال ۲۰۰۹ به عنوان پژوهشگر مهمان به گوگل پیوست.

رابین علاوه بر دریافت جایزه تورینگ در سال ۱۹۷۶، دستاوردهای علمی و جوایز متعدد دیگری را نیز به دست آورده است که از آن جمله می‌توان به موارد زیر اشاره کرد:

● جایزه «پاریس کانلاکیس» انجمن ماشین‌های رایانشی (ACM) در سال ۲۰۰۴

● عضو برجسته انجمن بین‌المللی پژوهش‌های رمزنگاری (IACR) در سال ۲۰۱۰

● جایزه «ادزگر دایکسترا» در زمینه رایانش توزیعی در سال ۲۰۱۵

● جایزه «بهترین استاد» از مؤسسه کورانت در سال ۱۹۷۰

● عضویت در آکادمی هنرها و علوم آمریکا در سال ۱۹۸۲

● عضویت در آکادمی علوم فرانسه به عنوان عضو وابسته خارجی در سال ۱۹۹۵

● عضویت در انجمن فلسفه آمریکا در سال ۱۹۸۸

دستاوردهای علمی رابین، بسیار گسترده و تقریباً تمام جنبه‌های حوزه رایانش را در بر می‌گیرد.

CACM, June 2026, Vol. 64, no. 6

منبع:

غیرقطعی^۵ را صورت‌بندی می‌کرد و نشان می‌داد که این ماشین‌ها از نظر محاسباتی با ماشین‌های قطعی هم‌ارز هستند. آنها همچنین هم‌ارزی ماشین‌های تورینگ چندنواره با ماشین‌های تک‌نواره را نشان دادند. این دو پژوهشگر به پاس همین دستاورد درخشان، در سال ۱۹۷۶ برنده یازدهمین جایزه تورینگ شدند. در متن تقدیرنامه آنها چنین آمده است: «معرفی ایده ماشین‌های غیرقطعی بسیار ارزشمند بوده و همواره الهام بخش پژوهش‌های بعدی در این حوزه گردیده است». پیش از ترک پرینستون، کورت گودل رابین را به مؤسسه مطالعات پیشرفته دعوت کرد و این دو در آنجا تا سال ۱۹۵۸ هر هفته با هم ملاقات می‌کردند. رابین آن جلسات را «فوق العاده روشنگر» توصیف کرده است. پس از آن رابین به برنامه پژوهشی تابستانی آی‌بی‌ام بازگشت و تمرکزش را از سؤال «چه چیزی محاسبه‌پذیر است؟» به این که محاسبه چقدر می‌تواند پرهزینه باشد، تغییر داد. او ثابت کرد که برای هر تابع محاسبه‌پذیر، تابع دیگری وجود دارد که می‌توان اثبات کرد محاسبه آن دشوارتر است و بدین ترتیب، یک سلسله مراتب واقعی از دشواری محاسباتی، مستقل از سخت‌افزار یا زبان برنامه‌نویسی ایجاد کرد. درک کلیدی این بود که پیچیدگی به‌طور ذاتی درون هر کار وجود دارد. بسیاری به انتشار مقاله رابین در سال ۱۹۶۰ به عنوان تولد چیزی که اکنون نظریه پیچیدگی نامیده می‌شود، اشاره می‌کنند.

رابین در سال ۱۹۵۸ به عنوان عضو هیئت علمی به دانشگاه عبری بازگشت و به پاس تأثیرگذاری در حوزه‌ای که در حال شکل‌گیری به‌عنوان رشته علوم رایانه بود، به سرعت ارتقاء مقام یافت. او در سال ۱۹۶۴ به ریاست مؤسسه ریاضیات منصوب شد و در سال ۱۹۷۰ به همراه الی شامیر، دانشکده علوم رایانه را تأسیس کرد و نخستین مدیر آن شد.

رابین در سال ۱۹۷۵ یک سال فرصت مطالعاتی را در آزمایشگاه علوم رایانه مؤسسه فناوری ماساچوست (ام‌آی‌تی) گذراند. او در آنجا با گری میلر آشنا شد، کسی که رساله دکتریش به بررسی روشی نوین برای تشخیص اول یا مرکب بودن یک عدد، بدون نیاز به تجزیه آن، می‌پرداخت. این روش بر یافتن اعدادی استوار بود که رابین بعدها آنها را «شاهد» نامید، اعدادی که مرکب بودن عدد موردنظر را اثبات می‌کردند.

رابین دریافت که تعداد «شاهد» بسیار بیشتر از چیزی است که میلر تصوّر می‌کرد. بدین ترتیب، او توانست آزمون میلر را که روشی کند و قطعی بود به آزمونی سریع و احتمالات محور تبدیل کند که قادر بود اول بودن یک عدد را با احتمالی بسیار بالا تعیین نماید. این رویکرد جدید که «آزمون اول بودن میلر - رابین» خوانده می‌شود در سال ۱۹۸۰ انتشار یافت و بعدها به مبنایی برای انتخاب اعداد اول برای استفاده در الگوریتم رمزنگاری کلید عمومی (RSA) تبدیل گردید.

5- Nondeterministic

برای دریافت اسلاید
سخنرانی‌های انجمن به
وبگاه انجمن مراجعه کنید
WWW.isi.org.ir

مدیریت فرایند تولید نرم افزار بخش هشتم: اجرای کار

سیدعلی آذرکار

مؤسسه ارزیابی کسب و کار آینده پژوه هوشمند

پست الکترونیکی: a.azarkar@sfbvi.com

(۱) مقدمه

- همسانی برای همه پروژه‌ها استفاده شود.
- فعالیت‌ها و پیش‌نیازهای فراموش شده، شناسایی و کشف می‌شود.
- در استفاده از منابع صرفه‌جویی می‌شود.
- شفافیت فرایند افزایش پیدا می‌کند.

در قسمت پیشین به بیان طرح‌های لازم برای اجرای بهتر یک پروژه نرم‌افزاری پرداخته شد. در ادامه، در این قسمت به طرح استقرار، که از طرح‌های مهم و پرریسک تولید یک نرم‌افزار است، پرداخته شده و به اجمال توصیف می‌شود.

(۱-۲) راهبردهای استقرار

در یک طیف کلی، استقرار می‌تواند به صورت انفجاری (Big-Bang) یا تدریجی (Incremental) انجام شود. هر کدام از این دو راهبرد مزایا و معایب خاص خود را داشته و این مدیریت پروژه است که باید حسب شرایط و نوع پروژه، یک راهبرد را انتخاب کرده و برنامه‌ریزی‌ها را بر اساس آن انجام دهد.

در راهبرد انفجاری، کل سامانه به یکباره به محیط عملیاتی منتقل می‌شود. این راهبرد نوعاً برای سامانه‌های کوچک نرم‌افزاری در سازمان‌های کوچک که ریسک شکست پروژه در آن‌ها کم بوده یا قابل تحمل است، توصیه می‌شود. از میان ویژگی‌های این راهبرد، موارد زیر قابل توجه است:

- انتقال یکباره همه کاربران به محیط عملیاتی جدید
- کوتاه بودن دوره گذار (یا انتقال)
- ریسک بالا
- پایان سریع‌تر پروژه
- نیازمند انجام آزمون‌های جامع و کافی سامانه (محصول) و نیز آمادگی کامل زیرساخت

(۲) طرح استقرار

استقرار یک سامانه نرم‌افزاری و اجرای آن در محیط عملیاتی، بدون شک حساس‌ترین و مهم‌ترین گام در چرخه عمر آن محسوب می‌شود. به همین دلیل، باید به جزء جزء مواردی که باید در زمان طرح‌ریزی این گام لحاظ شده تا استقرار بدون کم‌ترین مشکل به سرانجام برسد، توجه شود. روشن است که هر چه سامانه از حیث پیچیدگی، حساسیت و اندازه بزرگ‌تر باشد، این گام هم با ریسک‌های بیشتری مواجه بوده و نیازمند دقت و توجه ویژه‌ای خواهد بود.

هدف از تدوین طرح استقرار، تعیین فعالیت‌ها و منابع لازم برای استقرار سامانه نرم‌افزاری در محیط عملیاتی است. این طرح، دربرگیرنده موضوعاتی مانند راهبردها، منابع مورد نیاز، زمان‌بندی، نقاط عطف و نقاط مهم تصمیم‌گیری است. هدف این طرح، ایجاد هماهنگی بین همه عوامل درگیر (از مدیران پروژه تا تیم‌های فنی تا مدیران ارشد و تصمیم‌ساز تا کاربران نهایی) در فرایند استقرار است

در صورت تدوین این طرح:

در راهبرد تدریجی، سامانه یا مؤلفه‌ای آن به تدریج استقرار پیدا

- مراحل استقرار مستند و مکتوب شده و می‌تواند در آینده به شکل

● برنامه زمان‌بندی: این سرفصل شامل تعیین تاریخ‌هایی است که فعالیت‌های استقرار باید انجام شود. زمان‌بندی باید با توجه به شرایط عملیاتی و کسب‌وکاری سامانه انجام شود. زمان‌بندی باید به نحوی باشد که کم‌ترین اختلال در فعالیت‌های روزمره کارفرما ایجاد شود. این موضوع به‌ویژه در شرایطی که سامانه درگیر پاسخ‌گویی به درخواست‌های کاربران خارجی (مانند مشتریان و ارباب‌رجوع) است، حساسیت خاص خود را دارد و به‌راحتی می‌تواند نقطه شکست سامانه و پروژه باشد. بنابراین همه تمهیدات به منظور تنظیم زمان‌بندی استقرار با فعالیت‌های کسب‌وکاری اندیشیده شده باشد، به نحوی که امکان مشارکت و همکاری فعالانه همه عوامل اجرایی (به‌خصوص، کاربران نهایی) میسر و مقدور باشد. مهم‌ترین تاریخ در این زمان‌بندی، تاریخی است که سامانه باید به‌طور کامل عملیاتی شده یا به‌طور کامل جایگزین سامانه پیشین شود.

● راهبرد عقب‌گرد (Roll-Back): به منظور مدیریت بهتر ریسک‌های عملیاتی در زمان استقرار، بهتر است راهبرد/راهبردهای لازم برای زمانی که به هر دلیل استقرار با مشکل مواجه می‌شود، از پیش اندیشیده شده و مدون شده باشد. اهم موضوعاتی که در این بخش باید به آن توجه شود، عبارت است از: (۱) شرایط اجرای راهبرد، (۲) ساختار و سازمان تصمیم‌گیر برای انتخاب راهبرد مناسب؛ (۳) نحوه اطلاع‌رسانی و هماهنگی عوامل پروژه (پیش، حین و پس از استقرار). ● پایش استقرار: این فعالیت به دلیل حساسیت، نیازمند پایش مستمر، به‌ویژه حین و پس از استقرار، است. از موارد مهمی که می‌تواند در فرایند استقرار پایش شود، به این موارد می‌توان اشاره کرد: (۱) شرایط عملیاتی و کسب‌وکاری؛ (۲) شرایط زیرساخت نرم‌افزاری، سخت‌افزاری، ارتباطی و امنیتی اجرایی سامانه؛ (۳) آموزش کاربران؛ (۴) حضور تیم فنی در محل استقرار؛ (۵) ثبت درس‌آموخته‌ها؛ (۶) اصلاح مستندات؛ (۷) دریافت بازخوردهای کاربران؛ (۸) ثبت نرخ خطاها؛ (۹) میزان انطباق کارکرد سامانه با نیازهای کسب‌وکاری؛ (۱۰) میزان انطباق‌پذیری کاربر با سامانه؛ (۱۱) رفع خطاهای موردی.

● پیش‌نیازهای استقرار: در طرح باید پیش‌نیازها و منابع لازم برای استقرار تصریح و توافق شده باشد. علاوه بر پیش‌نیازهای فنی (مهم‌ها) کردن سامانه، ایجاد زیرساخت‌های لازم، انجام آزمون‌های لازم و مواردی از این دست، توجه به ایجاد آمادگی ذهنی و فرهنگی کاربران برای ورود به محیط جدید، بسیار مهم است. ضمن آنکه مواردی مانند تعیین و توافق برنامه زمان‌بندی، حضور نیروهای کارفرما و پیمانکار در زمان استقرار، نحوه اطلاع‌رسانی و هماهنگی بین عوامل از جمله مشتریان، باید مد نظر باشد.

● تهیه بازبینی استقرار: به منظور کنترل همه موارد لازم برای اجرای موفقیت‌آمیز استقرار، پیشنهاد می‌شود بازبینی‌ای برای این منظور توسط پیمانکار تهیه و حین/پس از استقرار به‌صورت پیوسته کنترل شود. مواردی که می‌تواند در این بازبینی گنجانده شود، شامل، ولی نه محدود به این موارد، است: تعیین زمان‌بندی، انجام آزمون‌های

می‌کند. از میان ویژگی‌های این راهبرد، موارد زیر قابل توجه است:

- امکان استقرار بر اساس عواملی مانند واحدهای کسب‌وکاری، گروه کاربران، منطقه جغرافیایی، ... یا ترکیبی از آن‌ها
- امکان شناسایی و اصلاح اشکالات احتمالی سامانه
- ریسک کمتر
- کنترل بهتر فرایند و امکان بهبود آن
- طولانی شدن زمان اجرای پروژه

یک راهبرد میانی، استفاده از دو محیط (محیط پیشین و محیط جدید) است؛ در این حالت، کاربران فعالیت‌های خود را (برای یک بازه زمانی مشخص و توافق شده) در دو محیط هم‌زمان به پیش می‌برند و پس از حصول اطمینان از درستی کارکرد سامانه جدید، سامانه پیشین کنار گذاشته می‌شود.

۲-۲) سرفصل‌های طرح استقرار

سرفصل‌هایی که باید در تدوین طرح استقرار مد نظر قرار گیرد، در ادامه توضیح داده شده است:

● محدوده: باید مشخص شود که قرار است چه محصولی و در چه محدوده‌ای (از حیث سازمانی، کسب‌وکاری، جغرافیایی، ...) استقرار یافته و وارد محیط عملیاتی شود.

● راهبرد استقرار: موضوع انتخاب راهبرد (با توجه به شرحی که آمد) باید در طرح با دقت مطرح و مستند شود؛ چرا که می‌تواند در آینده برای پروژه‌های دیگر هم استفاده شود.

● ساختار سازمانی: استقرار نیازمند درگیری و مشارکت فعالانه عواملی از سمت کارفرما و پیمانکار است. بنابراین، باید یک ساختار سازمانی جامع و فراگیر (از منظر اجرایی، مدیریتی و فنی) ایجاد شده تا نقش‌ها و مسئولیت‌های همه این عوامل در آن به روشنی تعیین و مشخص شده باشد.

● شاخص‌ها و اهداف: به منظور سنجش میزان اثربخشی و موفقیت فرایند استقرار پس از پایان آن، لازم است شاخص‌ها و اهداف مد نظر در طرح گنجانده شود. این شاخص‌ها می‌تواند دربرگیرنده شاخص‌های حوزه فنی (مانند کارایی سامانه) یا شاخص‌های کسب‌وکاری (تسریع در انجام فعالیت‌ها یا بهبود اثربخشی آن‌ها) باشد.

● آزمون‌های پیشینی: استقرار سامانه در محیط عملیاتی باید زمانی انجام شود که تیم پروژه از کارکرد صحیح آن اطمینان کامل داشته باشد. به همین منظور ضروری است انواع آزمون‌هایی که باید پیش از استقرار سامانه در محیط عملیاتی انجام شود، مشخص شده و نتایج آن به‌دقت بررسی و تحلیل شده باشد. از میان این آزمون‌ها، آزمون پذیرش کاربر (UAT) از اهمیت ویژه‌ای برخوردار است.

● ابزارها: اگر در فرایند استقرار از ابزاری خاصی باید استفاده شود، ذکر این ابزار (به همراه مشخصات کامل آن) در طرح ضروری است. روشن است که نحوه تأمین زیرساخت‌های لازم برای استفاده از این ابزار هم باید در طرح پیش‌بینی شده باشد.

نادرست، و خطاهای ناشی از تبدیل و انتقال داده‌ها. توصیه می‌شود که عوامل دست‌اندرکار استقرار (اعم از کارفرما و پیمانکار)، فعالیت‌هایی را هم پس از استقرار، پیش‌بینی کنند. هدف این فعالیت‌ها ارزیابی فرایند و نیز گردآوری و مدون کردن درس‌آموخته‌های استقرار است. از فعالیت‌هایی که می‌تواند پس از استقرار برنامه‌ریزی و اجرا شود، می‌توان به موارد زیر اشاره کرد:

- ارزیابی میزان انطباق‌پذیری کاربر با محیط جدید
- ایجاد میز خدمت به منظور بررسی و رسیدگی سریع به مشکلات/پرسش‌های کاربران
- ارزیابی وضعیت زیرساخت اجرایی و عملیاتی پروژه (از محصول تا زیرساخت‌های پردازشی، امنیتی، ارتباطی، ...)
- تعیین میانگین زمان رفع مشکلات
- بررسی ریسک‌های جدید و شناسایی نشده
- ارزیابی میزان رضایت‌مندی کاربران از سامانه
- ارزیابی میزان بهبود فرایندهای کسب‌وکاری پس از استقرار سامانه
- برپاسازی سازوکار ثبت و رفع مشکلات/خطاها

کارکردی، انجام تبدیل و انتقال داده، صحت‌سنجی داده‌های تبدیل‌یافته و منتقل شده، آموزش کاربران، آزمون امنیتی، برپاسازی سامانه عامل و مدیریت دادگان، برنامه حضور تیم فنی، هماهنگی با تیم تولید برای رفع خطاهای احتمالی در زمان استقرار، برنامه‌ریزی پایش استقرار، به اشتراک‌گذاری برنامه زمان‌بندی با همه عوامل، تعریف کاربران، تنظیم سطوح دسترسی و مجوزها، تنظیم مقادیر اولیه سامانه.

• مدیریت ریسک: اشاره شد که فرایند استقرار همواره با ریسک‌هایی مواجه است. از آنجا که در زمان استقرار، فرصت زیادی برای بررسی و نحوه مقابله با ریسک‌ها وجود ندارد، ضروری است که پیش از آغاز استقرار، این ریسک‌ها شناسایی شده، راهکارهای مقابله با آن‌ها اندیشیده شده و از همه مهم‌تر، ساختار و سازمانی به منظور تصمیم‌گیری در خصوص تصمیم‌گیری‌های لازم در زمان وقوع ریسک ایجاد شده باشد. ریسک‌ها، می‌توانند این مواردی از این دست را در برگیرند: محدودیت منابع، مقاومت در مقابل تغییر، اشکالات زیرساختی، تهدیدهای امنیتی، افت کارایی سامانه، مشکلات عملیاتی و اجرایی، خطا در گردش کار، آموزش ناکافی کاربران و ثبت داده‌های

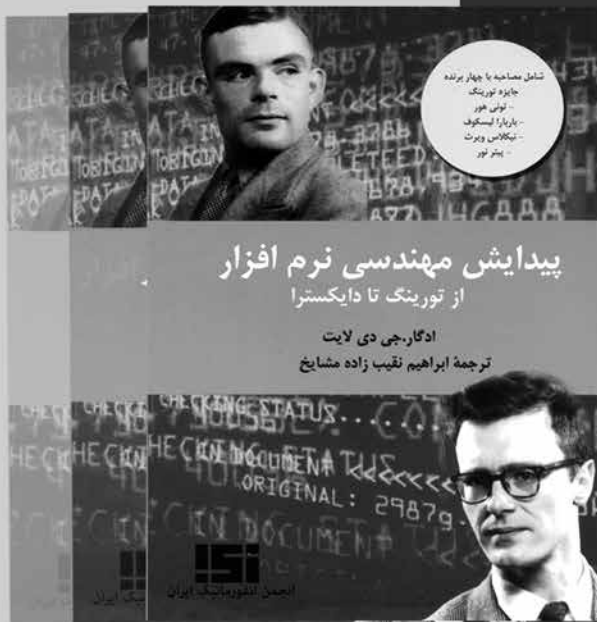
منتشر شد!

پیدایش مهندسی نرم افزار

ترجمه: ابراهیم نقیب‌زاده مشایخ

برای تهیه کتاب با دفتر انجمن انفورماتیک ایران

تماس بگیرید (۰۲۱-۶۶۴۱۲۸۶۱)

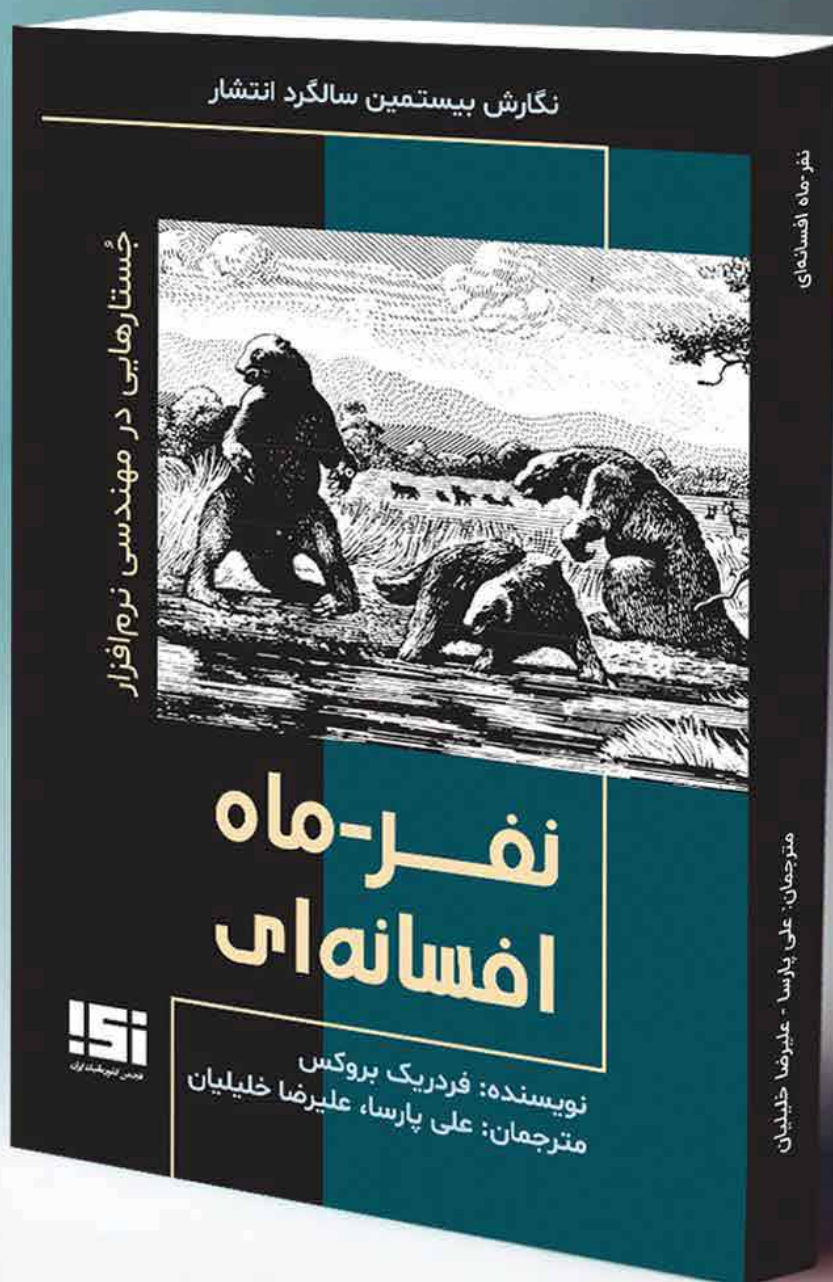




انجمن انفورماتیک ایران
INFORMATICS SOCIETY of IRAN

از انتشارات انجمن انفورماتیک ایران

منتشر شد!



گروه تخصصی هنر و تکنولوژی انجمن انفورماتیک ایران دعوت می کند فراخوان مشارکت در «روز جامعه پروسسینگ تهران ۲۰۲۶»

پدرام صادق بیکی

پست الکترونیکی: sadeghbeyki@isi.org.ir

درباره رویداد

سال ۲۰۲۶، بیست و پنجمین سال فعالیت پروسسینگ (Processing) است؛ پروژه‌ای که از ابتدا با این ایده شکل گرفت که برنامه‌نویسی می‌تواند چیزی فراتر از توسعه نرم‌افزار باشد و به ابزاری برای خلق، تجربه، آموزش و پژوهش تبدیل شود.

به مناسبت بیست و پنجمین سال فعالیت پروسسینگ، روز جامعه پروسسینگ (Processing Community Day) در نقاط مختلف جهان برگزار می‌شود؛ رویدادی جامعه‌محور برای گردهم آمدن هنرمندان، طراحان، برنامه‌نویسان، مدرسان، پژوهشگران، دانشجویان و همه کسانی که از کد به عنوان ابزاری برای ساختن، تجربه کردن و یادگیری استفاده می‌کنند.

روز جامعه پروسسینگ از سال ۲۰۱۷ آغاز شده و تاکنون بیش از صد رویداد در شش قاره برگزار شده است. شهرهایی مانند توکیو، نیویورک، ساوئوپاولو، بنگلور، کیتو، سئول، لس‌آنجلس و شانگهای از میزبانان این رویداد بوده‌اند.

<http://day.processing.org>

درباره پروسسینگ

محیط‌های برنامه‌نویسی Processing و p5.js محیط‌ها و ابزارهایی رایگان و متن‌باز هستند که هنرجویان، دانشجویان، مدرسان، هنرمندان و خالقان سراسر جهان از آن‌ها برای یادگیری، آزمایش و خلق هنر با کد استفاده می‌کنند.



کدنویسی در این فضا الزاماً به معنای برنامه‌نویسی حرفه‌ای یا توسعه نرم‌افزار نیست. یک طراح می‌تواند با الگوریتم فرم بسازد، یک هنرمند از داده برای خلق اثر استفاده کند، یک مدرس کدنویسی را آموزش

دهد و یک دانشجو نخستین تجربه‌های خود را با تصویر و کد شکل دهد.

آثاری که با Processing و p5.js و در فضای کدنویسی خلاق شکل گرفته‌اند، طیف گسترده‌ای از هنر مولد و تجربه‌های تعاملی تا تصویر، صدا، داده، هنر وب، چیدمان و اجرای زنده را در بر می‌گیرند و به موزه‌ها و فضاهای هنری مهم جهان، از جمله موزه هنر مدرن نیویورک، موزه ویتنی، موزه هنر تولدو و موزه تصاویر متحرک نیویورک راه یافته‌اند.

روز جامعه پروسسینگ تهران

تهران در سال ۲۰۲۶ برای نخستین بار میزبان روز جامعه پروسسینگ تهران (Processing Community Day Tehran) خواهد بود. این رویداد در مهرماه ۱۴۰۵ در موزه کامپیوتر ایران و با برنامه‌ریزی گروه هنر و تکنولوژی انجمن انفورماتیک ایران برگزار می‌شود. برنامه روز جامعه پروسسینگ تهران از طریق یک فراخوان عمومی و با مشارکت هنرمندان، طراحان، توسعه‌دهندگان، مدرسان، پژوهشگران، دانشجویان و فعالان حوزه کدنویسی خلاق شکل می‌گیرد.

فراخوان مشارکت

اگر با Processing یا p5.js کار کرده‌اید یا در هر سطحی از کدنویسی خلاق برای خلق اثر، آموزش یا پژوهش استفاده می‌کنید، می‌توانید پروژه، اثر یا تجربه خود را برای حضور در برنامه پیشنهاد دهید. برای مشارکت در نخستین روز جامعه پروسسینگ تهران لازم نیست یک برنامه‌نویس حرفه‌ای یا هنرمند باتجربه باشید. از یک تجربه کوچک و شخصی تا یک پروژه انجام‌شده، یک تجربه آموزشی، پژوهش، اثر هنری یا روشی که در جریان کار خود به آن رسیده‌اید، می‌تواند نقطه شروع یک گفت‌وگوی خوب باشد.

برای حضور در برنامه، چهار مسیر اصلی در نظر گرفته شده است:

- ارائه و سخنرانی (Presentation)
- کارگاه تجربی و آموزشی (Workshop)
- نمایش اثر هنری (Art Showcase)
- مقاله و پژوهش (Documentation)

چه در تهران باشید و چه در شهر یا کشوری دیگر، متناسب با نوع مشارکت خود می‌توانید پیشنهادتان را برای ما ارسال کنید.

۱. ارائه و سخنرانی || Presentation

اگر پروژه، تجربه یا ایده‌ای دارید که گفتن درباره آن می‌تواند برای دیگران مفید یا الهام‌بخش باشد، می‌توانید آن را در قالب یک ارائه کوتاه با دیگران به اشتراک بگذارید. موضوع ارائه می‌تواند شامل موارد زیر باشد:

- معرفی یک پروژه توسعه‌یافته یا در حال توسعه
- تجربه شخصی، آموزشی یا پژوهشی در کدنویسی خلاق
- بررسی یک تکنیک، ابزار یا فرایند کاری
- نمایش پروژه یا مطالعه موردی
- کدنویسی زنده یا ارائه اجرایی
- اجرای زنده صدا و تصویر با کد

نحوه ارائه: حضوری یا ویدئوی ضبط‌شده

زمان پیشنهادی: ۱۰ تا ۲۰ دقیقه

اطلاعات مورد نیاز: چکیده ارائه به همراه تصاویر یا پیوندهای مرتبط در قالب فایل پی‌دی‌اف

۲. کارگاه تجربی و آموزشی || Workshop

اگر تجربه‌ای دارید که فکر می‌کنید بهتر است مخاطبان خودشان آن را امتحان کنند، می‌توانید پیشنهاد برگزاری یک کارگاه را ارسال کنید.

کارگاه می‌تواند برای افراد تازه‌کار طراحی شود یا به موضوعی تخصصی‌تر بپردازد. موضوعات پیشنهادی می‌توانند شامل موارد زیر باشند:

- Processing
- p5.js
- WebGL
- GLSL و شیدرنویسی
- هنر مولد و تعاملی
- ابزارها و فناوری‌های مرتبط با کدنویسی خلاق

نحوه ارائه: حضوری

زمان پیشنهادی: ۱ تا ۲ ساعت

اطلاعات مورد نیاز: عنوان، طرح درس، سطح مخاطبان و پیش‌نیازهای احتمالی در قالب فایل پی‌دی‌اف

۳. نمایش اثر || Art Showcase

اگر با کد اثری هنری ساخته‌اید، می‌توانید آن را برای نمایش در رویداد ارسال کنید.

اولویت انتخاب آثار با کارهاییست که با یکی از برنامه‌های Processing و یا p5.js نوشته شده باشند نمونه آثار قابل ارسال:

- هنر مولد (Generative Art)
- هنر تعاملی (Interactive Art)
- هنر تحت وب (Web Art)
- هنر داده (Data Art)
- ویدئوآرت (Video Art)

زمان برگزاری و مهلت ارسال

مهلت ارسال آثار و پیشنهادهای: ۳۱ شهریور ۱۴۰۵

زمان برگزاری رویداد: مهرماه ۱۴۰۵

محل برگزاری: موزه کامپیوتر ایران

موزه کامپیوتر ایران فضایی برای ملاقات گذشته و آینده فناوری است؛ جایی که تاریخ ابزارهای محاسباتی و دنیای دیجیتال در کنار روایت‌های انسانی آن‌ها قرار می‌گیرد. این موزه با گردآوری ابزارها، اسناد، داستان‌ها و تجربه‌های مرتبط با تاریخ کامپیوتر تلاش می‌کند روایتی زنده و تعاملی از مسیر تحول فناوری ارائه دهد. نخستین روز جامعه پرسوسینگ تهران در موزه کامپیوتر ایران برگزار خواهد شد.

بخشی از روز جامعه پرسوسینگ تهران باشید.

اگر با کد اثری خلق کرده‌اید، چیزی یاد گرفته‌اید، تجربه‌ای دارید که می‌تواند به کار دیگران بیاید یا ایده‌ای دارید که فکر می‌کنید جای آن در چنین جمعی خالی است، می‌توانید بخشی از نخستین روز جامعه پرسوسینگ تهران باشید.

منتظر ایده‌ها، تجربه‌ها، پژوهش‌ها و پروژه‌های شما هستیم. برای ثبت‌نام از طریق لینک زیر اقدام فرمایید.

<https://www.atgiran.ir/PCDTehran2026>

سلب مسئولیت: «روز جامعه پرسوسینگ» یک ابتکار جهانی از سوی بنیاد پرسوسینگ است. رویداد «روز جامعه پرسوسینگ تهران» به‌صورت مستقل توسط گروه هنر و تکنولوژی انجمن انفورماتیک ایران برگزار می‌شود و هیچ‌گونه ارتباط سازمانی با بنیاد پرسوسینگ ندارد.

● سایر پروژه‌های هنری بر پایه کد

هنرمندانی که آثار خود را با Processing یا p5.js ساخته‌اند، می‌توانند اثر خود را در وبگاه OpenProcessing نیز ثبت کنند و با ارسال پیوند در فرم ثبت‌نام، از فرصت انتخاب و نمایش اثر در Curation این سایت بهره‌مند شوند.

۴. مقاله و مستند Documentation ||

اگر فعالیت شما بیشتر پژوهشی، تحلیلی یا مبتنی بر مستندسازی یک تجربه است، می‌توانید نوشتار یا مطالعه موردی خود را برای این بخش ارسال کنید.

موضوعات این بخش می‌توانند شامل موارد زیر باشند:

- پژوهش در زمینه هنر، فناوری و کدنویسی خلاق
- مطالعه موردی یک موضوع یا پروژه
- بررسی فرایند شکل‌گیری یا توسعه یک اثر
- مقاله و متون آموزشی
- تحلیل یک ابزار، تکنیک یا روش
- گزارش پژوهشی
- مستندسازی یک تجربه
- ویدئوی پژوهشی یا آموزشی

آثار متنی می‌توانند به‌صورت فایل پی‌دی‌اف و محتوای ویدئویی از طریق پیوند ارسال شوند. مقاله‌های منتخب این بخش، فرصت انتشار در نشریه «گزارش کامپیوتر» را خواهند داشت.

برای مشارکت لازم نیست تهران باشید

روز جامعه پرسوسینگ تهران فقط برای کسانی نیست که امکان حضور در محل برگزاری را دارند. هنرمندان و شرکت‌کنندگان خارج از تهران یا خارج از کشور نیز می‌توانند متناسب با نوع اثر یا ارائه خود، از طریق ارسال محتوای دیجیتال یا ویدئوی ضبط‌شده در رویداد مشارکت کنند.



من یک کامپیوتر مستقل هستم، اما احساس تنهایی می‌کنم و می‌خواهم عضوی از یک شبکه محبوب باشم

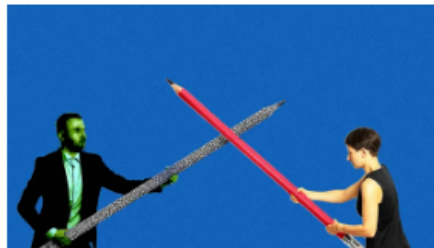
EMMA MADDEN CULTURE JUL 29, 2020 6:38 AM

More Typos, Fewer Em Dashes: Writers Are Creating an Anti- AI 'Literary Counterculture'

Novellists, journalists, and power LinkedIn posters are embracing first-person narratives and idiosyncrasies to avoid being mistaken for chatbots.



باز فکرم



غلط‌های املائی بیشتر، خط تیره‌های ام کمتر:
نویسندگان در حال ایجاد یک «ضدفرهنگ
ادبی» ضد هوش مصنوعی هستند

رمان‌نویس‌ها، روزنامه‌نگارها و کاربران قدرتمند
لینکدین، برای اینکه با ربات‌های چت اشتباه گرفته
نشوند، به روایت‌های اول شخص و ویژگی‌های خاص
خودشان روی آورده‌اند.

اما مدن

تردید در مردم سالاری رقمی
بذل کاری‌های جنبش نویسندگان علیه
هوش مصنوعی

از وبگاه بانی فیلم منتشره در ۱۴۰۵/۵/۱۰
wired.com به نقل از:

WIRED

سید ابراهیم ابطاحی

استادیار دانشکده مهندسی کامپیوتر دانشگاه صنعتی شریف

پست الکترونیکی: abtahi@Sharif.edu

هوش مصنوعی به پا خاسته‌اند؛ موضوعی که سایت ادبی وایرد
«Wired» به آن پرداخته است.
سال گذشته لورا بروک روبنسون - رمان‌نویس - هنگامی که نوشتن

رمان‌نویسان، روزنامه‌نگاران و نویسندگان مطالب پر بازدید برای این
که با «چت‌بات‌ها» اشتباه گرفته نشوند، با ایجاد جنبش روایت‌های
اول شخص و به کارگیری ویژگی‌های منحصر به فرد بر ضد نوشتارهای

کتاب «عشق یک الگوریتم است»، داستانی درباره عشق بین بنیان‌گذار یک نرم‌افزار هوش مصنوعی دوستیابی و موسیقی‌دانی که تلاش می‌کرد تا در عصر مدل‌های بزرگ زبانی سرپا بماند را آغاز کرد، تلاش مضاعفی به خرج داد تا نوشتارش شبیه هوش مصنوعی نشود. او جملاتی را به کار برد که پیش از ظهور مدل‌های زبانی هوش مصنوعی از نوشتن آنها اجتناب می‌کرد و تلاش کرد تا از استعاره‌های مبهم و صور خیالی اجتناب کند و مسیر دشوارتری را برگزید تا از دو خط تیره، برای نوشتن عبارات توصیفی و فهرست سه‌گانه (سه واژه یا صفت برای توصیف بیشتر؛ مانند: فرهنگی، هنری، اقتصادی) کمتر استفاده کند.

روبنسون خود می‌گوید که از زمان همگانی شدن مدل‌های زبانی بزرگ هوش مصنوعی، مطالبش را به‌عمد برخلاف جهت چیزی که او سبک پیش‌فرض هوش مصنوعی می‌نامد، نوشته و نوشته‌هایش بهتر شده‌اند.

او می‌گوید: «دریافتم که کمتر مثل ماشین می‌نویسم. تمایل زیادی پیدا کرده‌ام تا غلط‌های کوچک را تصحیح نکنم؛ مثل چشمک زدن مخفیانه به خواننده می‌ماند تا به او بفهمانم که یک انسان پشت این واژگان قرار دارد. مثلاً اگر عبارتی از خودم اختراع کنم چه می‌شود؟ یا اگر از علایم نگارشی به‌طور نابه‌جا استفاده کنم چه اتفاقی می‌افتد؟» روبنسون جزو گروه گسترده‌تری از نویسندگان است که از سبکی در نوشتار خود استفاده می‌کند که می‌توان آن را سبک نوشتن ضد هوش مصنوعی نامید. در زمینه‌های نقد و نویسندگی خلاق، همگرایی در میان نویسندگان شکل گرفته است تا به مجموعه‌ای از انتخاب‌ها برسند و به‌طور ضمنی در دام‌های نوشتارهای تولید شده توسط هوش مصنوعی گرفتار نشوند.

در این سبک بیشتر روش نفی، خلاف جریان و ضدیت به کار گرفته می‌شود؛ یعنی پرهیز از روش‌های آشکاری که هوش مصنوعی استفاده می‌کند؛ مانند دو خط تیره بلند، فهرست سه‌تایی توصیفی، ساختار جمله مجهول.

با این وجود نویسنده سعی می‌کند تا ویژگی‌های منحصر به فرد خود را حفظ کند و با بهره‌گیری از بازگویی‌های مشخص اول شخص و تشبیهات ادبی نامتعارف، تلاش یک انسان برای نوشتن را نشان دهند. در تمامی دنیای نشر، اتهاماتی بر کتاب‌ها وارد می‌شود، مبنی بر استفاده از هوش مصنوعی برای تولید کتاب که نگرانی‌های زیادی را موجب شده است. در ماه مارس، انتشارات «هاچت» هنگام انتشار رمان ترسناک «دختر خجالتی» نوشته میا بالارد در فضای مجازی با گمانه‌زنی‌هایی مبنی بر این که نویسنده به‌طور فزاینده‌ای از هوش مصنوعی استفاده کرده است، مواجه شد و کتاب را از فهرست انتشار خارج کرد؛ ادعایی که نویسنده آن را رد کرده است. بالارد که در حال

حاضر آثار جدیدی را برای انتشار دارد، ناچار شده است روش خود را تغییر دهد. او می‌گوید: «این که نوشته‌هایم اصیل و واقعی باشند، برایم یک اولویت است». وی اضافه می‌کند که آثارش را خودش ویرایش خواهد کرد؛ به ادعای او ویراستاری که استخدام کرده بود، کتاب را از طریق نرم‌افزارهای هوش مصنوعی ویرایش کرده است. بالارد تأکید می‌کند: «این اولین بار بود که چنین چیزی به ضرر من تمام شد». انتشارات «هاچت» هنوز به درخواست‌ها برای اظهارنظر پاسخی نداده است.

همین اتهام به استیون روزن‌بام- نویسنده کتاب «آینده حقیقت»- کتابی که به بررسی دگرگونی واقعیت توسط هوش مصنوعی می‌پردازد، وارد شد. گفته می‌شد که او از چندین نقل قول ساخته شده توسط هوش مصنوعی در اثر خود استفاده کرده است؛ اتهامی که بعدها توسط خود او تایید شد.

النا ساویچ- نویسنده و مدرس دانشگاه از استرالیا- که نوشتارش را با الهام از ادبیات عصر ویکتوریا با جملات پیچیده می‌نویسد، گفته است که «متأسفانه قربانی افکار مربوط به چت‌بات‌ها» شده است. او می‌گوید: «هوش مصنوعی چندین قواعد ساختاری را که برایم بسیار عزیز بودند از من دزدیده است. دلم برای استفاده از آنها بسیار تنگ شده است. دلم برای خط تیره و فهرست سه‌گانه تنگ شده است. حالا چه کنیم؟ قانون دوتایی؟ پنج‌تایی؟ هفت‌تایی؟ یا ده‌تایی؟»

این موضوع گسترش سبک‌های نوشتاری ضد دامنه زبانی هوش مصنوعی، روش‌های نگارش فردی و سردبیری را نیز دگرگون کرده است. ریچارد فیشر- استاد افتخاری کالج دانشگاهی لندن و دبیر مجله آنلاین «آیون»- می‌گوید که ورود هوش مصنوعی به حوزه نویسندگی او را بر آن داشته است نویسندگان را تشویق کند تا «انسانی‌تر» و به شکل محاوره‌ای‌تر، مثل زبان عمومی انسان‌ها بنویسند. او می‌گوید: «نباید سطح زمخت و ناصاف نوشتار را سیقل دهیم؛ زیرا هوش مصنوعی نمی‌تواند به تند و تیزی و منحصر به فردی کلام انسانی بنویسد. شاید برخی نویسندگان نخواهند روش کار خود را عوض کنند اما جاهایی ممکن است، مجبور به این کار شوند».

در جولای سال ۲۰۲۴ و در هنگامه فراگیر شدن مدل‌های زبانی هوش مصنوعی، مانو ساندارسان به سمت مدیریت محتوای تحریریه نشریه نقد ادبی «پیچفورک» منصوب شد. او که سابقه وبلاگ‌نویسی را در کارنامه خود داشت، نویسندگان را تشویق می‌کرد تا به جای پیروی از قواعد کلیشه‌ای رایج از «صدای منحصربه‌فرد خودشان» بهره ببرند؛ نکته‌ای که به گفته خودش، با مرور بر گذشته حرفه‌ای، «شاید تحت تأثیر همین موج هوش مصنوعی بوده باشد». او و دیگر سردبیران آن زمان یادداشتی برای نویسندگان فرستادند و در آن تأکید کردند که در نوشتن متون نباید از هوش مصنوعی استفاده شود

و آنها را به استفاد از زاویه دید شخص اول تشویق کردند. هوش مصنوعی بر فرایند مطالب سفارشی نیز تاثیر گذاشته است. فیلیپ پیکاردی - سردبیر یکی از نشریات تخصصی - هنگامی که در ماه می صندوق دریافت پیشنهادها را باز کرد، با حدود هزار مطلب ارسالی مواجه شد. او پس از بررسی‌های تخصصی متوجه شد که اکثر آنها با هوش مصنوعی تولید شده بودند. پیکاردی می‌گوید: «نکته کلیدی در این بود که ما در مجله به تجربه انسانی اول شخص نیاز داشتیم». از آن زمان او نویسندگان را به جسورانه‌تر نوشتن از نظر ساختار تشویق کرد و از اندیشمندانی چون میشل فوکو یا بل هاوکس که با امتناع از نوشتن حرف نخست خود با حرف بزرگ انگلیسی، زبان را به چالش کشیدند، مثال آورد. پیکاردی تاکید می‌کند: «ما به دنبال نوشتاری هستیم که هوش مصنوعی نتواند آن را بازتولید کند» یکی از بارزترین نمونه‌های سبک نوشتاری ضد چت‌بات‌ها تا به امروز، مقاله‌ای به قلم جسی مک‌گری است؛ او در مطلبی با عنوان «۵ دلیل برای دوست داشتن تابستان (بدون هوش مصنوعی)» نثری ارائه کرد، پر از غلط‌های نگارشی و پیشنهاد‌های ملتمسانه پی‌درپی تا ثابت کند که این مقاله را یک انسان نوشته است. مثلاً نوشته است: «وقتی تابستان فرا برسد می‌توانی از حال و هوای گرفته زمستان خلاص شوی و به آفتاب زیبا پناه ببری!»

مک‌گری می‌گوید که نوشتن اشتباهات واژگانی بسیار بامزه و کمی جسورانه بود و البته بسیار سخت؛ زیرا ناچار بودم به‌طور دائم دکمه عدم اعمال سیستم تصحیح خودکار کامپیوتر را فشار دهم.

این روش تقابل با مدل زبانی هوش مصنوعی به تدریج شکل گرفته و این روزها به اوج خود رسیده است. اکنون دیگر می‌توانید در پست‌های آموزشی گزینه چگونه از شبیه به هوش مصنوعی شدن اجتناب کنید را ببیند. ری استیلر بری - بنیان‌گذار موسسه بازاریابی کتاب دی‌اس‌ال - می‌گوید: «همیشه بر محتوای روایتی اول شخص تاکید می‌کنم؛ زیرا می‌خواهم سختی‌ای که برای نوشتن مطلب تحمل شده و عرق ریخته‌شده را ببینم». او می‌افزاید که این موسسه هر ماه جلساتی را برگزار کرده و روش‌هایی که در حال گسترش است را به‌روز رسانی می‌کند تا مطمئن شود متن‌هایی که انسان‌ها تولید می‌کنند، هرگز شبیه نوشته‌های هوش مصنوعی نباشند.

در ماه آوریل، بن هورویتز - سرمایه‌گذار و کارآفرین - وقتی مزایای این سبک مقابله با هوش مصنوعی را مشاهده کرد، نرم‌افزار «سینسری» را ساخت و آن را «ضد دستور زبان» نامید. این ابزار به عمد در ایمیل‌ها غلط املایی باقی می‌گذارد و پیام‌ها را با عبارت «ارسال شده از آیفون من» خاتمه می‌دهد. وقتی ایمیل‌های استاندارد ارسال شده را با ایمیل‌های تولید شده توسط نرم‌افزار خود مقایسه کرد، متوجه شد، تنها ایمیل‌هایی باز شده‌اند که با حروف کوچک و غلط املایی

نوشته شده بودند. هورویتز توضیح می‌دهد: «مردم اکنون عطش زیادی برای چیزهای متفاوت و فرار از یکنواختی دارند». او می‌گوید که نتیجه حاصله باعث شده است که او «واقعاً بیشتر قدر نوشته‌های خوب را بدانند».

هر چند گرایش و استقبال از نوشتار خوب می‌تواند یکی از مزایای عصر نوشتن با هوش مصنوعی نیز باشد. ماه آوریل گذشته، جان وارنر - مدرس باسابقه نویسندگی - کتاب «فراتر از واژگان: چگونه در عصر هوش مصنوعی به نوشتن فکر کنیم» را منتشر و در آن استدلال کرد که هوش مصنوعی ممکن است ارزش واقعی نوشتار خوب و سبک و فرایند نگارش مطلب با ارزش را دوباره احیا کند. او می‌گوید که سال‌ها تلاش کرد تا انشای معمولی پنج‌پاراگرافی دانش‌آموزان و آنچه او آن را مزخرفات شبه‌آکادمیک می‌خواند از میان بردارد؛ همان نوع متنی که اکنون توسط هوش مصنوعی تولید می‌شود.

مدل‌های زبانی بزرگ نوعی بدگمانی جدید را نسبت به سبک نگارش و نداشتن آن ایجاد کرده است که ممکن است در کوتاه‌مدت بد نباشد. اگرچه ممکن است این موضوع باعث نگرانی نویسندگان درجه یک جهانی نشود اما برای نویسندگان درجه دو و پایین‌تر می‌تواند برای پیشرفت، انگیزه ایجاد کند تا از استفاده زیاد از هوش مصنوعی دوری کنند.

هر چند، از آنجایی که هوش مصنوعی جا پای انسان می‌گذارد و همگام با رفتارهای نوشتاری انسانی آموزش می‌بیند، ممکن است به‌زودی راه‌اندازی جنبش ضد ضد ضد ضد سبک هوش مصنوعی موردنیاز باشد. مک‌گری می‌گوید: «هر سبک نوشتاری که در پاسخ و برخلاف سبک نوشتاری هوش مصنوعی ایجاد شود، هوش مصنوعی می‌تواند آن را تقلید کند؛ پس فقط نوشتار واقعی، نوشتار خوب است»



پی آیند به زبانی دیگر

مضمون اولین زنجیره : آداب هوش مصنوعی

قسمت پنجم - بخش دوم - فصل هفتم

کتابی از : پروفسور لوچیانو فلورییدی
با ترجمه : دکتر محسن رئیس قاسم

پروفسور لوچیانو فلورییدی معروف به فیلسوف اخلاق اطلاعات، استاد آکسفورد است و مولف کتب مهم بسیاری از انقلاب چهارم گرفته تا کتابی که می خوانید و عموماً در زمینه آداب فناوری اطلاعات



دکتر رئیس قاسم دانش آموخته علوم رایانه . کارشناسی دانشگاه شهید بهشتی (۱۳۶۷)، ارشد صنعتی شریف (۱۳۷۰) و دکتری از دانشگاه کارلتون کانادا (۱۳۷۷) و دارای حدود سی سال سابقه کار حرفه‌ای در ایران و کانادا <https://www.linkedin.com/in/mohsen-rais-ghasem-phd-3284753>

سید ابراهیم ابطحی

استادیار دانشکده مهندسی کامپیوتر دانشگاه صنعتی شریف
پست الکترونیکی: abtahi@Sharif.edu

پیشگفتار

به شکل پاورقی در شماره‌های متوالی نقل می‌کنیم. گاه ما متقاضی این ترجمه‌ها هستیم مثل مورد اول و شروع این پی‌آیند، که به ترجمه کتابی از لوچیانو فلورییدی^۱ فیلسوف اطلاعات اختصاص دارد که ترجمه آن را به درخواست من، دکتر محسن رئیس قاسم در کانادا

پی‌آیند به زبانی دیگر، پی‌آیند جدید ماست که با ترجمه کتاب‌هایی که مضمونی تازه از قلم فرهیختگانی نامدار را در بردارند که در انتظار نشر ترجمه تمام و کمال آنها نمی‌توان نشست و ضرورت ایجاب می‌کند با سرعت بیشتری به آنها دسترسی یافت را برایتان ترجمه و

1- Luciano Floridi



The Ethics of Artificial Intelligence
Principles, Challenges, and Opportunities
Luciano Floridi

اخلاق هوش مصنوعی

اصول، چالش‌ها و فرصت‌ها

لوچیانو فلورییدی

قسمت دوم

ارزیابی هوش مصنوعی

قسمت دوم به بررسی سرفصل‌های ویژه اخلاق هوش مصنوعی اختصاص دارد. رویکرد مورد استفاده در بررسی این سرفصل‌ها در قسمت اول تبیین شده بود؛ و می‌توان آن را ترکیبی از یک فرضیه و دو عامل دید. اولاً، هوش مصنوعی شکل جدیدی از هوشمندی نیست، بلکه شکل جدیدی از کارگزاری می‌باشد. هوش مصنوعی مصنوعی موفقیت‌اش را مدیون تفکیک کارگزاری از هوشمندی و نیز دربرگیری این کارگزاری منفک با محیط‌هایی است که دارای سازگاری فزاینده‌ای با هوش مصنوعی می‌باشند.

قسمت دوم از ده فصل تشکیل شده است که در مجموع کند و کاوی است در برخی از مبرم‌ترین موضوعات مطرح در زمینه اخلاق هوش مصنوعی. فصل ۴، با ارائه چارچوب یکپارچه‌ای از مبانی اخلاقی هوش مصنوعی، بر اساس تجزیه و تحلیل مقایسه‌ای از تأثیرگذارترین موارد پیشنهادی از سال ۲۰۱۷، سالی که برای اولین بار مبانی اخلاقی برای هوش مصنوعی در نشریات پژوهشی نشر یافتند، قسمت دوم را شروع می‌کند. فصل ۵، به بررسی تهدیداتی که ترجمه این مبانی به دستورالعمل‌های اجرایی را به چالش می‌کشد می‌پردازد. فصل ۶، بعد از چارچوب اخلاقی و تهدیدات اخلاقی، به تمشیت (governance) اخلاقی هوش مصنوعی می‌پردازد. من در این بخش، به معرفی مفهوم اخلاق نرم (soft ethics) به معنای اخلاق فراتر از الزامات قانونی (post-compliance) در اعمال مبانی اخلاقی می‌پردازم. در پی این سه فصل عمدتاً نظری، قسمت دوم توجه‌اش را به

انجام داده و می‌دهد و از شماره ۲۷۸ طی سیزده شماره آن را، و در این شماره در قسمت هفتم، فصل هفتم از بخش دوم آن را برای شما نقل می‌کنیم:

عنوان گسترده کتاب آداب هوش مصنوعی فلورییدی که طی بالغ بر دو سال ترجمه‌ای از همه فصول آن را خواهید خواند: اخلاق هوش مصنوعی، اصول، چالش‌ها و فرصت‌ها است. عنوان دو قسمت و سیزده فصل کتاب به شرح زیر است:

قسمت اول : درک هوش مصنوعی

فصل ۱- گذشته : پیدایش هوش مصنوعی، منتشره در شماره ۲۷۸
فصل ۲- حال : هوش مصنوعی به عنوان شکل جدیدی از عاملیت، نه هوش، منتشره در شماره ۲۷۹
فصل ۳- آینده : گسترش قابل پیش‌بینی هوش مصنوعی، منتشره در شماره ۲۸۰

قسمت دوم : ارزیابی هوش مصنوعی

فصل ۴- چارچوب یکپارچه‌ای برای مبانی اخلاقی هوش مصنوعی، منتشره در شماره ۲۸۱
فصل ۵- از اصول تا اعمال : تهدیدات اخلاقی نبودن، منتشره در شماره ۲۸۲
فصل ۶- اخلاق نرم و حکمرانی هوش مصنوعی، منتشره در شماره ۲۸۳

فصل ۷- نقشه آداب الگوریتم‌ها، منتشره در این شماره ۲۸۵
فصل ۸- راه کارهای ناپسند : استفاده نابجا از هوش مصنوعی در آسیب‌های اجتماعی
فصل ۹- راه کارهای پسندیده : استفاده مناسب از هوش مصنوعی در خدمت جامعه
فصل ۱۰- چگونه به جامعه خوب هوش مصنوعی برسیم: چند رهنمود
فصل ۱۱- حرکت حساب شده : تأثیر هوش مصنوعی روی تغییرات اقلیمی

فصل ۱۲- هوش مصنوعی و اهداف توسعه پایدار سازمان ملل متحد
فصل ۱۳- نتیجه گیری : سبز و آبی (اخلاق دیجیتال)
بجز مقدمه کتاب که به لحاظ ساختاری به سیاق توافق ویرایشی گزارش رایانه آن را ویراسته‌ام و موارد قلیلی از واژه‌گزینی، در دیگر موارد به احترام زحمات مترجم، همه مفروضات شکلی و ساختاری ایشان را با نقل وفادارانه در ترجمه رعایت کرده‌ام و در پایان می‌ماند سپاس ویژه من از مترجم فروتن این کتاب ارزشمند، که از اعضای قدیمی انجمن و دوست سالیان است.

پی آیند سیزده قسمتی آداب هوش مصنوعی لوچیانو فلورییدی به روایت دکتر محسن رئیس قاسم- قسمت ۷- بخش ۲- فصل ۷

جامعه پدیدار شده‌اند. این فصل، برآوردی از مباحث مطرح در این زمینه را، با هدف کمک به شناسایی و تحلیل پی‌آمدهای اخلاقی الگوریتم‌ها به دست می‌دهد. همچنین، تجزیه و تحلیلی به‌روز از نگرانی‌های معرفتی و هنجاری ارائه، و رهنمودهای عملی برای سامان‌دهی طراحی، توسعه و استقرار الگوریتم‌ها پیشنهاد شده است

۷.۰ درآمد: تعریفی کاری از الگوریتم

اجازه دهید با یک توضیح مفهومی شروع کنم. همانند «هوش مصنوعی»، در مورد تعریف الگوریتم توافق کمی توافق کمی در منابع مربوطه وجود دارد. از این اصطلاح بیشتر برای اشاره به تعریف صوری الگوریتم به عنوان ساختاری ریاضی با «ساختار کنترل ترکیبی، محدود، انتزاعی، موثر، با بیانی مودک، برای انجام هدفی مشخص با توجه به مفروضات داده شده» استفاده می‌شود (Hill, 2014, 4). اما همچنین از آن برای برداشتهای دامنه‌ای (domain-specific) که تمرکزشان بر پیاده‌سازی این ساختارهای ریاضی در فناوری‌هایی با کاربرد خاص می‌باشد استفاده می‌شود. در این فصل، به مسائل اخلاقی الگوریتم‌ها به عنوان سازه‌های ریاضی، به پیاده‌سازی آنها در قالب برنامه‌ها و پیکربندی‌ها (کاربردها)، و به روش‌هایی که از طریق آنها می‌توان به این مسائل پرداخت می‌پردازم^۳. در این برداشت، الگوریتم‌ها به عناصری کلیدی تبدیل شده‌اند که زیربنای خدمات و زیرساخت‌های حیاتی جوامع اطلاعاتی را شکل می‌دهند. برای نمونه، سیستم‌های توصیه‌گر را در نظر بگیرید، سیستم‌هایی الگوریتمی که پیشنهاداتی را، به‌صورت روزانه، درباره موارد مورد علاقه کاربر، از جمله انتخاب آهنگ، فیلم، محصول یا حتی دوست ارائه می‌دهد (Para-schakis, 2017; Perrault et al., 2019; Milano, Taddeo, and Floridi, 2020a; ridi, 2020b, 2021, 2019, 2020a; Obermeyer et al., 2019; Zhou et al., 2019; Morley, Machado, et al., 2020) موسسات مالی (Lee and Floridi, 2020; Lee, Floridi, 2020; Green and Aggarwal, 2020; ridi, and Denev, 2020)، و دادگاه‌ها (Chen, 2019; Yu and Du, 2019; Labati et al., 2016; banks 2017; Lewis, 2019; Hauer, 2019; Taddeo and Floridi, 2018b; Taddeo, McCutcheon, and Floridi, 2019)، همگی به‌طور فزاینده‌ای از الگوریتم‌ها در تصمیم‌گیری‌های مهم بهره می‌برند.

توان الگوریتم‌ها برای بهبود رفاه فردی و اجتماعی با تهدیدات اخلاقی قابل توجهی همراه است. همه می‌دانند که الگوریتم‌ها از نظر اخلاقی بی‌طرف نیستند. برای مثال، در نظر بگیرید که چگونه خروجی الگوریتم‌های ترجمه و موتورهای جستجو تا حد زیادی به‌عنوان عینی تصور می‌شوند، اما در واقعیت زبان مورد استفاده آنها مملو از کاربردهای جنسیتی است (Larson, 2017; Prates, Avelar, 2019; Lamb). همچنین سوگیری به‌طور گسترده‌ای گزارش شده است، مانند تبلیغات الگوریتمی برای فرصت‌هایی برای مشاغل پردرآمد و

پرسش‌های بیشتر کاربردی معطوف می‌کند. فصل ۷، نمایه به روز شده‌ای از مسائل اخلاقی اساسی‌ای که الگوریتم‌ها سبب می‌شوند به دست می‌دهد. در این فصل تلاش شده است بین فراگیری-در نظر گرفتن همه نوع الگوریتمی، و نه فقط آنهایی که در سیستم‌ها و برنامه‌های کاربردی هوش مصنوعی به کار می‌روند- و مشخص بودن تعادلی ایجاد کند. برای نمونه، در این فصل تمرکز فقط روی رباتیک و در نتیجه موارد اخلاقی ناشی از تن‌آیش (embodiment) هوش مصنوعی نیست. همان‌طور که در بخش ۸ در باب سوءاستفاده از هوش مصنوعی برای «شرارت اجتماعی» (social evil) نشان داده شده است، برخی از مشکلاتی که هم نیاز به برخورد قانونی (تمکین- com-pliance) و هم برخورد اخلاقی (اخلاق نرم) دارند، مشکلاتی ملموس و مبرم می‌باشند. با این همه، همان‌طور که در فصل ۹ در مورد هوش مصنوعی در خدمت جامعه (AI4SG - AI for social good) نشان داده شده است، هوش مصنوعی همچنین فرصت‌های بزرگی فراهم می‌آورد. هنگامی کار روی این دو فصل از آنها با «هوش مصنوعی بد» و «هوش مصنوعی خوب» یاد خواهیم کرد. هوش مصنوعی در هر دو فصل شامل رباتیک می‌باشد. فصل‌های ۱۰ تا ۱۲ با تمرکز بر چگونگی بروز هوش مصنوعی در خدمت جامعه، تأثیرات مثبت و منفی هوش مصنوعی بر تغییرات آب‌وهوایی، و تناسب هوش مصنوعی برای حمایت از اهداف توسعه پایدار سازمان ملل (SDGs) به توسعه هوش مصنوعی خوب بازمی‌گردد. سر انجام، فصل ۱۳ قسمت دوم و کتاب را با استدلال برای پیوند جدیدی بین سبز همه زیستگاه‌های ما و آبی همه فناوری‌های رقمی ما (به ویژه هوش مصنوعی) به پایان می‌برد. این فصل خاطر نشان می‌کند که کارگزاری، از جمله شکل جدید ارائه شده توسط هوش مصنوعی، نیازمند تمشیت است؛ و تمشیت متضمن فعالیت‌های اجتماعی-سیاسی است؛ و نظام رقمی کارگزاری اجتماعی-سیاسی را نیز دستخوش دگرگونی کرده، که موضوع کتاب بعدی، سیاست اطلاعات است.

۷

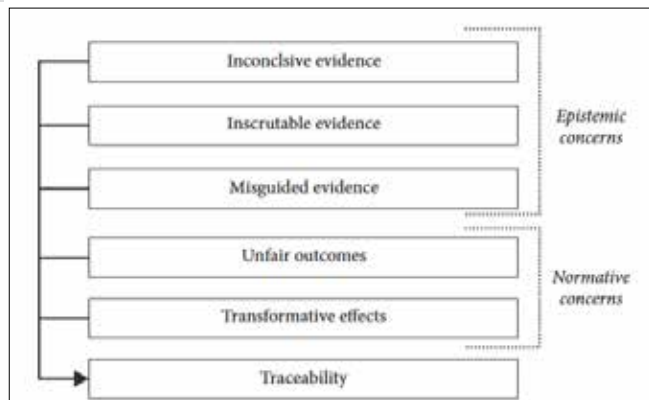
نگاشت جنبه‌های اخلاقی هوش مصنوعی

۷.۰ چکیده

پیش از این در فصل‌های ۴ تا ۶، به بررسی مبانی اخلاقی، تهدیدات، و تمشیت هوش مصنوعی پرداختیم. این فصل، با تمرکز بر بحث‌های جاری در مورد اخلاق الگوریتم‌ها، به بررسی چالش‌های اخلاقی ملموس هوش مصنوعی می‌پردازد. در فصل ۸ بیشتر در مورد رباتیک خواهیم گفت. پژوهش در زمینه اخلاق الگوریتم‌ها در سال‌های اخیر رشد قابل توجهی داشته است^۲. با توجه به رشد و کاربرد تصاعدی الگوریتم‌های یادگیری ماشینی (ML) در دهه گذشته، دشواری‌های اخلاقی و راه‌حل‌های جدیدی در ارتباط با کاربرد فراگیر آنها در

۲- ببینید (Floridi and Sanders, 2004) و (Floridi, 2013) برای ارجاعات قدیمی‌تر.

۳- این رویکرد به کار گرفته شده در Mittelstadt et al., 2016 و Samados et al., 2021 می‌باشد



شکل ۷.۱: نگرانی‌های شش گانه اخلاقی الگوریتم‌ها.
منبع: Mittelstadt et al. 2016

سبب می‌شوند، و بررسی راه‌کارهایی که در پژوهش‌های اخیر در این زمینه در میان گذاشته شده است، باقی می‌ماند. به همین دلیل، به عنوان نقطه آغازین بخش بعدی از آن استفاده شده است.

دو کار پژوهشی مذکور بنیان این فصل را شکل می‌دهند. در بخش‌های ۷.۳ تا ۷.۸، من فراتحلیلی (meta-analysis) از مباحث جاری در مورد اخلاق الگوریتم‌ها به‌دست می‌دهم و ارتباطی بین آنها و انواع دغدغه‌های اخلاقی قبلاً دیده شده برقرار می‌کنم. بخش ۷.۹ فصل را با مروری اجمالی و مقدمه‌ای برای فصل بعدی به پایان می‌رساند.

۷.۲ نگاشت اخلاق الگوریتم‌ها

از الگوریتم‌ها می‌توان در موارد زیر استفاده کرد:

- (۱) برای تبدیل داده‌ها به شواهد و دلایل (information Commissioner's Office) برای پیامدی مشخص،
- (۲) (۱) برای ایجاد انگیزه و آغاز عملی که می‌تواند دارای پیامدهای اخلاقی باشد مورد استفاده قرار می‌گیرد.
- الگوریتم‌های (تقریباً) خودمختار، مانند الگوریتم‌های یادگیری ماشین، می‌توانند (۱ و ۲) را انجام دهند. این امر، اقدام سوم، را پیچیده می‌کند:
- (۳) برای تشخیص مسئولیت تأثیرات اقداماتی که الگوریتم‌ها آغازگر آنها هستند.

موارد (۱ تا ۳)، به‌ویژه برای یادگیری ماشینی که شامل معماری‌های یادگیری عمیق است، جالب توجه است. سیستم‌های کامپیوتری که از الگوریتم‌های یادگیری ماشینی استفاده می‌کنند را می‌توان «خودمختار» یا «نیمه خودمختار» در نظر گرفت، چرا که خروجی‌شان ناشی از داده‌های [آموزشی] است و پیامدهای‌شان نامعین.

بر این مینا، نقشه مفهومی نشان داده شده در شکل ۷.۱ شش نگرانی اخلاقی را شناسایی می‌کند که در مجموع فضای مفهومی اخلاق الگوریتم‌ها را، به‌عنوان زمینه‌ای پژوهشی تعریف می‌کند. سه مورد از این نگرانی‌های اخلاقی به عوامل معرفتی اشاره دارند: شواهد

مشاغل در حوزه علم و فناوری که بیشتر برای مردان تبلیغ می‌شود تا زنان (Datta, Tschantz, Datta, 2015; Datta, Sen, Zick, 2016; Lambrecht, Tucker, 20). به همین منوال، الگوریتم‌های پیش‌بینانه مورد استفاده در مدیریت داده‌های سلامت میلیون‌ها بیمار در ایالات متحده، مشکلات موجود را تشدید می‌کنند، به‌گونه‌ای که بیماران سفیدپوست به‌طور محسوسی از بیماران سیاه‌پوست مشابه مراقبت‌های بهتری دریافت می‌کنند (Obermeyer et al., 2019). اگر چه راه‌کارهایی برای این مشکلات و مشکلات مشابه در دست بحث و طراحی است، تعداد سیستم‌های الگوریتمی که مشکلات اخلاقی از خود نشان می‌دهند همچنان در حال افزایش است.

در قسمت اول کتاب دیدیم که هوش مصنوعی حداقل از سال ۲۰۱۲ در حال تجربه «تابستان» جدیدی بوده است. این هم از نظر پیشرفت‌های فنی به‌دست آمده و هم از نظر توجهی است که این رشته از سوی دانشگاهیان، سیاست‌گذاران، دست‌اندرکاران فناوری، سرمایه‌گذاران (Perrault et al., 2019) و در نتیجه عموم مردم دریافت کرده است. در پی این «فصل» جدید از موفقیت‌ها، پژوهش‌های فزاینده‌ای در مورد پیامدهای اخلاقی الگوریتم‌ها، به‌ویژه در رابطه با انصاف، پاسخ‌گویی، و شفافیت، صورت گرفته است (Lee, 2018; Hoffmann et al., 2018; Shin and Park, 2019). در سال ۲۰۱۶، گروه پژوهشی ما در آزمایشگاه اخلاق رقی، گزارش جامعی با هدف ترسیم دورنمایی از این نگرانی‌های اخلاقی منتشر کرد (Mittelstadt et al., 2016). با این همه، این زمینه‌ای به سرعت در حال تغییر است. هم مشکلات اخلاقی جدید و هم راه‌هایی برای پرداختن به آنها پدیدار شده‌اند که بهبود و به‌روزرسانی آن گزارش را ضروری می‌سازد. از سال ۲۰۱۶، زمانی که دولت‌های ملی، سازمان‌های غیردولتی و شرکت‌های خصوصی نقش برجسته‌ای در گفتگوهای در مورد هوش مصنوعی و الگوریتم‌های «عدالانه» و «اخلاقی» پیدا کرده‌اند، کار روی اخلاق الگوریتم‌ها به‌طور قابل‌توجهی افزایش یافته است (Sandvig et al., 2016; Binns, 2018b, 2018a; Selbst et al., 2019; Ochigame, 2019; Wong, 2019; al., 2019). هم کمیت و هم کیفیت پژوهش‌های موجود در این زمینه افزایش زیادی یافته است. همه‌گیری COVID-19 هم به تشدید این موارد انجامید و هم آگاهی جهانی را در مورد آنها گسترش داد (Morley, Cowls, et al., 2020). با توجه به این تغییرات، ما دست به انتشار گزارش جدیدی زدیم (Morley et al., 2021)؛ همچنین نگاه کنید به (Morley et al., 2021). این گزارش، با افزودن نکته‌های جدید در مورد اخلاق الگوریتم‌ها، به‌روزرسانی تجزیه و تحلیل اولیه، افزودن ارجاعات به پژوهش‌هایی که در گزارش اولیه نادیده گرفته شد بودند، و گسترش موضوعات تحلیل شده (از جمله AI4SG) از کار (Mittelstadt et al., 2016) فراتر می‌رود. با این همه، نقشه مفهومی ۲۰۱۶ گزارش (شکل ۷.۱ را ببینید) چارچوبی پر بار برای بررسی بحث‌های جاری در مورد اخلاق الگوریتم‌ها، شناسایی دشواری‌های اخلاقی که الگوریتم‌ها

ارتباطات سببی (causal) را تشخیص نمی‌دهند. به این خاطر، آنها می‌توانند به ارتباط پنداری ۵ منجر شوند: «یافتن الگوهایی که وجود خارجی ندارند، صرفاً به این دلیل که یافتن ارتباطاتی در هر جهت با در دست داشتن مقادیر عظیمی از داده‌ها دشوار نیست» (Boyd and Crawford, 2012, 668).

این مسئله در دسرافرین است چرا که الگوهای شناسایی شده توسط الگوریتم‌ها می‌تواند نتیجه ویژگی‌های ذاتی سیستم مدل‌سازی شده توسط داده‌ها، مجموعه داده‌ها (یعنی خود مدل و نه سیستم زیربنایی آن) یا دستکاری ماهرانه مجموعه داده‌ها (ویژگی‌هایی که نه ربطی به مدل دارد و نه به سیستم) باشد. این مورد پارادوکس سیمپسون ۶ است که زمانی روی می‌دهد که روندهای مشاهده‌شده در گروه‌های مختلف داده‌ها، هنگام ترکیب داده‌ها معکوس می‌شوند (Blyth, 1972). در هر دو مورد آخر، کیفیت پایین داده‌ها به شواهد ناکافی برای پشتیبانی از تصمیمات انسانی منجر می‌شود.

پژوهش‌های اخیر این نگرانی را تقویت کرده‌اند که شواهد ناکافی می‌تواند به خطرات اخلاقی جدی منتهی شود. به عنوان مثال، تمرکز بر شاخص‌های غیرسببی می‌تواند توجه‌مان را از دلایل زیربنایی یک مشکل خاص منحرف کند (Floridi et al., 2020). حتی با استفاده از روش‌های سببی، داده‌های در دسترس همیشه دارای اطلاعات کافی برای توجیه اقدامی یا منصفانه کردن تصمیمی نیست (Olhede, 2018a; Wolfe, 2018b). کیفیت داده‌ها، از جمله، مناسبت زمانی، کامل بودن و درستی مجموعه داده‌ای، پرسش‌هایی را که می‌توان با استفاده از آن مجموعه داده‌ای پاسخ داد، محدود می‌کند (Olteanu et al., 2016). علاوه بر اینها، دریافت‌ها و اطلاعات قابل استخراج از مجموعه داده‌ها، اساساً به وابسته به فرضیاتی هستند که خود فرآیند جمع‌آوری داده‌ها را هدایت کرده‌اند (Diakopoulos and Koliska, 2017). برای مثال، الگوریتم‌های پیش‌بینی نتایج [درمان] بیماران در محیط‌های بالینی طراحی شده‌اند، کاملاً به داده‌های ورودی‌ای که امکان کمی‌سازی آنها وجود دارد (مثلاً علایم حیاتی و میزان موفقیت درمان‌های مقایسه‌ای پیشین) متکی هستند. این بدان معناست که عوامل احساسی و عاطفی (مثلاً تمایل به زندگی) که تأثیرات قابل توجهی بر درمان بیمار دارند، نادیده گرفته می‌شوند، که از دقت پیش‌بینی‌های الگوریتمی می‌کاهد (Buhmann, Paßmann, and Fieseler, 2019). این مثال به خوبی نشان می‌دهد که چگونه دستاوردهای پردازش داده‌ای الگوریتمی می‌توانند نامشخص، ناقص و مقطعی باشند (Diakopoulos and Koliska, 2017).

می‌توان رویکردی خام و استقرایی (inductivist) اختیار کرد و فرض را بر این گذاشت که می‌توان از شواهد ناکافی و فاقد قطعیت اجتناب کرد به این شرط که الگوریتم‌ها به داده‌های کافی دسترسی داشته باشند، حتی اگر نتوان توضیحی سببی برای این نتایج ارائه داد.

ناکافی (inconclusive)، درک ناشدنی (inscrutable) و گمراه‌کننده (misguided). دو مورد به صراحت هنجاری هستند: پیامدهای ناعادلانه (unfair) و اثرات تحول‌آفرین (transformative). یکی، توان ردیابی (traceability)، هم به اهداف معرفتی و هم به اهداف هنجاری مربوط می‌شود. عوامل معرفتی مورد اشاره، اهمیت کیفیت و دقت داده‌ها ۴ در توجیه نتایج الگوریتم‌ها را برجسته می‌سازند. این نتایج به نوبه خود می‌توانند تصمیمات اخلاقی‌ای را شکل دهند که افراد، جوامع و محیط‌زیست را تحت تأثیر قرار می‌دهند. دو نگرانی هنجاری در نقشه مفهومی بالا، به روشنی به تأثیر اخلاقی اقدامات و تصمیمات مبتنی بر الگوریتم، از جمله عدم شفافیت ناروشنی (opacity) در فرآیندهای الگوریتمی، نتایج ناعادلانه و پیامدهای ناخواسته اشاره دارند. نگرانی‌های معرفت‌شناختی و هنجاری - همراه با پراکندگی طراحی، توسعه و استقرار الگوریتم‌ها - ردیابی دنباله رویدادها و عواملی که منجر به پیامدی مشخص می‌شوند را دشوار می‌کند. این دشواری، امکان شناسایی دلایل بروز پیامدی، و در نتیجه، انتساب مسئولیت اخلاقی برای آن را مختل می‌کند (Floridi, 2012b). این همان چیزی است که ششمین دغدغه اخلاقی، یعنی توان ردیابی، به آن اشاره دارد.

تاکید بر این نکته اهمیت دارد که این نقشه مفهومی را می‌توان در سطح انتزاعی اخلاق خرد و کلان تفسیر کرد (Floridi, 2008a). این نقشه، در سطح اخلاق خرد، به باز کردن دشواری‌های اخلاقی الگوریتم‌ها می‌پردازد. با برجسته کردن جدایی‌ناپذیری این موارد از موارد مربوط به داده‌ها و مسئولیت‌ها، این نقشه مفهومی همچنین ضرورت رویکرد کلان اخلاقی برای پرداختن به اخلاق الگوریتم‌ها به عنوان بخشی از فضای مفهومی گسترده‌تری، اخلاق دیجیتال، را نشان می‌دهد. من و تادئو در گذشته استدلال کرده‌ایم که: اگرچه اخلاق داده‌ها، الگوریتم‌ها و راه‌کارها خطوط پژوهشی متمایزی را دنبال می‌کنند، این پژوهش‌ها آشکارا در هم تنیده شده‌اند ... اخلاق [رقمی] باید به کل فضای مفهومی و از این رو هر سه محور پژوهشی، حتی اگر با اولویت‌ها و مرکزیت‌های متفاوت، بپردازد (Floridi and Taddeo, 2016, 4).

در ادامه این فصل، به ترتیب به هر یک از این شش دغدغه اخلاقی خواهیم پرداخت و تحلیلی به‌روز از منابع پژوهشی اخلاق الگوریتم‌ها (در سطح خرد)، و با هدف ارتقاء بحث اخلاق رقمی (در سطح کلان)، ارائه خواهیم داد.

۷.۳ اقدامات ناموجه ناشی از شواهد ناکافی

پژوهش در مورد شواهد ناکافی اشاره به الگوریتم‌های یادگیری ماشینی نامعین (non-deterministic) دارد که خروجی‌هاشان دارای ماهیتی احتمالی است (James et al., 2013; Valiant, 1984). این نوع الگوریتم‌ها عموماً در شناسایی پیوند (association) و همبستگی (correlation) بین متغیرها را در داده‌های داده شده موفقند، اما

۵- در اصل apophenia ببینید <https://en.wikipedia.org/wiki/Apophenia>
۶- در اصل Simpson's paradox ببینید [https://en.wikipedia.org/wiki/Simpson's paradox](https://en.wikipedia.org/wiki/Simpson%27s_paradox)

۴- در مورد بحث و فلسفه کیفیت داده‌ها به Floridi and Illari, 2014 رجوع کنید.

که مشخصه اغلب الگوریتم‌ها (به‌ویژه الگوریتم‌ها و مدل‌های یادگیری ماشینی) می‌باشد، زیرساخت‌های اجتماعی - فنی آنها، و تصمیماتی که می‌گیرند می‌پردازد. فقدان شفافیت می‌تواند به دلیل محدودیت‌های خود فناوری، ناشی از تصمیمات طراحی، و مبهم‌سازی داده‌های زیربنایی باشد (Lepri et al., 2018; Dahl 2018, Ananny and Crawford, 2018; Weller, 2019)، یا از محدودیت‌های قانونی در رابطه به مالکیت معنوی سرچشمه بگیرد. در هر حال، فقدان شفافیت، در بیشتر موارد، به فقدان بررسی موشکافانه و پاسخگویی (Oswald 2018; Webb et al., 2019) می‌انجامد، و به فقدان «قابلیت اعتماد» ختم می‌گردد (Bennett 2019, AI HLEG).

طبق مطالعات اخیر (Diakopoulos and Koliska, 2017; Stilgoe, 2018; Zerilli et al., 2019; Buhmann, Paßmann, and Fieseler, 2019)، عوامل دخیل در فقدان کلی شفافیت الگوریتمی از این قرارند:

- از نظر شناختی، تفسیر مدل‌ها و مجموعه‌های داده‌های الگوریتمی عظیم برای انسان‌ها ممکن نیست

- نبود ابزارهای مناسب برای ترسیم و ردیابی برنامه‌های رایانه‌ای بزرگ و داده‌های بسیار ساختار برنامه‌های رایانه‌ای و داده‌ها آنچنان ضعیف است که سر در آوردن از آنها شدنی نیست و

- به‌روزرسانی‌های مداوم و تأثیرات انسانی روی مدل‌ها

عدم شفافیت همچنین یکی از ویژگی‌های ذاتی الگوریتم‌های خودآموز است که منطق تصمیم‌گیری‌شان (برای ساخت قواعد جدید) در طول فرآیند یادگیری تغییر می‌کند. بنابراین، سازندگان این الگوریتم‌ها برای درک دقیق چرایی برخی تغییرات مشخص با دشواری مواجه‌اند (Burrell, 2016; Buhmann, Paßmann, and Fieseler, 2019). با این حال، این‌ها لزوماً به معنی پیامدهای ناروشتن نیست؛ سازندگان این الگوریتم‌ها حتی بدون درک هر مرحله، می‌توانند فرایارمترها (hyperparameters).

- پارامترهایی که فرآیند آموزش را کنترل می‌کنند) را برای آزمایش خروجی‌های مختلف تنظیم کنند. (Martin, 2019)، تأکید دارد که اگرچه توضیح خروجی‌های الگوریتم‌های یادگیری ماشینی قطعاً دشوار است، اما مهم است که اجازه ندهیم این دشواری‌ها انگیزه‌ای بشود برای سازمان‌ها که برای شانه خالی کردن از مسئولیت به توسعه سیستم‌های پیچیده روی آورند.

عدم شفافیت همچنین می‌تواند ناشی از انعطاف‌پذیری الگوریتم‌ها باشد. الگوریتم‌ها را می‌توان به صورت پیوسته، توزیع‌شده و پویا دوباره برنامه‌ریزی کرد (Sandvig et al., 2016)، که آنها را انعطاف‌پذیر می‌کند. انعطاف‌پذیری الگوریتمی به سازندگان اجازه می‌دهد تا الگوریتم‌های از قبل مستقر شده را رصد و بهبود بخشند. اما می‌توان از این ویژگی برای پاک کردن سابقه تغییرات الگوریتم‌ها سوءاستفاده کرد، که سبب ایجاد شک و تردید کاربران نهایی در مورد توانایی‌های آنها می‌گردد (Ananny and Crawford, 2018). برای مثال، الگوریتم جستجوی اصلی گوگل را در نظر بگیرید. انعطاف‌پذیری آن، شرکت را

اما، پژوهش‌های اخیر این دیدگاه را رد می‌کند. به‌طور مشخص، پژوهش‌هایی که به خطرات اخلاقی نمانگاری نژادی (racial profiling) با استفاده از سیستم‌های الگوریتمی می‌پردازند، با برجسته ساختن این واقعیت که نابرابری‌های ساختاری دیرینه اغلب در داده‌های مورد استفاده الگوریتم‌ها نهادینه شده‌اند، به روشنی محدودیت‌های چنین رویکردی را نشان می‌دهند. این نابرابری‌ها به ندرت تصحیح می‌شوند و در بیشتر مواقع دست نخورده باقی می‌مانند (Hu, 2017; Turner Lee, 2018; Noble, 2018; Benjamin, 2019; Richardson, 2020; Schultz, and Crawford, 2019; Abebe et al., 2020). افزایش داده‌ها به خودی خود به دقت بیشتر یا بازنمایی بهتر منجر نمی‌شود. برعکس، می‌تواند با افزایش امکان پیدا کردن همبستگی‌های صوری، دشواری‌های مربوط به شواهد غیرقطعی و ناکافی را تشدید کند. به گفته روها بنجامین (Ruha Benjamin)، «عمق محاسباتی بدون عمق تاریخی یا جامعه‌شناختی صرفاً یادگیری سطحی است آن‌ه یادگیری عمیق»^۷.

کوتاهی‌های یاد شده، توجیه‌پذیری خروجی‌های الگوریتمی را با محدودیت‌های جدی مواجه می‌کند. استنتاج‌های درجه دو می‌تواند روی افراد یا کل جمعیت تأثیرات منفی بگذارد. در مورد علوم طبیعی، این امر حتی می‌تواند کفه را به نفع یا علیه «نظریه علمی خاصی» تغییر دهد (Ras, van Gerven, and Haselager, 2018). به همین دلیل است که اعتبارسنجی مستقل داده‌های ورودی به الگوریتم‌ها، و نیز کسب اطمینان از وجود کاربست‌های لازم برای حفظ و تکرارپذیری داده‌ها به منظور کاستن از پیامدهای اقدامات ناموجه ناشی شواهد ناکافی و غیرقطعی، در کنار فرایندهای نظارتی برای شناسایی پیامدهای ناعادلانه و نتایج ناخواسته، از اهمیت بسیاری برخوردار است (Henderson et al., 2018; Rahwan, 2018; Davis, 2020; and Marcus, 2019; Brundage et al., 2020). خطر شواهد ناکافی و فاقد قطعیت و آموزه‌های قابل پی‌گیری نادرست، همچنین ناشی از تصور عینیت مکانیکی درباره تجزیه و تحلیل‌های رایانه‌ای می‌باشد (Karppi, 2018; Lee, 2018; Buhmann, Paßmann, and Fieseler, 2019). این می‌تواند سبب شود تصمیم‌گیرندگان انسانی ارزیابی‌های مبتنی بر تجربیات ارزشمند خود را نادیده بگیرند - موسوم به «سوگیری اتوماسیون»^۸ (Cummings, 2012) - یا حتی از مسئولیت خود در قبال تصمیمات اخذ شده شانه خالی کنند (به بخش ۷.۸ در زیر مراجعه کنید) (Grote and Berens, 2020). همان‌طور که در بخش‌های ۷.۴ و ۷.۸ خواهیم دید، نداشتن درک روشنی از چگونگی تولید خروجی‌ها توسط الگوریتم‌ها، این مشکل را تشدید می‌کند.

۷.۴ ناروشتن ناشی از شواهد درک ناشدنی

شواهد درک ناشدنی به دشواری‌های مربوط به عدم شفافیت

7-<https://aas.princeton.edu/news/ruha-benjamin-deep-learning-computational-depth-without-sociological-depth-superficial>

۸- https://en.wikipedia.org/wiki/Automation_bias ببینید

امکان اعمال روش‌های اخلاقی مبنایی را فراهم آورد- کمکی به حل مشکلات اخلاقی موجود در ارتباط با کاربرد الگوریتم‌ها نمی‌کند. در واقع، این سردرگمی می‌تواند مشکلات جدیدی بیافریند. به همین دلیل است که تشخیص و تمایز عوامل مختلفی که مانعی در راه شفافیت الگوریتم‌ها هستند، شناسایی علل آنها، و روشن کردن موارد مورد نیاز، و چگونگی پرداختن به آنها در لایه‌های مختلف سیستم‌های الگوریتمی در فراخوان شفافیت، اهمیت دارند (Diako- poulos and Koliska, 2017).

روش‌های مختلفی برای پرداختن به دشواری‌های ناشی از عدم شفافیت وجود دارد. برای مثال، (Gebru et al., 2020) پیشنهاد می‌کند که از رویه‌های مستندسازی استاندارد، شبیه آنچه در صنعت الکترونیک به کار می‌رود، برای برطرف کردن نسبی دشواری‌های شفافیتی ناشی از انعطاف‌پذیری الگوریتم‌ها استفاده گردد. در این رویه‌ها، «هر قطعه، صرف‌نظر از سادگی یا پیچیدگی آن، دارای برگه اطلاعاتی همراهی است که ویژگی‌های کارکردی، نتایج آزمایش، کاربرد توصیه شده، و سایر اطلاعات آن را شرح می‌دهد» (Gebru et al., 2020, 2). متأسفانه، در حال حاضر، مستندسازی‌ای که در دسترس عموم باشد برای سیستم‌های الگوریتمی رایج نیست و هیچ ساختار جافاده‌ای هم در مورد مستندسازی مبدا مجموعه داده‌ها نیز وجود ندارد (Arnold et al., 2019; Gebru et al., 2020).

اگرچه نورس و هنوز در حال رشد، رویکرد بالقوه امیدوارکننده دیگری برای کسب شفافیت الگوریتمی، استفاده از ابزارهای فنی برای آزمایش و ممیزی سیستم‌های الگوریتمی و تصمیم‌گیری است (Mokander and Floridi, 2021). هم آزمایش این که آیا الگوریتمی گرایش‌های منفی (مانند تبعیض ناعادلانه) از خود نشان می‌دهد، و هم ممیزی دقیق پیش‌بینی‌ها یا زنجیره تصمیم‌ها می‌تواند سطح بالایی از شفافیت را ممکن سازد (Weller, 2019; Malhotra, Kot- Brundage et al., 2020; Wal, and Dalal, 2018). در این راستا، چارچوب‌های برهانی (discursive) برای کمک به کسب‌وکارها و سازمان‌های بخش دولتی برای درک تأثیرات بالقوه الگوریتم‌های ناروشن، تشویق راه‌کارهای خوب تدوین شده‌اند (ICO, 2020). به‌عنوان مثال، موسسه AI Now دانشگاه نیویورک، راهنمایی برای ارزیابی تأثیرات الگوریتمی تهیه کرده است که هدفش افزایش آگاهی و بهبود گفتگو در مورد زبان‌های بالقوه الگوریتم‌های یادگیری ماشینی است (Reisman et al., 2018). هدف از افزایش آگاهی و گفتگو، فراهم آوردن شرایط لازم برای طراحی الگوریتم‌های یادگیری ماشینی شفاف‌تر و در نتیجه قابل اعتمادتر است. همچنین، تلاش می‌شود درک عموم مردم از الگوریتم‌ها و کنترل‌شان روی آنها بهبود یابد. در همین راستا، (Diakopoulos and Koliska, 2017) فهرستی فراگیری از «عوامل شفافیت» را در سطح چهار لایه سیستم‌های الگوریتمی ارائه داده‌اند: داده، مدل، استنتاج، و واسطه (interface). این عوامل، از جمله، «عدم قطعیت (مثلاً حاشیه خطا)، مناسب

قادر می‌سازد تا همواره تغییراتی در آن اعمال کند، که نشان‌گر نوعی بی‌ثباتی دایمی است (Sandvig et al., 2016). بدین‌خاطر، افرادی که تحت تأثیر این الگوریتم‌ها هستند، می‌باید همواره این تغییرات را رصد کرده و فهم خود از آنها را به‌روز کنند، کاری که برای اکثر افراد غیرممکن است (Ananny and Crawford, 2018).

همان‌طور که من و توریلی در گذشته اشاره کرده‌ایم، شفافیت به خودی خود یک اصل اخلاقی نیست، بلکه یک شرط اخلاقی برای ممکن ساختن یا تضعیف سایر راه‌کارها یا اصول اخلاقی است (Turilli and Floridi, 2009, 105). به عبارت دیگر، شفافیت یک ارزش ذاتی اخلاقی نیست، بلکه وسیله‌ای ارزشمند برای رسیدن به اهداف اخلاقی است. شفافیت بخشی از شرایط وجوبی رفتار اخلاقی است که من در جای دیگر آن را به‌عنوان «اخلاق ساخت» (infraethics) تعریف کرده‌ام (infrastructural ethics) - اخلاق زیرساختی، (Floridi, 2017a). گاهی اوقات، ناروشنی می‌تواند مفیدتر باشد. به‌عنوان مثال، می‌تواند در خدمت پوشیده نگاه داشتن آرا و تمایلات سیاسی شهروندان، یا رقابت در مزایده‌های خدمات عمومی باشد. حتی در زمینه‌های الگوریتمی، شفافیت کامل می‌تواند به خودی خود باعث مشکلات اخلاقی خاصی گردد (Ananny and Crawford, 2018). شفافیت کامل می‌تواند اطلاعات مهمی در مورد ویژگی‌ها و محدودیت‌های الگوریتمی در اختیار کاربران قرار دهد، اما همچنین می‌تواند کاربران را در اطلاعات زیادی غرق کند، و سبب ناروشن شدن الگوریتم شود (Kizilcec, 2016; Ananny and Crawford, 2018).

پژوهشگران دیگری تأکید دارند که تمرکز بیش از حد بر شفافیت می‌تواند برای نوآوری مضر باشد و به انحراف بی‌جهت منابعی که می‌توانستند برای بالابردن ایمنی، کارکرد، و دقت به کار روند بیانجامد (Danks and London 2017, Oswald, 2018; Ananny and Weller, 2019). اولویت دادن به شفافیت (و تبیین‌پذیری) به ویژه در زمینه الگوریتم‌های پزشکی بحث‌های زیادی برانگیخته است (Robbins, 2019). شفافیت همچنین می‌تواند به افراد امکان دستکاری و بهره‌برداری از سیستم را بدهد (Martin, 2019; Magalhães, 2018; Cows et al., 2019). دانستن منبع یک مجموعه داده، فرضیات مورد استفاده در نمونه‌برداری، یا چگونگی مرتب‌سازی ورودی‌های جدید توسط یک الگوریتم، می‌تواند در یافتن راه‌هایی برای بهره‌برداری از آن الگوریتم مورد استفاده قرار گیرد (Szegedy et al., 2014; Yampolskiy, 2018). با این حال، توانایی دستکاری الگوریتم‌ها فقط در دسترس بخش‌های محدودی از جامعه می‌باشد (به‌ویژه کسانی که مهارت‌های رقی پیشرفته، و منابع لازم برای استفاده از آنها را دارند)، که نوع دیگری از نابرابری‌های اجتماعی را تشکیل می‌دهد (Martin, 2019; Bambauer and Zarsky, 2018). اشتباه گرفتن شفافیت به عنوان هدفی به خودی خود - به‌جای عاملی اخلاقی (Floridi, 2017) که باید به درستی طراحی شود تا

به‌طور کامل درک کنند (Hutson, 2019). این امر پژوهشگران را بر آن داشته که بر ضرورت سرمایه‌گذاری بیشتر در آموزش عمومی برای افزایش سواد محاسباتی و داده‌ای به منظور مقابله با پیچیدگی‌های فنی تاکید ورزند (Lepri et al., 2018). به نظر می‌رسد چنین اقداماتی، در بلندمدت، نقشی پررنگ در حل مسایل چندلایه‌ای ناشی از الگوریتم‌های همیشه‌حاضر ایفا کند. به‌علاوه، از نرم‌افزارهای متن‌باز اغلب به‌عنوان عواملی اساسی در حل این مشکلات یاد می‌شود (Lepri et al., 2018). بنابراین، می‌تواند منجر به مداخلات الگوریتمی شود که تلاش می‌کنند «بی‌طرف» باشند، اما در انجام این کار، خطر تثبیت شرایط اجتماعی موجود را به همراه دارند (Green and Viljoen, 2020, 20) و در عین حال توهم دقت را ایجاد می‌کنند (Karppi, 2018; Selbst et al., 2019). به همین دلایل، استفاده از الگوریتم‌ها در برخی از محیط‌ها به‌طور کلی مورد سوال قرار می‌گیرد (Selbst et al., 2019; Mayson, 2019; Katell et al., 2020; Abebe et al., 2020).

۷.۵ پیش‌داوری‌های ناخواسته ناشی از شواهد گمراه‌کننده

تمرکز سازندگان سیستم‌های الگوریتمی، در درجه اول، روی این است که الگوریتم‌هاشان از پس اجرای وظایف تعریف شده برآیند. بنابراین، برای درک پدیده سوگیری و پیش‌داوری (bias) در الگوریتم‌ها و تصمیم‌گیری الگوریتمی، آشنایی با تفکری مدنظر سازندگان ضروری است. برخی از پژوهشگران اشاره دارند که مشخصه تفکر غالب در حوزه توسعه الگوریتم‌ها، «فرمالیسم الگوریتمی»، یا پایبندی به قوانین و فرم‌های از پیش تعیین‌شده، می‌باشد (Green and Viljoen, 2020, 21). اگرچه این رویکردی مفید در کارکردهای انتزاعی و تعریف فرآیندهای تحلیلی است، اما پیچیدگی‌های اجتماعی دنیای واقعی را نادیده می‌گیرد (Katell et al., 2020). در نتیجه، چنین رویکردی می‌تواند به مداخلات الگوریتمی‌ای منجر گردد که اگرچه در تلاش است که «بی‌طرف» باشند، اما در عمل خطر تحکیم شرایط اجتماعی موجود را به همراه دارد (Green and Viljoen, 2020, 20)، در حالی که توهم دقت را القا می‌کند (Karppi 2018, Selbst et al., 2019). بنابراین، استفاده از الگوریتم‌ها در برخی موارد به‌طور کلی جای پرسش دارد (Selbst et al., 2019; Mayson, 2019; Katell et al., 2020; Abebe et al., 2020).

برای مثال، تعداد روزافزونی از پژوهشگران استفاده از ابزارهای ارزیابی تهدیدات مبتنی بر الگوریتم‌ها را در روندهای حقوقی به زیر سوال می‌برند (Berk et al., 2018; Abebe et al., 2020). برخی از پژوهشگران با تاکید بر محدودیت‌های روش‌های انتزاعی در مورد سوگیری‌های ناخواسته در الگوریتم‌ها، استدلالاتی در مورد ضرورت توسعه چارچوبی اجتماعی-فنی برای پرداختن و بهبود عدالت و انصاف در الگوریتم‌ها مطرح می‌کنند (Edwards and Veale, 2017; Selbst et al., 2019; Wong, 2019; Katell et al., 2020; Abebe et al., 2020).

زمانی (مثلاً زمان گردآوری داده‌ها)، کامل یا ناقص بودنداده‌ها، روش نمونه‌گیری، یافت‌گاه (مثلاً منابع) و حجم (مثلاً داده‌های آموزشی مورد استفاده در یادگیری ماشین) را در بر می‌گیرد (Diakopoulos and Koliska., 2017, 818).

هر رویه شفافیت‌سازی به احتمال زیاد (و در واقع از روی ضرورت) شامل توضیحی تفسیرپذیر (interpretable explanation) از فرآیندهای داخلی این سیستم‌ها می‌باشد. (Paßmann, and Fie-seler, 2019). اگرچه عدم شفافیت ویژگی ذاتی بسیاری از الگوریتم‌های یادگیری ماشینی است، اما این بدان معنا نیست که نمی‌توان آنها را بهبود بخشید. شرکت‌هایی مانند گوگل و آبی‌ام با عرضه عمومی ابزارهایی مانند Explainable AI, AI Explain-ability 360 و What-If Tool، تلاش‌های خود را برای تفسیرپذیرتر و فراگیرتر کردن الگوریتم‌های یادگیری ماشینی افزایش داده‌اند. قطعاً استراتژی‌ها و ابزارهای بیشتری ساخته خواهند شد. این ابزارها واسطه‌های بصری تعاملی‌ای هستند که به کارشناسان هوش مصنوعی و عموم مردم امکان می‌دهند که نتایج مختلف مدل‌ها را بررسی کنند، از استدلال‌ات مبتنی بر مورد ۹ و قواعد قابل تفسیر بهره گیرند، و حتی سوگیری‌های ناخواسته را در مجموعه داده‌ها و مدل‌های الگوریتمی شناسایی و تاثیراتشان را کاهش دهند (Mojsilovic, 2018; Wexler, 2018). با این حال، توضیح الگوریتم‌های یادگیری ماشینی، بسته به نوع توضیحات موردنظر، این واقعیت که تصمیمات در بیشتر موارد ماهیتی چند بعدی دارند، و این که کاربران مختلف ممکن است به توضیحات متفاوتی نیاز داشته باشند، دارای محدودیت‌هایی هستند (Edwards and Veale, 2017; Watson et al., 2019). روش‌های مناسب برای ارائه توضیحات از اواخر دهه ۱۹۹۰ به عنوان یک مشکل شناخته شده است (Tickle et al., 1998)، اما تلاش‌های جاری را می‌توان به دو رویکرد اصلی تقسیم کرد: توضیحات محور و توضیحات مدل محور (Doshi-Velez and Kim, 2017; Lee, 2017; Kim, and Lizarondo, 2017; Baumer, 2017; Buhmann, Paßmann, and Fieseler, 2019). در رویکرد اول، دقت و طول توضیحات متناسب با کاربران و تعاملات خاص آنها با یک الگوریتم مشخص تنظیم می‌شود (برای مثال، به Green and Viljoen, 2020 و مدل بازی‌مانندی که دیوید واتسون و من در Watson and Floridi, 2020 پیشنهاد دادیم مراجعه کنید). در رویکرد دوم، توضیحات مربوط به کل مدل است و به مخاطب آنها بستگی ندارد.

توضیح‌پذیری، به‌ویژه با توجه به سرعت رو به رشد تعداد مدل‌ها و مجموعه‌های داده‌ای متن‌باز (open-source) و به راحتی در دسترس، اهمیت بیشتری پیدا می‌کند. افراد غیرمتخصص به‌طور فزاینده‌ای در حال به‌کارگیری مدل‌های الگوریتمی پیشرفته‌ای هستند که به آسانی از طریق مجموعه‌ها یا سکوها برخطی، مانند گیت‌هاب، در دسترس‌شان هست، بدون این که محدودیت‌ها و ویژگی‌های آنها را

۹-در اصل case based reasoning ببینید https://en.wikipedia.org/wiki/Case-based_reasoning

اگر بتوان محافظت از گروه‌های اجتماعی خاصی را در یک الگوریتم کدگذاری کرد، همیشه امکان سوگیری‌هایی وجود دارد که از پیش در نظر گرفته نشده بودند. از این روست که مدل‌های زبانی متونی را که به شدت بر مردان تمرکز دارد، بازتولید می‌کنند (Fuster et al., 2017; Doshi-Velez and Kim, 2017). حتی اگر سوگیری پیش‌بینی و متغیرهای محافظت‌شده از داده‌ها حذف شده باشند، متغیرهای جایگزین (proxy) دور از انتظار می‌توانند در بازسازی آن سوگیری‌ها به کار روند؛ که به «سوگیری به واسطه جایگزین» منجر می‌شود که تشخیص و اجتناب از آن دشوار است (Fuster et al., 2017; Gillis and Spiess, 2019). نمونه نوعی چنین پیشداوری‌هایی، پیش‌داوری بر مبنای کدپستی است.

با این همه، ممکن است دلایل خوبی برای استفاده از تخمین‌گرهای آماری سوگیرانه در پردازش‌های الگوریتمی وجود داشته باشد، مثلاً می‌توان از آنها برای کاهش سوگیری داده‌های آموزشی استفاده کرد. در این روش تلاش می‌شود، بین یک نوع سوگیری الگوریتمی در درس‌ساز با نوع دیگری از سوگیری الگوریتمی، یا با افزودن سوگیری جبرانی هنگام تفسیر خروجی‌های الگوریتمی، تعادل ایجاد شود (Danks and London, 2017). آزمایش الگوریتم‌ها در زمینه‌های مختلف و با مجموعه داده‌های متنوع رویکردی ساده‌تر برای کاهش سوگیری در داده‌ها است (Shah, 2018).

فراهم آوردن امکان بررسی مدل‌ها، مجموعه داده‌ها و فراداده‌های (خواستگاه) آنها توسط افراد بیرونی، می‌تواند به اصلاح سوگیری‌های نادیده یا ناخواسته نیز کمک کند (Shah, 2018). همچنین شایان ذکر است که داده‌های ساختگی (synthetic)، یا داده‌های تولید شده توسط هوش مصنوعی به شیوه یادگیری تقویتی یا شبکه‌های تقابلی مولد (GAN - Generative Adversarial Networks)، فرصت ارزشمندی برای پرداختن به مشکلات سوگیری در داده‌ها ارائه می‌دهند و همان‌طور که در فصل ۳ اشاره کردم، می‌تواند نشانه مهمی از توسعه مهم هوش مصنوعی در آینده باشند. تولید داده‌های منصفانه توسط شبکه‌های تقابلی مولد می‌تواند به ایجاد تنوع در مجموعه داده‌های مورد استفاده در الگوریتم‌های بینایی رایانه‌ای کمک کند (Xu et al., 2018). به‌عنوان مثال، سیستم StyleGAN2 می‌تواند تصاویر با کیفیتی از چهره‌های انسانی موهوم تولید کند و نشان داده شده است که به‌ویژه در تولید مجموعه‌های داده‌ای متنوع از چهره‌های انسان مفید است - چیزی که بسیاری از سیستم‌های الگوریتمی تشخیص چهره در حال حاضر فاقد آن هستند (Obermeyer et al., 2019; Kortylewski et al., 2019; Harwell, 2020).

استقرار (deployment) نابجای الگوریتم‌ها نیز می‌تواند سبب پدید آمدن سوگیری ناخواسته گردد. سوگیری انتقال زمینه (transfer context bias) یا سوگیری در درس‌آفرینی که هنگام استفاده از الگوریتمی در محیطی جدید بروز می‌کند، را در نظر بگیرید. به عنوان مثال، اگر از الگوریتم مراقبت‌های بهداشتی یک بیمارستان پژوهشی

(et al., 2020)، در این راستا، (Selbst et al., 2019, 3-60)، به پنج «تله» انتزاعی یا در نظر نگرفتن زمینه‌های اجتماعی الگوریتم‌ها اشاره دارد. به دلیل نبود چارچوبی اجتماعی - فنی این، تله‌ها در طراحی الگوریتم‌ها تداوم می‌یابند:

- ناتوانی در مدل‌سازی سیستمی که در پی اعمال معیاری اجتماعی، مانند انصاف و عدالت، می‌باشد.
- شکست در درک این که کاربرد راه‌حل‌های الگوریتمی طراحی شده برای یک زمینه اجتماعی در زمینه اجتماعی دیگری می‌تواند گمراه‌کننده، نادرست و مضر باشد.
- ناتوانی در در نظر گرفتن ابعاد مختلف مفاهیم اجتماعی، مانند انصاف، که می‌تواند رویه‌ای، زمینه‌ای، و مناقشه‌انگیز باشد، که با ساختارهای ریاضی سازگاری ندارد.
- ناتوانی در درک این که چگونه رفتارها و ارزش‌های نهفته در سیستمی، با ورود فناوری به آن سیستم اجتماعی تغییر می‌یابند و ناتوانی در تشخیص این احتمال که شاید فناوری بهترین راه حل برای مشکلی نباشد.

واژه «سوگیری یا پیش‌داوری» (bias) در بیشتر موارد بار معنایی منفی‌ای دارد، اما در اینجا از آن برای اشاره به «کج‌روی از هنجاری» استفاده شده است (Danks and London, 2017, 4692). «کج‌روی می‌تواند در هر یک از مراحل طراحی، توسعه و استقرار روی دهد. داده‌های مورد استفاده در آموزش الگوریتم‌ها یکی از منابع اصلی بروز پیش‌داوری است (Shah, 2018)، چه به واسطه داده‌های نمونه‌گیری‌شده گزینشی و چه داده‌هایی که منعکس‌کننده سوگیری‌های اجتماعی موجود می‌باشد (Diakopoulos and Ko-liska, 2017; Danks and London, 2017; Binns, 2018a, 2018b, Malhotra, Kotwal, and Dalal, 2018). به عنوان مثال، نابرابری‌های ساختاری و از نظر اخلاقی مسئله‌دار که برخی از گروه‌های قومی را در شرایط نابرابر قرار می‌دهد ممکن است به وضوح در داده‌ها دیده نشود و بنابراین اصلاح نگردد (Noble, 2018; Benjamin, 2019). علاوه بر این، داده‌های مورد استفاده برای آموزش الگوریتم‌ها به ندرت «بر اساس طرح آزمایشی مشخصی» گردآوری می‌شوند (Olhede and Wolfe, 2018a, 3). آنها مورد استفاده قرار می‌گیرند، حتی اگر نادرست، یک‌سویه، یا به طور سیستماتیک پیش‌داورانه باشند و نمایانگر بدی از جمعیت مورد مطالعه باشند (Richardson, Schultz, and Crawford, 2019).

کنار گذاشتن برخی از متغیرهای داده‌ای هنگام آموزش سیستم‌های تصمیم‌گیری الگوریتمی می‌تواند تأثیرات این مشکل را کاهش دهد. در واقع، قوانین ضد تبعیض و حفاظت از داده‌ها، با هدف کاهش خطرات تبعیضات ناعادلانه، معمولاً پردازش متغیرهای حساسیت‌برانگیز یا «محافظت‌شده» (مانند جنسیت یا نژاد) که از نظر آماری اهمیت دارند، را محدود یا ممنوع می‌کنند. متأسفانه، حتی

مثال، انصاف الگوریتمی را می‌توان هم در رابطه با گروه‌ها و هم در رابطه با افراد تعریف کرد (Doshi-Velez and Kim, 2017). به همین دلیل و دلایل مشابه، به تازگی چهار تعریف اصلی از انصاف الگوریتمی توجهات زیادی را به خود جلب کرده‌اند (برای مثال، به Kleinberg, Mullainathan, and Raghavan, 2016; Corbett-Davies and Goel,

2020; Lee and Floridi, 2020 مراجعه کنید):

● ضد دسته‌بندی (anti-classification) که به مورد استفاده قرار نگرفتن رده‌بندی‌های محافظت‌شده، مانند نژاد و جنسیت، و نماینده‌هایشان (proxies) در تصمیم‌گیری‌ها اشاره دارد؛

● یکسانی دسته‌بندی (classification parity)، که یک مدل را در صورتی منصفانه می‌داند که معیارهای رایج پیش‌بینی، از جمله نرخ‌های مثبت و منفی کاذب، در بین گروه‌های محافظت‌شده برابر باشند؛

● واسنجی (calibration)، که انصاف را به‌عنوان معیاری از کیفیت واسنجی الگوریتمی بین گروه‌های محافظت‌شده در نظر می‌گیرد؛ و

● یکسانی آماری (statistical parity)، که انصاف را به‌عنوان برابری میانگین احتمالات برای همه اعضای گروه‌های محافظت‌شده تعریف می‌کند.

با این حال، هر یک از این تعاریف رایج از انصاف دارای اشکالاتی است و عموماً با یکدیگر ناسازگارند (Kleinberg, Mullainathan, and Raghavan, 2016). به‌عنوان مثال، در ضد دسته‌بندی، همان‌طور که در بالا ذکر شد (Gillis and Spiess, 2019)، حذف کردن ویژگی‌های محافظت‌شده مانند نژاد، جنسیت و مذهب از داده‌های آموزشی، با هدف جلوگیری از تبعیض، به سادگی ممکن نیست. نابرابری‌های ساختاری سبب می‌گردند که داده‌هایی که به خودی خود تبعیض‌آمیز نیستند، مانند کدپستی، می‌توانند به‌طور نیابتی (و به‌طور عمدی یا غیرعمدی) برای استنباط ویژگی‌های محافظت‌شده مانند نژاد عمل کنند (Edwards and Veale, 2017). به‌علاوه، در موارد مهمی، استفاده از ویژگی‌های محافظت‌شده برای تصمیم‌گیری‌های عادلانه خالی از اشکال است. به‌عنوان مثال، درصد پایین‌تر تکرار جرم برای زنان به این معنی است که حذف جنسیت به‌عنوان ورودی در الگوریتم‌های تکرار جرم، زنان را به شکل نامتناسبی در گروه‌های پرخطر قرار می‌دهد (Corbett-Davies and Goel, 2018). به همین دلیل، Binns (2018b) بر اهمیت در نظر گرفتن زمینه‌های تاریخی و جامعه‌شناختی، که امکان بازنمایی‌شان در داده‌های ورودی الگوریتم‌ها نیست، اما می‌توانند رویکردهای مناسب زمینه‌ای برای عدالت در الگوریتم‌ها را رهنمون باشند، تأکید دارد. همچنین توجه به این نکته ضروری است که مدل‌های الگوریتمی اغلب می‌توانند به نتایج غیرمنتظره‌ای برسند که برخلاف شهود انسانی و سبب آشفتگی فکری است؛ (Grgić-Hlača et al., 2018). به روشنی نشان می‌دهد که چگونه استفاده از ویژگی‌هایی که مردم آنها را منصفانه می‌دانند، در برخی موارد، می‌توان سبب افزایش نژادپرستی بروز داده شده

در یک کلینیک روستایی استفاده شود و فرض شود که کلینیک روستایی به همان اندازه بیمارستان پژوهشی دسترسی به منابع دارد، تصمیمات تخصیص منابع مراقبت‌های بهداشتی الگوریتم، نادرست و ناقص خواهد بود (Danks and London 2017). در همین راستا، Grgić-Hlača et al. (2018)

در مورد چرخه‌های معیوب زمانی که الگوریتم‌ها اقدام به ارزیابی‌های زنجیره‌ای نادرست می‌کنند، هشدار می‌دهند. در الگوریتم ارزیابی تهدیدات COMPAS، هنگام برآورد احتمال تکرار جرم، یکی از معیارهای مورد استفاده سابقه کیفری دوستان متهم است. از این رو، داشتن دوستانی با سابقه کیفری، چرخه معیوبی ایجاد می‌کند که در آن، متهمی که دوستان محکوم دارد، احتمال ارتکاب جرم و در نتیجه محکومیت به زندان بیشتری خواهد داشت. این امر، تعداد افراد دارای سابقه کیفری در یک گروه خاص را، صرفاً بر اساس یک همبستگی آماری، افزایش می‌دهد. (Richardson, Grgić-Hlača et al. 2018; Schultz, and Crawford 2019).

مواردی از سوگیری الگوریتمی پر سر و صدا در سال‌های اخیر - حداقل از زمان گزارش‌های افشاگرانه (Angwin et al., 2016) پیرامون سیستم COMPAS - منجر به تمرکز فزاینده‌ای بر مسایل مربوط به انصاف الگوریتمی شده است. تعریف و عملیاتی‌سازی انصاف الگوریتمی به «امری فوری در دانشگاه و صنعت» تبدیل شده است (Shin and Park, 2019)، که به افزایش قابل توجهی در تعداد مقالات، کارگاه‌ها و کنفرانس‌های اختصاص داده شده به «انصاف، پاسخگویی و شفافیت» (FAT - fairness, accountability, and transparency) منتهی شده است (Hoffmann et al., 2018; Ekstrand and Levy, 2019; Shin and Park, 2018). من در بخش بعدی به بررسی مباحث و مشارکت‌های کلیدی در این زمینه می‌پردازم.

۷.۶ تبعیض ناشی از شواهد ناعادلانه

در مورد نیاز به عدالت و انصاف الگوریتمی، به‌ویژه برای کاهش تهدید تبعیض‌های مستقیم و غیرمستقیم (به ترتیب «رفتار متفاوت» و «تأثیر متفاوت»، در قانون ایالات متحده) ناشی از تصمیمات الگوریتمی، توافق گسترده‌ای وجود دارد (Barocas and Selbst, 2016; Grgić-Hlača et al., 2018; Green and Chen, 2019). با این همه، بین پژوهش‌گران در مورد تعریف، ارزیابی، و هنجارهای انصاف الگوریتمی، اختلاف نظر وجود دارد (Gajane and Pechenizkiy, 2018; Saxena et al., 2019; Lee, 2018; Milano, Taddeo, and Floridi, 2019, 2020a). در (Wong, 2019) به بیست و یک تعریف از انصاف که در منابع پژوهشی پیدا می‌شود اشاره شده است، اما این تعاریفی اغلب با هم ناسازگار هستند (Doshi-Velez and Kim, 2017) و به نظر می‌رسد به این زودی‌ها نباید انتظار تعریفی جامع را داشت واقعیت این است که تفاوت‌های ظریف زیادی در تعریف، سنجش و کاربرد استانداردهای مختلف انصاف الگوریتمی وجود دارد. به عنوان

الگوریتم‌ها و کاهش دقت شد.

(Veale and Binns, 2017) و (Katell et al., 2020) دو رویکرد در مورد روش‌های بهبود انصاف الگوریتمی پیشنهاد می‌کنند. رویکرد اول، مداخله شخص ثالثی را در نظر دارد که در آن نهادی سوای تهیه‌کنندگان الگوریتم، داده‌های مربوط به ویژگی‌های حساسیت‌برانگیز یا محافظت‌شده را در اختیار دارد. آن نهاد آنگاه تلاش می‌کند تا تبعیض‌های ناشی از داده‌ها و مدل‌ها را شناسایی و کاهش دهد. رویکرد دوم، روش مشارکتی دانش محوری را پیشنهاد می‌کند که تمرکز بر منابع داده‌ای گروه‌های اجتماعی است که دربردارنده تجربیات عملی یادگیری ماشینی و مدل‌سازی می‌باشد. این دو رویکرد با هم تضادی ندارند؛ و بسته به زمینه‌های کاربریشان، می‌توانند مزایای متفاوتی داشته باشند، و همچنین ترکیب آنها خالی از فایده نیست.

با توجه به تأثیرات قابل توجه تصمیمات الگوریتمی بر زندگی مردم و اهمیت زمینه در انتخاب معیارهای مناسب برای سنجش انصاف، کمی تلاش‌ها برای جلب نظر عموم در مورد انصاف الگوریتمی تعجب برانگیز است (M. K. Lee, Kim, and Lizarondo, 2017; Saxena et al., 2019; Binns, 2018a). در پژوهشی در مورد برداشتهای مردم از تعاریف مختلف عدالت الگوریتمی، (Saxena et al., 2019)، ۳ نشان می‌دهد که چگونه در زمینه تصمیمات مربوط به گرفتن وام، مردم «تعریف واسنجی انصاف» با گزینش بر مبنای شایستگی را به «برخورد مشابه با افراد مشابه» ترجیح می‌دهند، و به هواداری از تبعیض جبرانی (affirmative action) می‌پردازند. در مطالعه‌ای مشابه، (Lee, 2018) شواهدی ارائه می‌دهد حاکی از آن که در مورد وظایفی که نیاز به مهارت‌های منحصر به فرد انسانی دارند، مردم تصمیمات الگوریتمی را کمتر منصفانه و الگوریتم‌ها را کمتر قابل اعتماد می‌دانند.

در گزارشی از آزمایشات انجام‌شده در مورد تفسیرپذیری و شفافیت الگوریتم‌ها، (Webb et al., 2019) نشان می‌دهد که ارجاعات اخلاقی، به‌ویژه در مورد انصاف، در بین شرکت‌کنندگانی که ترجیحات خود در مورد الگوریتم‌ها را به بحث می‌گذارند، ثابت است. این مطالعه خاطرنشان می‌کند که افراد تمایل دارند از ترجیحات شخصی خود فراتر بروند و به‌جای آن بر «رفتار درست و نادرست» تمرکز کنند؛ که نشانگر لزوم درک زمینه استقرار الگوریتم، در کنار دشواری‌های درک آن و پیامدهایش می‌باشد (Webb et al., 2019). در مورد سیستم‌های توصیه‌گر، (Burke, 2017) رویکردی چند جانبه و دارای دست‌اندرکاران متعدد برای تعریف انصاف پیشنهاد می‌کند که از تعاریف کاربر محور فراتر می‌رود و دربردارنده منافع سایر دست‌اندرکاران سیستم نیز می‌باشد.

روشن است که درک دیدگاه عمومی در مورد عدالت الگوریتمی به متخصصان فناوری کمک می‌کند تا الگوریتم‌هایی با رعایت اصول انصاف توسعه دهند که با احساسات عموم مردم در مورد مفاهیم رایج

عدالت همسو است (Saxena et al., 2019, 1). در نظر داشتن «دلایل قابل قبول از سوی آنانی که بیشترین تأثیرات منفی متوجه‌شان است» و «آمادگی برای جرح و تعدیلات موجه» هنگام طراحی الگوریتم‌ها از سوی سازندگان (Wong, 2019, 15) نقشی مهم در بهبود اثرات اجتماعی الگوریتم‌ها دارد. با این همه، توجه به این نکته که معیارهای انصاف در بیشتر موارد کاملاً ناکافی هستند اهمیت دارد، به خصوص هنگامی که در پی اعتبارسنجی مدل‌هایی هستند که در مورد گروه‌های اجتماعی‌ای به کار گرفته می‌شوند که به دلیل خاستگاه، سطح درآمد، یا گرایش جنسی خود در جامعه از قبل در محرومیت قرار دارند. به سادگی نمی‌توان پویایی‌های توان اقتصادی، اجتماعی و سیاسی موجود را «بهینه‌سازی» کرد (Winner 1980, Benjamin 2019).

۷.۷ چالش‌های استقلال و حریم خصوصی اطلاعاتی ناشی از اثرات تحول‌آفرین

اثرات جمعی الگوریتم‌ها، به بحث‌هایی در مورد استقلال کاربران نهایی دامن زده است (Ananny and Crawford 2018, Beer, 2017; Möller et al., 2018; Malhotra, Kotwal, and Dalal, 2018; Shin and Park, 2019; Hauer, 2019). خدمات مبتنی بر الگوریتم‌ها به‌طور فزاینده‌ای «درون اکوسیستمی از مسایل پیچیده اجتماعی-فنی» جای دارند (Shin and Park, 2019)، که می‌تواند سد راه استقلال کاربران شود. محدودیت‌های استقلال کاربران از سه منبع سرچشمه می‌شوند:

- توزیع فراگیر و پیش‌کنشگری (proactivity) الگوریتم‌های (یادگیری) در آگاه‌سازی در مورد گزینه‌های کاربران (Yang et al., 2018; Taddeo and Floridi, 2018a)؛
- درک محدود کاربران از الگوریتم‌ها؛
- فقدان قدرت مرتبه دوم (یا جذابیت) بر نتایج الگوریتمی (Rubel, Castro, and Pham, 2019).

ضمن بررسی چالش‌های اخلاقی هوش مصنوعی، (Yang et al., 2018, 11) به «تأثیر الگوریتم‌های مستقل و خودآموز بر استقلال فردی انسان» اشاره کرده و تأکید دارد که «قدرت پیش‌بینی هوش مصنوعی و تشویق و ترغیب بی‌وقفه آن، حتی اگر عمدی نباشد، باید کرامت و حق تعیین سرنوشت انسان را تقویت کند، نه تضعیف». خطرات سیستم‌های الگوریتمی که با شکل دهی گزینه‌های کاربران، مانعی سر راه استقلال انسان هستند، به‌طور گسترده در نوشته‌های پژوهشی بررسی شده‌اند و در اکثر اصول اخلاقی سطح بالا برای هوش مصنوعی، از جمله اصول EGE کمیسیون اروپا و کمیته هوش مصنوعی مجلس اعیان بریتانیا (Floridi and Cows, 2019) جایگاه ویژه‌ای را به خود اختصاص داده است. پیش از این (Floridi and Cows, 2019)، در بررسی این اصول سطح بالا، اشاره کردم این که الگوریتم‌ها از استقلال افراد پشتیبانی کنند کافی نیست. در

عظیم داده‌های حساسیت‌برانگیز به کار رفته در نمایه‌سازی (profil-ing) و پیش‌بینی‌های الگوریتمی، که بخش جدایی‌ناپذیر سیستم‌های توصیه‌گر هستند، مسایل متعددی را در مورد حریم خصوصی اطلاعات فردی ایجاد می‌کنند.

نمایه‌سازی الگوریتمی در طی دوره زمانی نامحدودی صورت می‌گیرد که در آن افراد بر اساس منطق داخلی سیستمی دسته‌بندی می‌شوند. با کسب اطلاعات جدید، نمایه‌های فردی به‌روزرسانی می‌شوند. این اطلاعات معمولاً یا مستقیماً بر مبنای تعاملات فرد با سیستم مشخصی، یا به‌طور غیرمستقیم از گروهی از افراد که به‌صورت الگوریتمی گرد آمده‌اند، استنباط می‌گردد (Paraschakis, 2017, 2018). در واقع، نمایه‌سازی الگوریتمی از اطلاعات گردآوری‌شده در مورد افراد و گروه‌های دیگری که شبیه فرد موردنظر تشخیص داده شده‌اند، نیز استفاده می‌کند. این اطلاعات دربردارنده ویژگی‌هایی چون موقعیت جغرافیایی و سن و نیز اطلاعات رفتاری و ترجیحات خاص، از جمله این که فرد در پیکرافزار (platform) خاصی بیشتر به‌دنبال چه نوع محتوایی است، می‌شود (Chakraborty et al., 2019). اگرچه این به نوعی یادآور مشکل شواهد ناکافی است، اما همچنین نشان می‌دهد که اگر حریم خصوصی گروهی (Floridi, and Taylor, 2014c; Floridi, 2016) تضمین نشود، ممکن است افراد هرگز نتوانند خود را از فرآیند نمایه‌سازی و پیش‌بینی‌های الگوریتمی خارج کنند (Milano, 2019, 2020a). به عبارت دیگر، حریم خصوصی اطلاعاتی فردی بدون تضمین حریم خصوصی گروهی (group privacy) دست‌یافتنی نیست.

شاید کاربران همیشه از نوع اطلاعاتی که در مورد آنها نگهداری می‌شود و اینکه از این اطلاعات در چه مواردی استفاده می‌شود، آگاهی نداشته باشند، یا نتوانند در مورد آن آگاهی کسب کنند. با توجه به نقش سیستم‌های توصیه‌گر در شکل‌دهی پویای هویت‌های فردی از طریق مداخله در گزینه‌های افراد، فقدان کنترل بر اطلاعات فردی به معنای از دست دادن استقلال است. فراهم آوردن امکان مشارکت افراد در طراحی سیستم‌های توصیه‌گر می‌تواند به ایجاد نمایه‌های دقیق‌تری کمک کند که ویژگی‌ها و گروه‌های اجتماعی‌ای را در نظر بگیرد که در غیر این صورت در برچسب‌گذاری مورد استفاده سیستم برای دسته‌بندی کاربران جایی نداشتند. اگرچه بهبود نمایه الگوریتمی در همه زمینه‌ها مطلوب نیست، بهبود طراحی الگوریتم‌ها براساس پیشنهادات دست‌اندرکاران مختلف الگوریتم‌ها، با پژوهش‌های اشاره شده در مورد RRI همگام بوده و توانایی کاربران در تعیین سرنوشت خود را بهبود می‌بخشد (Whitman, Hsiang, and Roark, 2018).

آگاهی از این که چه کسی مالک داده‌های شخصی است و در چه مواردی از آنها استفاده می‌شود، می‌تواند به ایجاد بده-بستان بین حریم خصوصی اطلاعاتی و مزایای پردازش اطلاعات نیز کمک کند (Sloan and Warner, 2018, 21). برای مثال، در زمینه‌های پزشکی،

عوض، استقلال الگوریتم‌ها باید محدود و برگشت‌پذیر باشد. فراتر از جوامع غربی، اصول هوش مصنوعی پکن - که توسط کنسرسیومی از شرکت‌ها و دانشگاه‌های پیشرو چین برای هدایت تحقیق و توسعه هوش مصنوعی توسعه یافته است - نیز دربردارنده مباحثی در مورد استقلال انسان می‌باشد (Roberts, Cowls, Morley, et al., 2021). ناتوانی در درک اطلاعات یا تصمیم‌گیری‌های نیز می‌تواند استقلال انسان را محدود کند. به بیان شین و پارک، الگوریتم‌ها «فاقد توانشی هستند که امکان درک، یا شناسایی بهترین روش بکارگیری‌شان برای دستیابی به اهداف مشخص را برای کاربران فراهم می‌آورد» (Shin and Park, 2019, 279). یکی از مسایل کلیدی که در مباحث مربوط به استقلال کاربران، دشواری ایجاد توازن مناسب بین تصمیم‌گیری خود افراد و تصمیم‌گیری‌هایی است که به الگوریتم‌ها واگذار می‌کنند. دیدیم که این مشکل به دلیل عدم شفافیت در فرآیند تصمیم‌گیری که از طریق آن تصمیمات خاص به الگوریتم‌ها واگذار می‌شود، پیچیده‌تر می‌گردد. (Ananny and Crawford, 2018) خاطرنشان می‌کند که این فرآیند، در بیشتر موارد، همه دست‌اندرکاران را در نظر نمی‌گیرد و از نابرابری‌های ساختاری مصون نیست.

اغلب از «طراحی مشارکتی» به دلیل تمرکز آن بر طراحی الگوریتم‌ها با هدف ارتقای ارزش‌های کاربران و حفاظت از استقلال آنها، به عنوان روشی برای پژوهش و نوآوری مسئولانه (responsible research and innovation - RRI)، یاد می‌شود (Whitman, Hsiang, and Roark, 2020, Katell et al., 2018). هدف طراحی مشارکتی «لحاظ کردن دانش ضمنی و تجربه عینی شرکت‌کنندگان در فرآیند طراحی» می‌باشد (Whitman, Hsiang, and Roark, 2018, 2). به عنوان مثال، چارچوب مفهومی «جامعه در جریان» (society-in-the-loop) در پی توانمندسازی دست‌اندرکاران مختلف در جامعه برای طراحی سیستم‌های الگوریتمی پیش از استقرار، و اصلاح یا معکوس کردن تصمیمات سیستم‌های الگوریتمی است که از قبل زمینه‌ساز فعالیت‌های اجتماعی هستند. این چارچوب با هدف برقراری «قرارداد اجتماعی الگوریتمی» کارسازی که به عنوان «پیمانی بین دست‌اندرکاران مختلف انسانی، با واسطه ماشین‌ها» تعریف می‌شود، طراحی شده است (Rahwan, 2018, 1) و این مهم را با شناسایی و حصول توافق در مورد ارزش‌های دست‌اندرکاران مختلف تحت تاثیر سیستم‌های الگوریتمی، به عنوان مبنایی برای نظارت بر پایبندی به قرارداد اجتماعی، انجام می‌دهد.

حریم خصوصی اطلاعاتی ارتباط نزدیکی با استقلال کاربر دارد (Co-hen, 2000; Rossler, 2015). حریم خصوصی اطلاعاتی، آزادی فردی برای تفکر، برقراری ارتباط و ایجاد رابطه، از جمله سایر فعالیت‌های ضروری انسانی را تضمین می‌کند (Rachels, 1975; Allen, 2011). با این حال، افزایش تعامل افراد با سیستم‌های الگوریتمی، توانایی آنها را در کنترل این که چه کسی به اطلاعات آنها دسترسی دارد و با آن چه می‌کند، را به شدت کاهش داده است. بنابراین، حجم

به علاوه، به دلیل ساختار و کارکرد بازارهای واسطه‌ای داده‌ها، در بسیاری از موارد «ردیابی داده‌ای تا منبع اصلی آن» پس از معرفی آن به بازار غیرممکن است (Crain, 2018, 93). حفاظت از اسرار تجاری، بازارهای پیچیده‌ای که فرآیند جمع‌آوری داده‌ها را از فرآیند خرید و فروش «جدا» می‌کنند، و ترکیبی از حجم زیادی از اطلاعات تولید شده به صورت محاسباتی بدون «منبع تجربی واقعی» در ترکیب با داده‌های واقعی، از جمله دلایل این امر می‌باشند (Crain, 2018, 94). پیچیدگی‌های فنی و پویایی الگوریتم‌های یادگیری ماشینی، به نگرانی‌هایی در مورد «عامل شویی» (agency laundering) برمی‌انگیزد: خطایی اخلاقی شامل دست شستن از اقدامات مشکوک اخلاقی، صرف نظر از این که آیا این اقدامات عمدی بوده‌اند یا خیر، و انداختن گناه به گردن الگوریتم‌ها است (Rubel, Castro, and Pham, 2019). این کار هم توسط سازمان‌ها و هم توسط افراد انجام می‌شود. (Rubel, Castro, and Pham, 2019) شامل مثالی ساده و تکان‌دهنده از عامل شویی توسط فیس‌بوک می‌باشد:

تیم پروپابلیکا (ProPublica) با استفاده از سیستم خودکار فیس‌بوک، یک دسته‌بندی ایجاد شده توسط کاربران به نام «متنفر از یهودیان» با بیش از ۲۲۰۰ عضو پیدا کرد... برای کمک به پروپابلیکا در یافتن مخاطبان بیشتر (و در نتیجه خرید آگهی‌های بیشتر)، فیس‌بوک تعدادی دسته‌بندی اضافی پیشنهاد داد ... پروپابلیکا از این پیکرافزار برای انتخاب سایر نمایه‌هایی که دسته‌بندی‌های یهود ستیزانه داشتند استفاده کرد و فیس‌بوک آگهی‌های پروپابلیکا را با تغییرات جزئی تأیید کرد. هنگامی که پروپابلیکا دسته‌بندی‌های یهود ستیزانه را فاش کرد و سایر رسانه‌های خبری نیز دسته‌بندی‌های نفرت‌انگیز مشابهی را گزارش کردند، فیس‌بوک در پاسخ توضیح داد که الگوریتم‌ها این دسته‌بندی‌ها را بر اساس اهدافی که کاربران مشخص کرده‌اند ایجاد کرده‌اند [و] «ما هرگز چنین کاربردی را در نظر نداشتیم یا پیش‌بینی نکرده بودیم که این قابلیت به این شکل استفاده شود». (Rubel, Castro, and Pham, 2019, 1024-5)

امروزه، کوتاهی در درک اثرات ناخواسته پردازش و تجاری‌سازی انبوه داده‌های شخصی، که مشکلی آشنا در تاریخ فناوری است (Wie-ner, 1950, 1989; Klee, 1996; Benjamin, 2019)، با توضیحات محدودی که اکثر الگوریتم‌های یادگیری ماشینی ارائه می‌دهند، همراه شده است. این رویکرد، خطر مسئولیت‌گریزی از نوع «کامپیوتر گفت» را به همراه دارد (Karppi, 2018). این امر می‌تواند سبب گردد متخصصان میدانی، مانند دست‌اندرکاران بخش بهداشت، از به زیر سوال بردن پیشنهادات الگوریتمی، حتی اگر عجیب به نظرشان برسد، پرهیز کنند (Watson et al., 2019). تأثیرات متقابل بین متخصصان میدانی و الگوریتم‌های یادگیری ماشینی می‌تواند به «سستی‌های معرفت‌شناختی» (Grote and Berens, 2020)، مانند

افراد تمایل بیشتری به اشتراک‌گذاری اطلاعاتی دارند که می‌تواند به تشخیص خود یا دیگران کمک کند. اما این موضوع در زمینه گزینش‌های شغلی کمتر صدق می‌کند. (Sloan and Warner, 2018) استدلال می‌کند که هنجارهای هماهنگی اطلاعات می‌تواند به تضمین انطباق صحیح این بده-بستان‌ها در زمینه‌های مختلف، و پرهیز از گذاشتن مسئولیت و تلاش بیش از حد بر دوش افراد کمک کند. اطلاعات شخصی باید در زمینه‌های اعمال قانون، در مقایسه با فرآیند گزینش شغلی، به شیوه‌ای متفاوت جریان یابد. GDPR اتحادیه اروپا نقش مهمی در ایجاد مبانی چنین هنجارهایی ایفا کرده است (Sloan and Warner, 2018).

پژوهش‌های روز افزون اخیر در مورد حریم خصوصی تفاضلی (differential privacy) روش‌های جدیدی در اختیار سازمان‌هایی قرار می‌دهد که به دنبال محافظت از حریم خصوصی کاربران خود هستند و در عین حال علاقمند به حفظ کیفیت مدل و معقول نگاه داشتن هزینه‌ها و پیچیدگی‌های نرم‌افزاری هستند. هدف این روش‌ها برقراری توازن بین سودمندی و حریم خصوصی است (Abadi et al., 2016; Wang et al., 2017; Xian et al., 2017). پیشرفت‌های فنی این چنینی، سازمان‌ها را قادر می‌سازد تا مجموعه داده‌ای را در دسترس عموم بگذارند و در عین حال اطلاعات مربوط به افراد را مخفی نگه دارند (از شناسایی مجدد جلوگیری کنند)؛ آنها همچنین می‌توانند محافظت از حریم خصوصی در مورد داده‌های حساسیت‌برانگیز، مانند داده‌های ژنتیکی، را تضمین کنند (Wang et al., 2017). به تازگی، Social Science One و فیس‌بوک از حریم خصوصی تفاضلی برای انتشار ایمن یکی از بزرگ‌ترین مجموعه‌های داده‌ای (۳۸ میلیون URL که به صورت عمومی در فیس‌بوک به اشتراک گذاشته شد) برای پژوهش‌های دانشگاهی در مورد تأثیرات رسانه‌های اجتماعی در جامعه، استفاده کردند (King and Persily, 2020).

۷.۸ مسئولیت‌های اخلاقی ناشی از توان ردیابی

محدودیت‌های فنی الگوریتم‌های مختلف یادگیری ماشینی، مانند فقدان شفافیت و توضیح‌ناپذیری، واری‌های آنها را دشوار کرده و لزوم رویکردهای جدید برای ردیابی مسئولیت اخلاقی و پاسخگویی در قبال عملکرد الگوریتم‌های یادگیری ماشینی را به روشنی نشان می‌دهد. در مورد مسئولیت اخلاقی، (Reddy, Cakici, and Ballester, 2019) به کم‌رنگ شدن مرز بین محدودیت‌های فنی الگوریتم‌ها و مرزهای گسترده‌تر قانونی، اخلاقی و نهادی که در آن عمل می‌کنند، اشاره دارند. حتی در مورد الگوریتم‌های بدون یادگیری، مفهوم‌سازی‌های خطی مرسوم در مورد مسئولیت، کارکرد محدودی در زمینه‌های اجتماعی-فنی حاضر دارند. ساختارهای اجتماعی-فنی گسترده‌تر، ردیابی مسئولیت عملکرد سیستم‌های پخش‌شده (distributed) دربردارنده انسان و کارگزاران مصنوعی را دشوار می‌سازد (Floridi, 2012b; Crain, 2018).

(که توصیه Buhmann, 2019, 13, Paßmann, and Fieseler بود). در همین راستا، Binns (2018b fv) مفهوم سیاسی-فلسفی «خرد عموم» (public reason) تمرکز می‌کند. با توجه به تفاوت‌های مسئولیت دهی اقدامات الگوریتم‌ها از نظر ماهیت و دامنه در بخش‌های دولتی و خصوصی، (Binns 2018b) خواستار ایجاد چارچوب مشترک عمومی می‌شود (همچنین ببینید Dignum et al., 2018). در این چارچوب، تصمیمات الگوریتمی می‌باید همان سطح از نظارت عمومی، که تصمیم‌گیری‌های انسانی دریافت می‌کنند، را دریافت کنند. چنین رویکردی توسط گروه زیادی از دیگر پژوهشگران نیز توصیه شده است (Ananny and Crawford, 2018; Blacklaws, 2018; Buhmann, Paßmann, and Fieseler, 2019).

دشواری‌های «عامل شویی» و «طفره رفتن اخلاقی»، همانطور که در فصل ۵ دیدیم، از ناکافی بودن چارچوب‌های مفهومی موجود برای ردیابی و اطلاق مسئولیت اخلاقی ناشی می‌شود. همانطور که در گذشته هنگام بررسی سیستم‌های الگوریتمی و تأثیر اقدامات آنها اشاره کرده‌ام:

ما با [اقدامات اخلاقی پخش‌شده] ناشی از تعاملات از نظر اخلاقی خنثی شبکه‌های (بالقوه ترکیبی) کارگزاران مواجهیم. به عبارت دیگر، چه کسی مسئول (مسئولیت اخلاقی پخش شده ...) برای [اقدامات اخلاقی پخش‌شده] است؟ (Floridi, 2016a, 2)

در آنجا، پیشنهادم این بود که مسئولیت اخلاقی کامل «به‌طور پیش‌فرض و با بالاترین ارجحیت» به تمام کارگزاران اخلاقی (مثل انسان یا مبتنی بر انسان، مانند شرکت‌ها) در شبکه که به‌طور سببی با کارکرد مشخصی در شبکه مرتبط هستند، اطلاق شود. رویکرد پیشنهادی ریشه در مفاهیم انتشار پس‌گرد (back-propagation) در نظریه شبکه‌ها، مسئولیت محض در علم حقوق، و دانش عمومی از منطق شناخت‌شناختی دارد. شایان ذکر است که این رویکرد، مسئولیت اخلاقی را از قصد و عمد کنشگران و از ایده مجازات و پاداش برای انجام عملی مشخص جدا می‌کند. در عوض، بر لزوم اصلاح اشتباهات (انتشار پس‌گرد) و بهبود عملکرد اخلاقی همه کارگزاران اخلاقی در شبکه، در راستای آنچه در قانون به عنوان مسئولیت محض شناخته می‌شود، تمرکز می‌کند.

۷.۹ فرجام: کاربرد خوب و ضرورانه الگوریتم‌ها

اخلاق الگوریتم‌ها، از سال ۲۰۱۶، به موضوع محوری بحث‌هایی بین پژوهشگران، ارائه‌دهندگان فناوری، و سیاست‌گذاران تبدیل گشته است. این بحث‌ها همچنین به دلیل بروز به اصطلاح تابستان هوش مصنوعی (به فصل ۱ مراجعه کنید) و همگام با استفاده فراگیر الگوریتم‌های یادگیری ماشینی مورد توجه قرار گرفته‌اند. ارگان‌هایی مانند EGE کمیسیون اروپا، کمیته هوش مصنوعی مجلس اعیان

جزماندیشی یا ساده‌لوحی (Hauer, 2019) بیانجامد، و سدی در راه انتساب مسئولیت در سیستم‌های پخش‌شده بشود (Floridi, 2016a). در تحلیلی، (Shah, 2018) نشان می‌دهد که چگونه می‌توان خطر کوتاهی در مسئولیت برخی از دست‌اندرکاران را، به عنوان مثال، با ایجاد نهادهای جداگانه برای نظارت اخلاقی بر الگوریتم‌ها، برطرف کرد. یکی از این نهادها، DeepMind Health است که هیئت بررسی مستقلی با دسترسی نامحدود (Murgia, 2018) تشکیل داد، تا این که شرکت گوگل در سال ۲۰۱۹ به فعالیت آن پایان داد. با این حال، انتظار این که یک نهاد نظارتی واحد مانند کمیته اخلاق پژوهشی یا هیئت بررسی نهادی «به تنهایی مسئول تضمین دقت، سودمندی و درستی کلان‌داده‌ها باشد» غیرواقع‌بینانه است (Lipworth et al., 2017). در واقع، برخی استدلال کرده‌اند که این ابتکارات فاقد هرگونه انسجام بوده و حتی می‌توانند به «آبی‌شویی اخلاقی» بیانجامد، که در فصل ۵ آن را به عنوان «اجرای اقدامات صوری در راستای ارزش‌ها و مزایای اخلاقی فرآیندها، محصولات، خدمات یا سایر کارکردهای رقمی به منظور این که از نظر اخلاقی رقمی بهتر از آنچه هستند به نظر آیند» تعریف کردیم. در مواجهه با رژیم‌های قانونی سختگیرانه، کنشگران کارکشته ممکن است به «واگذاری اخلاقی» روی آورند، که همان‌طور که در فصل ۵ دیدیم، در آن «فرآیندها، محصولات یا خدمات» غیراخلاقی به کشورهایی با چارچوب‌ها و سازوکارهای اجرایی ضعیف‌تر صادر شده و سپس، نتایج چنین فعالیت‌های غیراخلاقی به کشور مبدأ بازگردانده می‌شود.

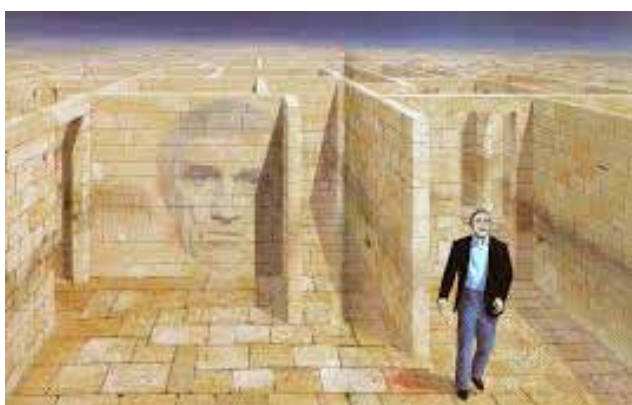
رویکردهای مشروح زیادی پاسخگو کردن الگوریتم‌ها وجود دارد. اگرچه میزانی از مداخلات فنی برای بهبود توضیح‌پذیری الگوریتم‌های یادگیری ماشینی ضروری است، بیشتر رویکردها بر مداخلات هنجاری متمرکز هستند. به باور آنانی و کرافورد، سازندگان الگوریتم‌ها، دست کم، می‌باید اقداماتی در راستای تسهیل گفت‌وگو عمومی در مورد فناوری خود انجام دهند (Ananny and Crawford, 2018). در پاسخ به اقدامات اخلاقی موردی (ad hoc)، پیشنهاد شده است که پاسخگویی باید در درجه اول به عنوان امری عرفی مورد توجه قرار گیرد (Dignum et al., 2018; Reddy, Cakici, and Ballestero, 2019). در راستای پُر کردن «شکاف» عرفی، و بر مبنای اصول هفت‌گانه الگوریتم‌های انجمن ماشین‌های محاسباتی (ACM)، در مقاله‌ای (Buhmann, Paßmann, and Fieseler, 2019) مطرح می‌کنند که سازمان‌ها باید از طریق آگاهی از الگوریتم‌ها، اعتبارسنجی و آزمایش آنها، مسئولیت الگوریتم‌های خود را، صرف‌نظر از عدم شفافیت‌شان، بر عهده بگیرند (Malhotra, Kotwal, and Da-lal, 2018). هنگام استقرار الگوریتم‌ها، باید عواملی چون مطلوبیت و زمینه گسترده‌تر کارکردشان را مدنظر داشت، که این همه می‌باید به «فرهنگ الگوریتمی» پاسخگوتری بیانجامد (Vedder and Naudts, 2017). برای گردآوری چنین ملاحظات، «مجامع و فرآیندهای تعاملی و گفت‌وگویی» با دست‌اندرکاران مربوطه می‌تواند سودمند باشد



دومینیک آپیا نقاش سوئسی فرا واقع گرابود. نقاشی‌های او سناریوهای تخیلی از شهرها، مناظر و فضاهای داخلی را به تصویر می‌کشد که عناصر آشنا مانند رویا در کنار هم قرار گرفته‌اند: یک کلیسای جامع تبدیل به ایستگاه راه‌آهن می‌شود، دریای مدیترانه از متروی پاریس فوران می‌کند و اتاقی با کودکان ناپدید شده در نقاط مختلف

تاریخ تولد: ۲۹ ژوئیه ۱۹۲۶، ژنو، سوئیس

فوت: ۸ ژانویه ۲۰۱۷، ژنو، سوئیس



نگاره یا نگاره‌هایی از دومینیک آپیا که در پایان هر قسمت این پی‌آیند، بی‌تکرار به شما تقدیم خواهیم کرد



بریتانیا، گروه کارشناسی HLEGAI کمیسیون اروپا و OECD (ببینید OECD 2019)، در قالب دستورالعمل‌ها و مبانی اخلاقی ملی و بین‌المللی (به فصل‌های ۴ و ۶ مراجعه کنید)، به بسیاری از پرسش‌های اخلاقی تحلیل شده در این فصل و متون پژوهشی بررسی شده، پرداخته‌اند.

تمرکز فزاینده بر استفاده از الگوریتم‌ها، هوش مصنوعی، و فناوری‌های رقمی به‌طور گسترده‌تر برای ارائه پیامدهای مثبت اجتماعی (Hager et al., 2019; Cowls et al., 2019; Cowls et al., 2021b)، یکی از جنبه‌هایی است که اگرچه به روشنی در نقشه مفهومی اولیه به آن اشاره نشده بود، در حال تبدیل شدن به یکی از موضوعات محوری مباحث در ادبیات مربوطه است. این موضوع فصل ۸ می‌باشد. این درست است، حداقل از نظر اصولی، که هر ابتکاری که هدفش کاربرد الگوریتم‌ها برای منافع اجتماعی است باید به‌طور رضایت‌بخشی به تهدیدات شش‌گانه مورد اشاره نیز و آنهایی که در فصل ۵ دیدیم بپردازد. اما بحث رو به گسترشی نیز در مورد اصول و معیارهای تعیین‌کننده در طراحی و تمشیت الگوریتم‌ها و فناوری‌های رقمی به‌طور کلی، و به‌طور مشخص در راستای منافع اجتماعی (همان‌طور که در فصل ۴ دیدیم) در جریان است. برای کاهش تهدیدات و در عین حال استفاده از توان مثبت این فناوری‌ها تحلیل‌های اخلاقی ضرورت دارند. چنین تحلیل‌هایی در خدمت اهداف دوگانه شفاف‌سازی ماهیت تهدیدات اخلاقی و توان مثبت این الگوریتم‌ها و فناوری‌های رقمی می‌باشند. تحلیل‌های اخلاقی سپس این درک را به رهنمودهای متقن و عملی برای تمشیت طراحی و کاربرد دست‌ساخته‌های رقمی تبدیل می‌کنند. اما پیش از حرکت در راستای این جنبه‌های مثبت‌تر، درک همه جانبه این که چگونه هوش مصنوعی می‌تواند برای اهداف شیطانی و غیرقانونی استفاده شود ضروری است. این موضوع فصل بعد است ...

ادامه دارد

هوش (دانایی) مصنوعی: راهنمایی برای انسان‌های متفکر

نویسنده مقاله: هوش (دانایی) مصنوعی، مدل بزرگ زبانی، چت‌جی‌پی‌تی، جی‌پی‌تی-۵، ۵

ویرایشگر: علیرضا خلیلیان، استادیار مهندسی نرم‌افزار

رایانشانی: akhalilian@gmail.com

شناسهٔ آرکید: ۷۰۳۸-۴۰۷۹-۰۰۰۲-۰۰۰۰

زین واقعه مدهوشم باهوشم و بی‌هوشم

هم ناطق و خاموشم هم لوح خاموشانم

«مولوی، دیوان شمس، غزلیات»

داشته است. برخلاف بسیاری از نویسندگان کتاب‌های تخصصی هوش مصنوعی که تمرکز خود را بر الگوریتم‌ها و جنبه‌های فنی قرار می‌دهند، میچل تلاش کرده است تا تصویری واقع‌بینانه از قابلیت‌ها، محدودیت‌ها و آینده هوش مصنوعی ارائه کند.



شکل ۱: تصویر ملانی میچل

میچل تحصیلات خود را در رشته علوم رایانش به پایان رساند و از شاگردان برجستهٔ داکلاس هافستادر^۲، دانشمند مشهور علوم شناختی و نویسندهٔ کتاب «گودل، اِشر و باخ: آمیزه‌ای گرانگ و ابدی» بوده است. همکاری علمی با هافستادر تأثیری عمیق بر دیدگاه او درباره هوش، شناخت و یادگیری گذاشت. بخش مهمی از پژوهش‌های وی به این پرسش اختصاص دارد که چگونه می‌توان ماشین‌هایی ساخت

3-Douglas Hofstadter

یادداشت ویرایشگر

این مقاله توسط هوش (دانایی) مصنوعی تولید شده است و ویرایشگر آن را بازخوانی، اصلاح و ویرایش کرده و شکل‌هایی به آن افزوده است. موضوع آن معرفی کتابی است با عنوان «هوش (دانایی) مصنوعی: راهنمایی برای انسان‌های متفکر» که به وسیلهٔ ملانی میچل^۱ نوشته شده است. پیش از ورود به مقاله لازم است دو نکته را توضیح دهیم. نخست گذاشتن «دانایی» داخل پرانتز اشاره به پیشنهاد ویرایشگر دارد که در مقاله‌ای در شمارهٔ ۲۸۲ گزارش کامپیوتر آن را به جای «هوش» برای شکل فعلی فناوری هوش (دانایی) مصنوعی نهاده است. دوم این که به نظر می‌رسد منظور نویسنده از واژهٔ «متفکر» در عنوان کتاب دعوت به سنجش‌گرانه‌اندیشی^۲ در مورد وضعیت فعلی فناوری هوش (دانایی) مصنوعی و پرهیز از شیفتگی یا هراس شتاب‌زده از آن است. اما ابتدا بگوییم ملانی میچل کیست.

معرفی نویسنده

ملانی میچل یکی از شناخته شده‌ترین پژوهشگران معاصر در حوزه هوش مصنوعی، علوم شناختی و سامانه‌های پیچیده است. آثار علمی و کتاب‌های او طی دو دههٔ اخیر نقش مهمی در تبیین توانایی‌ها و محدودیت‌های هوش مصنوعی برای پژوهشگران و عموم علاقه‌مندان

1-Melanie Mitchell

۲-Critical thinking: این معادل در کتاب «شفاف اندیشیدن» از نشر کرگدن دیده شد. ویرایشگر.

معرفی کتاب

کتاب «هوش (دانایی) مصنوعی: راهنمایی برای انسان متفکر» یکی از شناخته شده‌ترین کتاب‌های عمومی و تخصصی در حوزه هوش مصنوعی است که نخستین بار در سال ۲۰۱۹ منتشر شد و در سال ۲۰۲۵ با «پیشگفتار جدید» متناسب با تحولات سریع هوش مصنوعی مولد و مدل‌های بزرگ زبانی تجدید چاپ شد. این کتاب از سوی بسیاری از صاحب‌نظران به‌عنوان یکی از بهترین منابع برای آشنایی عمیق با هوش مصنوعی معرفی شده است؛ زیرا بدون گرفتار شدن در هیجان‌زدگی یا بدبینی، تصویری متوازن از دستاوردها، محدودیت‌ها و آینده این فناوری ارائه می‌دهد.

برخلاف بسیاری از کتاب‌های آموزشی که تنها به الگوریتم‌ها، برنامه‌نویسی یا ریاضیات هوش مصنوعی می‌پردازند، این کتاب می‌کوشد به پرسش‌های بنیادی‌تری پاسخ دهد؛ از جمله:

- هوش مصنوعی واقعاً چیست؟
 - آیا ماشین‌ها «می‌فهمند» یا صرفاً الگوها را پردازش می‌کنند؟
 - چرا برخی سامانه‌های هوش مصنوعی در انجام وظایف پیچیده بسیار موفق هستند، اما در انجام کارهای ساده انسانی دچار خطا می‌شوند؟
 - آیا دستیابی به هوش فراگیر مصنوعی امکان‌پذیر است؟
- به همین دلیل، کتاب نه تنها یک معرفی فناوری، بلکه تحلیلی از ماهیت «هوش» و تفاوت آن با محاسبه ماشینی نیز به شمار می‌آید.



شکل ۲: تصویر کتاب

هدف نویسنده از نگارش کتاب

ملانی میچل در مقدمه کتاب توضیح می‌دهد که هدف اصلی او، ارائه تصویری دقیق، علمی و واقع‌بینانه از وضعیت کنونی هوش مصنوعی است. او معتقد است که در رسانه‌ها و حتی در برخی محافل علمی، دو دیدگاه افراطی درباره هوش مصنوعی وجود دارد:

- گروهی معتقدند هوش مصنوعی به‌زودی جایگزین انسان خواهد شد؛
- گروهی دیگر توانایی‌های کنونی هوش مصنوعی را بیش از حد ناچیز می‌دانند.

که نه تنها اطلاعات را پردازش کنند، بلکه بتوانند مفاهیم را درک کرده، استدلال کنند و همانند انسان با موقعیت‌های جدید سازگار شوند.

میچل سال‌ها به‌عنوان پژوهشگر در مؤسسه‌های معتبر علمی فعالیت کرده و عضو هیئت علمی مؤسسه سانتافه بوده است. این مؤسسه یکی از مراکز پیشرو جهان در مطالعه سامانه‌های پیچیده، شبکه‌ها، هوش مصنوعی و علوم میان‌رشته‌ای محسوب می‌شود. پژوهش‌های او عمدتاً در مرز میان علوم رایانش، زیست‌شناسی، علوم شناختی و ریاضیات قرار دارد و از همین رو آثارش برای طیف گسترده‌ای از پژوهشگران قابل استفاده است. زمینه‌های اصلی پژوهشی او عبارتند از: هوش مصنوعی، یادگیری ماشین، علوم شناختی، سامانه‌های پیچیده، الگوریتم‌های تکاملی، پردازش مفاهیم، قیاس و استدلال و هوش فراگیر مصنوعی.

یکی از ویژگی‌های برجسته فعالیت‌های علمی میچل، تلاش برای پاسخ به این پرسش بنیادی است که «آیا ماشین واقعاً می‌تواند بفهمد؟». این پرسش محور بسیاری از مقاله‌ها و کتاب‌های او بوده است. از مهمترین کتاب‌های منتشر شده توسط وی می‌توان به موارد زیر اشاره کرد:

- هوش (دانایی) مصنوعی: راهنمایی برای انسان‌های متفکر
- گشتی هدفمند و برنامه‌ریزی‌شده در پیچیدگی
- مقدمه‌ای بر الگوریتم‌های ژنتیک

هر یک از این آثار در زمان انتشار خود از منابع معتبر آموزشی در زمینه هوش مصنوعی و سامانه‌های پیچیده محسوب شده‌اند. کتاب «مقدمه‌ای بر الگوریتم‌های ژنتیک» سال‌ها به‌عنوان یکی از منابع اصلی درس الگوریتم‌های ژنتیک در دانشگاه‌های جهان مورد استفاده قرار گرفته است. همچنین کتاب «گشتی هدفمند و برنامه‌ریزی‌شده در پیچیدگی» یکی از جامع‌ترین آثار عمومی درباره نظریه پیچیدگی و سامانه‌های پیچیده به شمار می‌آید.

از دیدگاه علمی، میچل علاوه بر تألیف کتاب، مقاله‌های پژوهشی متعدد در مجلات و گردهمایی‌های معتبر بین‌المللی منتشر کرده است. بسیاری از این مقاله‌ها به دلیل نگاه میان‌رشته‌ای، بارها توسط پژوهشگران علوم رایانش، زیست‌شناسی، علوم شناختی و حتی فلسفه ذهن استناد شده‌اند. در مجموع، جایگاه علمی ملانی میچل را می‌توان در سه محور خلاصه کرد:

- پژوهشگر برجسته در حوزه هوش مصنوعی و سامانه‌های پیچیده؛
- نویسنده آثار مرجع برای آموزش مفاهیم هوش مصنوعی و علوم شناختی؛
- مروج نگاه انتقادی و واقع‌بینانه نسبت به توانایی‌ها و محدودیت‌های هوش مصنوعی.

موفقیت‌های اخیر هوش مصنوعی، افزایش قدرت پردازش، دسترسی به داده‌های عظیم و پیشرفت الگوریتم‌های یادگیری عمیق را بررسی می‌کند.

بخش دوم: دیدن و ادراک. این بخش چهار فصل دارد و به بینایی ماشین و یادگیری عمیق اختصاص یافته است.

- فصل چهارم: تشخیص اشیاء، صحنه‌ها و ادراک تصویری
- فصل پنجم: شبکه‌های هم‌آمیزی^۴ و مجموعه داده مشهور ایمیجنت^۵
- فصل ششم: بررسی دقیق‌تر نحوه یادگیری ماشین‌ها
- فصل هفتم: اعتمادپذیری، عدالت و اخلاق در هوش مصنوعی

بخش سوم: یادگیری بازی. این بخش سه فصل را در بر می‌گیرد و در آن توضیح داده می‌شود که چگونه ماشین‌ها از طریق آزمون و خطا و دریافت پاداش، راهبردهای مناسب را فرا می‌گیرند و این روش چگونه در رباتیک، کنترل و تصمیم‌گیری به کار گرفته می‌شود.

- فصل هشتم: یادگیری تقویتی و پاداش
- فصل نهم: موفقیت سامانه‌هایی مانند آلفاگو^۶ در بازی‌ها
- فصل دهم: کاربردهای یادگیری تقویتی فراتر از بازی

بخش چهارم: هوش مصنوعی و زبان طبیعی. این بخش سه فصل دارد و در آن نویسنده به چالش‌های پردازش زبان طبیعی، مدل‌سازی معنا و توانایی ماشین‌ها در تولید و درک زبان می‌پردازد. هر چند نسخه جدید کتاب مقدمه‌ای درباره موج مدل‌های بزرگ زبانی نیز افزوده است، ولی ساختار اصلی فصل‌ها حفظ شده است.

- فصل یازدهم: نمایش معنای واژه‌ها و مدل‌های زبانی
- فصل دوازدهم: ترجمه ماشینی و فرایند رمزگذاری و رمزگشایی
- فصل سیزدهم: سامانه‌های پرسش و پاسخ و گفت‌وگو

بخش پنجم: مانع معنا. این بخش مهمترین و فلسفی‌ترین قسمت کتاب است و سه فصل پایانی را در بر می‌گیرد. در این بخش، نویسنده استدلال می‌کند که اگرچه سامانه‌های امروزی در انجام وظایف خاص عملکردی چشمگیر دارند، اما هنوز فاقد درک عمیق، دانش عمومی و عقل سلیم، توانایی انتزاع و قیاس انسانی هستند. از دیدگاه او، این فاصله یکی از مهمترین موانع دستیابی به هوش فراگیر مصنوعی است.

- فصل چهاردهم: مفهوم «فهم» در ماشین
- فصل پانزدهم: انتزاع، دانش و قیاس در هوش مصنوعی
- فصل شانزدهم: جمع‌بندی، پرسش‌های باز و آینده هوش مصنوعی

نویسنده می‌کوشد میان این دو دیدگاه تعادل برقرار کند و نشان دهد که هوش مصنوعی امروز در برخی وظایف بسیار موفق است، اما هنوز با هوش انسانی فاصله قابل‌توجهی دارد. او تأکید می‌کند که ارزیابی این فناوری باید بر پایه شواهد علمی باشد، نه تبلیغات یا پیش‌بینی‌های اغراق‌آمیز.

مخاطبان کتاب

یکی از ویژگی‌های برجسته این اثر، گستردگی جامعه مخاطبان آن است. کتاب به گونه‌ای نوشته شده که افراد با پیشینه‌های مختلف بتوانند از آن استفاده کنند. مخاطبان اصلی کتاب عبارتند از: دانشجویان مهندسی و علوم رایانش، دانشجویان هوش مصنوعی و علم داده، پژوهشگران علوم‌شناختی، استادان دانشگاه، متخصصان فناوری اطلاعات، سیاست‌گذاران حوزه فناوری، علاقه‌مندان به فلسفه ذهن و اخلاق هوش مصنوعی، خوانندگان عمومی که می‌خواهند با هوش مصنوعی آشنا شوند. یکی از نقاط قوت کتاب این است که نویسنده حتی هنگام توضیح مفاهیم نسبتاً پیچیده مانند شبکه‌های عصبی عمیق، یادگیری تقویتی یا پردازش زبان طبیعی، از زبانی روان و مثال‌های قابل فهم استفاده می‌کند. از این‌رو، کتاب برای دانشجویان مقطع کارشناسی نیز قابل استفاده است و در عین حال برای پژوهشگران نیز نکات تحلیلی ارزشمندی دارد.

ساختار کلی کتاب

کتاب در یک پیشگفتار و پنج بخش اصلی تنظیم شده است. این پنج بخش در مجموع شامل شانزده فصل هستند و خواننده را از تاریخچه هوش مصنوعی تا مباحث پیشرفته‌ای مانند فهم، معنا، انتزاع و هوش فراگیر مصنوعی هدایت می‌کنند. ساختار کلی کتاب به شرح زیر است:

پیشگفتار: نویسنده در آغاز کتاب به فضای هیجان، نگرانی و ابهامی می‌پردازد که پیرامون هوش مصنوعی شکل گرفته است. او توضیح می‌دهد که رسانه‌ها معمولاً میان دو تصویر کاملاً متفاوت از آینده هوش مصنوعی در نوسان هستند؛ آینده‌ای آرمانی یا آینده‌ای آخرازمایی. هدف کتاب، فاصله گرفتن از این دو نگاه و بررسی واقعیت‌های علمی است.

بخش اول: پیشینه. این بخش شامل سه فصل نخست کتاب است و به معرفی تاریخچه هوش مصنوعی اختصاص دارد.

● **فصل اول:** ریشه‌های هوش مصنوعی. در این فصل تاریخ شکل‌گیری هوش مصنوعی از دهه ۱۹۵۰، گردهمایی دارتموث، نقش پژوهشگرانی مانند آلن تورینگ، جان مک‌کارتی، ماروین مینسکی و دیگر پیشگامان این حوزه بررسی می‌شود.

● **فصل دوم:** شبکه‌های عصبی و ظهور یادگیری ماشین. در این فصل تحول شبکه‌های عصبی مصنوعی، یادگیری ماشین و نقش داده‌های بزرگ در پیشرفت هوش مصنوعی تشریح می‌شود.

● **فصل سوم:** بهار هوش مصنوعی. نویسنده در این فصل عوامل

4-Convolutional networks

5-ImageNet

6-AlphaGo

خلاصه‌ای از محتوای کتاب

تاریخچه هوش مصنوعی و سیر تحول آن

نویسنده کتاب را با این پرسش آغاز می‌کند که هوش مصنوعی چگونه متولد شد؟ او توضیح می‌دهد که اندیشه ساخت ماشین‌های هوشمند قدمتی چند صد ساله دارد، اما شکل‌گیری هوش مصنوعی به‌عنوان یک رشته علمی به گردهمایی دارتموث در سال ۱۹۵۶ بازمی‌گردد. در این گردهمایی، پژوهشگرانی مانند جان مک‌کارتی، ماروین مینسکی، کلود شانون و هربرت سایمون برای نخستین بار اصطلاح «آرتیفیشیال اینتلیجنس»^۷ را به کار بردند و این ایده را مطرح کردند که می‌توان رفتارهای هوشمندانه انسان را با استفاده از رایانه شبیه‌سازی کرد.

میچل سپس روند تاریخی پیشرفت این حوزه را بررسی می‌کند. در دهه‌های نخست، پژوهشگران بر این باور بودند که اگر قوانین منطقی و نمادهای کافی در اختیار ماشین قرار گیرد، رایانه قادر خواهد بود همانند انسان استدلال کند. این رویکرد که به «هوش مصنوعی نمادین» مشهور است، موفقیت‌هایی در حل مسایل ریاضی و بازی‌هایی مانند شطرنج به‌دست آورد، اما در مواجهه با مسایل واقعی و پیچیده با محدودیت‌های جدی روبه‌رو شد.

در ادامه، نویسنده به دوره‌ای اشاره می‌کند که به «زمستان هوش مصنوعی» مشهور است؛ دوره‌ای که به دلیل تحقق نیافتن وعده‌های اولیه، سرمایه‌گذاری‌ها و حمایت‌های علمی کاهش یافت. با این حال، پیشرفت سخت‌افزار، افزایش توان پردازشی رایانه‌ها، ظهور اینترنت و تولید حجم عظیمی از داده‌ها، زمینه را برای احیای دوباره این رشته فراهم کرد.

از دیدگاه نویسنده، مهمترین تحول در تاریخ هوش مصنوعی، گذار از برنامه‌نویسی مبتنی بر قوانین به یادگیری از داده‌ها است. در این رویکرد جدید، به جای آن که برنامه‌نویس همه قوانین را به ماشین آموزش دهد، الگوریتم‌ها از طریق مشاهده نمونه‌های فراوان، الگوهای موجود در داده‌ها را استخراج می‌کنند. این تغییر بن‌انگاره، زمینه‌ساز موفقیت‌های بزرگ یادگیری ماشین و یادگیری عمیق شد.

پیام اصلی نویسنده در این بخش آن است که پیشرفت هوش مصنوعی نتیجه یک مسیر خطی نبوده، بلکه حاصل دهه‌ها آزمون و خطا، شکست و نوآوری است. او تأکید می‌کند که شناخت تاریخ این حوزه برای درک واقع‌بینانه توانایی‌ها و محدودیت‌های آن ضروری است.

یادگیری ماشین و یادگیری عمیق؛ موتور محرک هوش مصنوعی امروز

در بخش دوم، نویسنده به معرفی فناوری‌هایی می‌پردازد که امروزه بیشترین سهم را در موفقیت هوش مصنوعی دارند؛ یعنی یادگیری ماشین و یادگیری عمیق. به اعتقاد میچل، تفاوت اصلی سامانه‌های هوشمند امروزی با نسل‌های نخست هوش مصنوعی در این است

که ماشین‌ها دیگر تنها مجری دستورالعمل‌های برنامه‌نویس نیستند، بلکه قادرند از داده‌ها بیاموزند. هرچه داده‌های آموزشی بیشتر و متنوع‌تر باشند، دقت مدل نیز افزایش می‌یابد.

نویسنده سپس شبکه‌های عصبی مصنوعی را معرفی می‌کند و توضیح می‌دهد که این شبکه‌ها با الهام از ساختار نورون‌های مغز انسان طراحی شده‌اند، هر چند شباهت آنها با مغز بسیار محدود است. او تأکید می‌کند که بسیاری از رسانه‌ها در بیان این شباهت دچار اغراق شده‌اند.

یکی از مثال‌های مهم کتاب، پروژه ایمپجنت است. نویسنده نشان می‌دهد که چگونه در اختیار داشتن میلیون‌ها تصویر برچسب‌خورده و استفاده از شبکه‌های عصبی عمیق، انقلابی در بینایی ماشین ایجاد کرد و دقت سامانه‌های تشخیص تصویر را به میزان چشمگیری افزایش داد.

نمونه دیگری که در کتاب بررسی می‌شود، سامانه آلفاگو از شرکت دیپ‌ماینده^۸ است. این سامانه توانست قهرمان جهانی بازی گو را شکست دهد؛ موفقیتی که سال‌ها غیرممکن تصور می‌شد. با این حال، میچل توضیح می‌دهد که موفقیت آلفاگو به معنای برخورداری آن از هوش فراگیر نیست، بلکه نتیجه آموزش بسیار گسترده در یک مسئله کاملاً مشخص است.

از دیدگاه نویسنده، یکی از بزرگترین کژفهمی‌ها درباره هوش مصنوعی این است که مردم موفقیت در یک وظیفه خاص را با «هوشمندی فراگیر و عمومی» اشتباه می‌گیرند. ماشینی که در تشخیص تصاویر یا بازی گو عملکرد فوق‌العاده‌ای دارد، ممکن است در انجام ساده‌ترین وظایف روزمره که برای یک کودک آسان است، ناتوان باشد.

پیام اصلی این بخش آن است که موفقیت‌های کنونی هوش مصنوعی عمدتاً حاصل قدرت یادگیری آماری از داده‌های عظیم است، نه نتیجه دستیابی ماشین به درک واقعی از جهان.

کاربردهای هوش مصنوعی در بینایی، بازی و پردازش زبان

در ادامه، نویسنده کاربردهای مهم هوش مصنوعی را بررسی می‌کند و نشان می‌دهد که چگونه فناوری‌های یادگیری ماشین در حوزه‌های مختلف به کار گرفته شده‌اند.

در زمینه بینایی ماشین، کتاب توضیح می‌دهد که رایانه‌ها می‌توانند اشیاء، چهره‌ها و الگوهای تصویری را با دقت بسیار بالا شناسایی کنند. این فناوری در خودروهای خودران، تصویربرداری پزشکی، سامانه‌های امنیتی و رباتیک کاربرد گسترده‌ای یافته است. با این حال، نویسنده مثال‌هایی ارائه می‌کند که نشان می‌دهد تغییرهای بسیار کوچک در تصویر می‌تواند موجب خطاهای عجیب در عملکرد سامانه شود.

در بخش مربوط به یادگیری تقویتی، کتاب توضیح می‌دهد که ماشین‌ها چگونه از طریق آزمون و خطا و دریافت پاداش، راهبردهای مناسب را فرا می‌گیرند. موفقیت آلفاگو نمونه‌ای از این نوع یادگیری است، اما نویسنده یادآور می‌شود که چنین سامانه‌هایی معمولاً در

8-DeepMind

7-Artificial Intelligence

آینده هوش مصنوعی و پیام نهایی کتاب

در بخش پایانی، نویسنده به آینده هوش مصنوعی و امکان دستیابی به هوش فراگیر مصنوعی می‌پردازد. او معتقد است که اگرچه پیشرفت‌های اخیر بسیار چشمگیر بوده‌اند، اما هنوز شواهد کافی برای نزدیک بودن به هوش فراگیر مصنوعی وجود ندارد.

از دیدگاه میچل، تحقق هوش فراگیر مستلزم حل مسایلی مانند درک معنا، انتزاع، استدلال، دانش عمومی و عقل سلیم و یادگیری انعطاف‌پذیر است؛ مسایلی که هنوز پاسخ روشنی برای آنها وجود ندارد. او پیشنهاد می‌کند که پژوهشگران به جای تمرکز صرف بر افزایش اندازه مدل‌ها یا حجم داده‌ها، باید به دنبال نظریه‌های عمیق‌تری درباره ماهیت هوش باشند.

پیام نهایی کتاب آن است که هوش مصنوعی یکی از مهمترین فناوری‌های قرن بیست‌ویکم است و می‌تواند در پزشکی، آموزش، صنعت، حمل‌ونقل و علوم مختلف تحولات بزرگی ایجاد کند؛ اما بهره‌برداری مسئولانه از این فناوری مستلزم شناخت دقیق قابلیت‌ها و محدودیت‌های آن است. نویسنده خواننده را به نگاهی سنجشگرانه و علمی دعوت می‌کند و یادآور می‌شود که آینده هوش مصنوعی نه با تبلیغات و پیش‌بینی‌های اغراق‌آمیز، بلکه با پژوهش‌های بنیادی، همکاری میان‌رشته‌ای و توسعه مسئولانه رقم خواهد خورد.

محیط‌های کاملاً تعریف‌شده و دارای قوانین ثابت آموزش می‌بینند و انتقال دانش آنها به مسایل جدید بسیار دشوار است.

یکی از جذاب‌ترین بخش‌های کتاب به پردازش زبان طبیعی اختصاص دارد. نویسنده توضیح می‌دهد که چگونه رایانه‌ها متن را به نمایش‌های عددی تبدیل می‌کنند، روابط میان واژه‌ها را می‌آموزند و در ترجمه ماشینی، خلاصه‌سازی یا پاسخ به پرسش‌ها از این دانش استفاده می‌کنند. با وجود پیشرفت‌های چشمگیر، او تأکید می‌کند که تولید متن روان لزوماً به معنای فهم واقعی زبان نیست.

در این بخش، نویسنده بارها تفاوت میان تشخیص الگو و درک معنا را برجسته می‌کند. از نظر او، بسیاری از سامانه‌های موفق امروزی بیش از آن که «معنا» را بفهمند، در شناسایی الگوهای آماری مهارت دارند. پیام اصلی این بخش آن است که کاربردهای کنونی هوش مصنوعی بسیار گسترده و ارزشمند هستند، اما نباید این موفقیت‌ها را با هوش انسانی یکسان دانست.

محدودیت‌های بنیادین هوش مصنوعی

این بخش، مهمترین و متمایزکننده‌ترین قسمت کتاب است. برخلاف بسیاری از آثار که تنها بر موفقیت‌های هوش مصنوعی تأکید دارند، ملانی میچل توجه خود را به محدودیت‌های بنیادین این فناوری معطوف می‌کند.

به اعتقاد او، مهمترین ضعف سامانه‌های کنونی، نبود دانش عمومی و عقل سلیم است. انسان از دوران کودکی مجموعه عظیمی از دانسته‌های بدیهی را درباره جهان کسب می‌کند؛ دانسته‌هایی که در تصمیم‌گیری‌های روزمره نقش اساسی دارند. برای مثال، انسان می‌داند که اگر لیوان واژگون شود، آب روی زمین می‌ریزد یا اگر جسمی از ارتفاع رها شود، سقوط خواهد کرد. این نوع دانش برای ماشین‌ها بسیار دشوار است.

نویسنده همچنین مفهوم شکنندگی را مطرح می‌کند. بسیاری از سامانه‌های هوش مصنوعی در شرایطی که اندکی با داده‌های آموزشی متفاوت باشد، عملکرد خود را از دست می‌دهند. در حالی که انسان می‌تواند دانش خود را به موقعیت‌های جدید تعمیم دهد، ماشین‌ها اغلب در انتقال آموخته‌های خود به مسایل جدید ناتوان هستند.

میچل همچنین به مسئله گرایه^۹ در داده‌ها می‌پردازد و توضیح می‌دهد که اگر داده‌های آموزشی دارای تبعیض یا خطا باشند، سامانه هوشمند نیز همان گرایه‌ها را بازتولید خواهد کرد. بنابراین، توسعه هوش مصنوعی صرفاً یک مسئله فنی نیست، بلکه ابعاد اجتماعی و اخلاقی نیز دارد.

پیام اصلی این بخش این است که هنوز فاصله قابل‌توجهی میان سامانه‌های هوشمند کنونی و هوش انسانی وجود دارد و نباید توانایی ماشین در انجام وظایف خاص را با فهم، آگاهی یا استدلال انسانی اشتباه گرفت.

برای دریافت اسلاید

سخنرانی‌های انجمن به

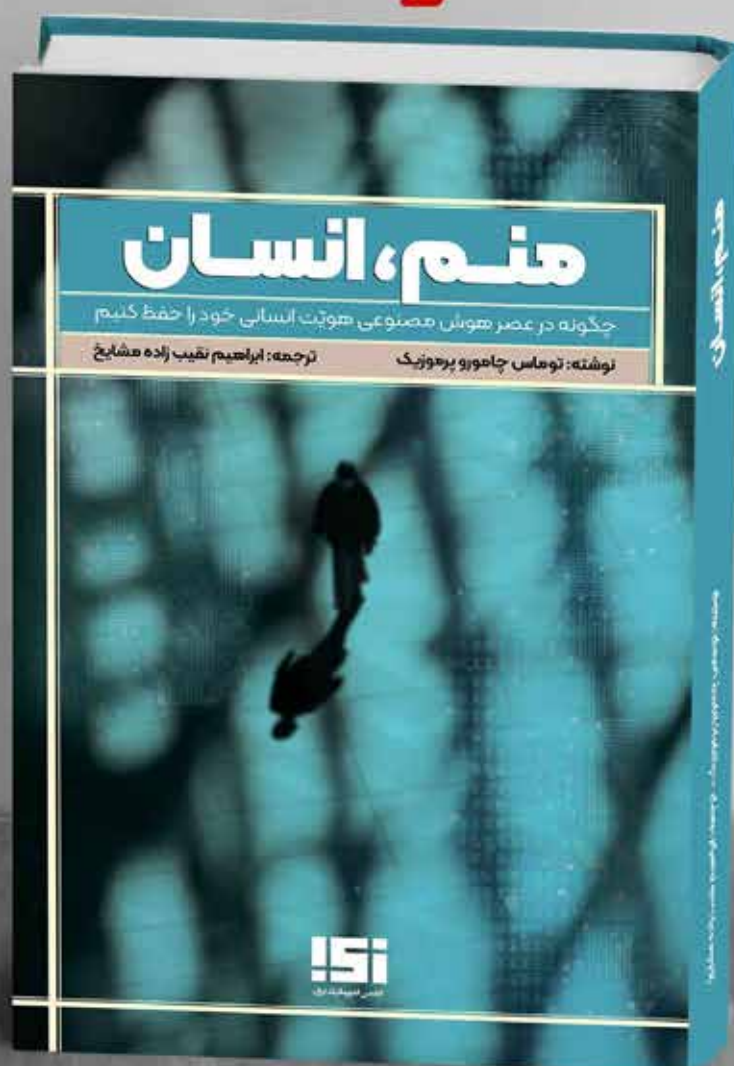
وبگاه انجمن مراجعه کنید

WWW.isi.org.ir

از انتشارات

انجمن انفورماتیک ایران

منتشر شد!



تهیه کتاب از دفتر

انجمن انفورماتیک ایران

۰۲۱-۶۶۴۱۲۸۶۱

انتشار شماره ۴۱ مجله علوم رایانشی

پیشنهادهای در صحت و بهبود قابل توجه در معیار یادآوری است که برای شناسایی نقص‌های واقعی اهمیت دارد. این روش با ارائه پایداری بالا در اعتبارسنجی ده‌تایی، توضیح‌پذیری قوی، و بهبود معیارهای کلیدی مانند یادآوری، رویکردی جامع برای پیش‌بینی نقص نرم‌افزار ارائه می‌دهد و می‌تواند در کاربردهای صنعتی و پژوهشی مورد استفاده قرار گیرد.

مقایسه عملکرد الگوریتم‌های DQN و Double DQN در محیط‌های یادگیری تقویتی با منابع محدود

سما نوری و راغولی، کاویان گنج کار، سینا صمدی قره‌ورن، کیمیا شیرینی

چکیده

یادگیری تقویتی عمیق طی سال‌های اخیر به‌عنوان یکی از رویکردهای پیشرفته در حل مسائل تصمیم‌گیری پیچیده مطرح شده است و توانسته در حوزه‌هایی مانند بازی‌های ویدئویی، رباتیک و کنترل خودکار عملکردی چشمگیر نشان دهد. با این حال، بررسی کارایی این الگوریتم‌ها در محیط‌هایی با منابع محاسباتی محدود کمتر مورد توجه قرار گرفته است. در این مقاله، دو الگوریتم مطرح الگوریتم یادگیری تقویتی عمیق شامل شبکه کیو عمیق Deep Q-Network (DQN) و شبکه کیو عمیق دوگانه مورد مقایسه و ارزیابی قرار گرفتند. هر دو الگوریتم با بهره‌گیری از شبکه‌های عصبی هم‌آمیختگی (CNN) و در محیط ساده و کم منبع گوگل کولب پیاده‌سازی شدند. آزمایش‌ها بر روی بازی Breakout انجام شد و با تنظیم حداقلی ابرپارامترها، عملکرد هر الگوریتم تحلیل گردید. نتایج نشان دادند که شبکه کیو عمیق دوگانه در برخی شبیه‌سازی‌ها دقت و پایداری بیشتری دارد، در حالی که vanilla DQN با وجود سادگی ساختار خود، همچنان کارایی قابل‌قبولی ارائه می‌دهد. به‌طور کلی، می‌توان نتیجه گرفت که در محیط‌های با محدودیت منابع، استفاده از رویکردهای

شماره ۴۱ مجله علوم رایانشی، نشریه علمی انجمن انفورماتیک ایران، در تابستان ۱۴۰۵ منتشر شد. در این شماره، ۷ مقاله به چاپ رسیده است که عنوان و چکیده آن‌ها برای اطلاع خوانندگان گزارش کامپیوتر در زیر آمده است. متن کامل مقالات از طریق وبگاه مجله علوم رایانشی به نشانی csj.isi.org.ir دسترس علاقه‌مندان قرار دارد

پیش‌بینی نقص نرم‌افزار با استفاده از یادگیری ماشین ترکیبی خودکار و تکامل ویژگی تطبیقی: رویکردی نوین با دقت و توضیح‌پذیری بالا

ابوالفضل دیباجی، احسان ناظر فرد، امیر کلباسی، علیرضا باقری

چکیده

پیش‌بینی نقص نرم‌افزار یکی از موضوعات مهم در مهندسی نرم‌افزار است که به بهبود کیفیت و کاهش هزینه‌های توسعه کمک می‌کند. در این پژوهش، روشی به‌نام یادگیری ماشین ترکیبی خودکار با تکامل ویژگی تطبیقی معرفی شده که با ترکیب پیش‌پردازش داده‌ها شامل حذف ویژگی‌های بی‌واریانس، پر کردن مقادیر گمشده با میانه و نرمال‌سازی، تکامل پویای ویژگی‌ها، متعادل‌سازی پیشرفته با Ne-arMiss، بهینه‌سازی چند-هدفه با Optuna، و ایجاد مدل گروهی جدید، دقت و توضیح‌پذیری را در پیش‌بینی نقص نرم‌افزار بهبود می‌بخشد. این روش روی مجموعه داده KC1 از مخزن PROMISE ناسا آزمایش شده و با تولید یازده ویژگی تطبیقی شامل نسبت‌های متریک، ترکیب‌های تعاملی، و تبدیل‌های لگاریتمی، و همچنین متعادل‌سازی کلاس‌ها، به صحت ۹۶/۹ درصد، معیار F1 برابر ۹۶/۸ درصد، یادآوری ۹۳/۸ درصد، و AUC برابر ۹۷/۲ درصد دست یافته است. تحلیل توضیح‌پذیری با SHAP TreeExplainer نشان داد که ویژگی‌های تعداد خطوط کد (loc) و معیار t نقش کلیدی در پیش‌بینی‌ها دارند. مقایسه با مدل‌های پایه (مانند رگرسیون لجستیک و جنگل تصادفی) و چندین مطالعه پیشین نشان‌دهنده برتری روش

ساده‌تر الگوریتم یادگیری تقویتی عمیق می‌تواند گزینه‌ای بهینه‌تر و کارآمدتر باشد.

تفکر رایانشی به مثابه پلی میان علوم رایانه و ریاضیات

پویا کریمی، ابوالفضل رفیع‌پور

چکیده

تفکر رایانشی در دو دهه اخیر به عنوان یکی از مفاهیم بنیادین علوم رایانه، جایگاهی محوری در توسعه فناوری‌های نوین یافته است. این مفهوم که در ایده‌های سیمور پاپرت ریشه دارد و توسط جنت وینگ به صورت نظام‌مند مطرح شد، فزاینده‌تر از مهارت برنامه‌نویسی، رویکردی شناختی برای تحلیل و حل مسائل پیچیده در حوزه‌های مختلف علمی و اجتماعی تلقی می‌شود. با وجود گسترش پژوهش‌ها، مرور ادبیات نشان می‌دهد که شکاف قابل توجهی میان بنیان‌های نظری علوم رایانه مانند طراحی و تحلیل الگوریتم‌ها، نظریه محاسبه، کلان داده، محاسبات موازی و هوش مصنوعی و کاربردهای آموزشی آن در ریاضیات وجود دارد. پژوهش‌های فنی عمدتاً بر توسعه چارچوب‌های محاسباتی متمرکز بوده‌اند، در حالی که مطالعات آموزشی بیشتر بر ابزارهایی نظیر برنامه نویسی اسکرچ، پایتون و لوگو یا تلفیق سطحی تفکر رایانشی در برنامه‌های درسی متمرکز داشته‌اند. مطالعه حاضر یک مقاله مروری است که با هدف پر کردن این شکاف، به بررسی جایگاه تفکر رایانشی در علوم رایانه و تبیین ظرفیت‌های آن در آموزش ریاضی می‌پردازد. یافته‌ها نشان می‌دهد که ریاضیات، به واسطه تأکید بر الگوریتم، تجرید و بازنمایی، بستری طبیعی برای پرورش تفکر رایانشی فراهم می‌کند و ادغام اصول فنی علوم رایانه در آموزش ریاضی می‌تواند درک عمیق‌تر مفاهیم، توسعه مهارت‌های حل مسئله، و ارتقای خلاقیت و تفکر انتقادی دانش‌آموزان را به همراه داشته باشد. در نهایت، این مقاله بر ضرورت طراحی چارچوب‌های میان‌رشته‌ای تأکید می‌کند که بر پایه‌ی بنیان‌های اصیل علوم رایانه استوار بوده و در عین حال متناسب با نیازهای آموزشی به‌ویژه در سطح متوسطه طراحی شوند. بدین ترتیب، پرورش تفکر رایانشی در آموزش ریاضی نه تنها به بهبود یادگیری منجر خواهد شد، بلکه نقشی کلیدی در آماده‌سازی نسل آینده برای مواجهه با چالش‌های پیچیده عصر دیجیتال ایفا خواهد کرد.

کشف تقلب در کارت‌های اعتباری با استخراج ویژگی از شبکه عصبی و طبقه‌بندی تقویت گرادیان سبک وزن

مائده نصیری، ابوالفضل تازی مرزآباد

چکیده

با افزایش استفاده از کارت‌های اعتباری، تشخیص خودکار

تراکنش‌های متقلبانه اهمیت بیشتری یافته است. در این پژوهش، دو مدل ترکیبی مبتنی بر شبکه عصبی پیچشی و پرسپترون چندلایه برای استخراج ویژگی به کار گرفته شد و ویژگی‌های استخراج شده به الگوریتم تقویت گرادیان سبک وزن منتقل گردید. برای انتخاب بهترین لایه، پنج معیار کمی شامل شاخص‌های کالینسکی- هاراباز، دیویس- بولدین، اطلاعات متقابل، میزان افزونگی و ابعاد لایه برای همه لایه‌ها محاسبه و لایه بهینه براساس میانگین نرمال شده معیارها تعیین شد. نتایج نشان داد مدل‌های ترکیبی عملکرد بهتری نسبت به مدل منفرد دارند و مدل مبتنی بر شبکه پیچشی بیشترین کارایی را ارائه کرده است. این مدل موجب افزایش امتیاز F1، کاهش قابل توجه خطای مثبت کاذب و افزایش اندک خطای منفی کاذب شد؛ اما این افزایش در برابر کاهش محسوس خطای مثبت کاذب قابل پذیرش بوده و تعادل میان دقت و بازیابی را حفظ می‌کند. آزمون ویلکاکسون تفاوت عملکرد را از نظر آماری معنادار نشان نداد، با این حال بهبود معیارهای عملکردی اهمیت استفاده از مدل ترکیبی مبتنی بر شبکه هم‌آمیختگی را در سناریوهای واقعی تشخیص تقلب برجسته می‌سازد. تحلیل SHAP نیز نشان داد ویژگی دوم استخراج شده از شبکه هم‌آمیختگی بیشترین نقش را داشته و تحت تأثیر متغیر «مقدار تراکنش» شکل گرفته است.

کارایی سازوکار توجه در شبکه‌های عصبی عمیق برای پیش‌بینی قیمت سهام: شواهد تجربی از بورس تهران

رامین خوچانی، یونس نادمی، مجید ابتیاع

چکیده

هدف از این پژوهش بررسی کارایی سازوکار توجه در ساختار ترکیبی شبکه‌های CNN و BiLSTM برای پیش‌بینی قیمت سهام بورس تهران است. در مدل پیشنهادی (CNN-BiLSTM-Atten-tion) ابتدا با لایه‌های هم‌آمیختگی ویژگی‌های غیرخطی محلی استخراج می‌شوند و سپس با BiLSTM به صورت دوطرفه پردازش می‌شوند. سازوکار توجه نیز به صورت خودکار به بخش‌های مهم توالی قیمتی وزن‌دهی می‌کند. داده‌های تجربی شامل قیمت‌ها و شاخص‌های تکنیکی پنج سهم منتخب (فولاد، فملی، وبملت، فارس، خودرو) در بورس تهران در بازه ۱۴۰۰/۱/۱ تا ۱۴۰۴/۵/۱۵ (مجموعاً ۹۲۹ رکورد برای آموزش و آزمون) استفاده شد. داده‌ها به صورت ۸۰ درصد برای آموزش و ۲۰ درصد برای آزمون تقسیم شدند. نتایج کلیدی نشان داد مدل CNN-BiLSTM-Attention در مقایسه با سایر مدل‌های مرجع به طور معنی‌داری دقت پیش‌بینی بالاتری ارائه می‌کند. در میانگین عملکرد برای پنج سهم منتخب، مدل پیشنهادی در معیار RMSE به مقدار ۰/۰۳۵ و در معیار MAPE به مقدار ۰/۰۴۶ دست یافت، در حالی که بهترین مدل

می‌باشد، نشان‌دهنده بهبود مناسبی در عملکرد مدل پیش‌بینی نقص نرم‌افزار می‌باشد.

نحوه تعامل مخاطبان بازی‌های رایانه‌ای با بررسی دستگاه‌های ورودی و خروجی در یونیتی

امیرحسین اردلان

چکیده

تعامل کاربر با محیط بازی، هسته اصلی تجربه بازیکن را تشکیل می‌دهد و این تعامل از طریق دستگاه‌های ورودی و خروجی متنوعی صورت می‌گیرد. این مقاله به بررسی جامع نحوه تعامل مخاطب در بازی‌های رایانه‌ای با تمرکز بر عملکرد دستگاه‌های ورودی و خروجی و پیاده‌سازی آنها در موتور بازی‌سازی یونیتی (Unity) می‌پردازد. ابتدا، دستگاه‌های ورودی و خروجی رایج و حیاتی در بوم‌سازگان بازی‌های رایانه‌ای (مانند صفحه کلید، موش و صفحه نمایش) معرفی شده و اصول فنی کارکرد آنها تشریح می‌گردد. سپس، با بررسی کدهای نمونه و API‌های مرتبط در یونیتی، چگونگی دریافت، پردازش، و تفسیر این ورودی‌های فیزیکی و ترجمه آنها به عملیات درون بازی مورد تحلیل قرار می‌گیرد و شیوه ارتباط آنها با سخت‌افزار توضیح داده می‌شود. هدف این پژوهش، ایجاد پلی میان مبانی سخت‌افزاری دستگاه‌های تعاملی و پیاده‌سازی نرم‌افزاری آنها در یک موتور بازی‌سازی قدرتمند مانند یونیتی است. درک این چرخه کامل، از لحظه ورود فیزیکی تا ترجمه آن به عملیات درون بازی، برای توسعه‌دهندگان بازی و پژوهشگران این حوزه ضروری بوده و می‌تواند به طراحی رابط‌های کاربری بهینه‌تر و خلق تجربه‌های بهتری در بازی‌ها منجر شود.

مرجع (CNN-LSTM) به ترتیب مقادیر ۰/۰۵۲ برای RMSE و ۰/۰۶۶ برای MAPE را ثبت کرد. این بهبود دقت که در معیار ضریب تعیین (R) نیز آشکار است (رسیدن به میانگین ۰/۸۶۵ در برابر ۰/۸۰۳ در مدل مرجع)، اثربخشی هم‌کردار هم‌آمیزی، BiLSTM و سازوکار توجه را تأیید می‌کند. اگر چه پیچیدگی محاسباتی این ساختار نسبتاً بیشتر است، اما دقت بسیار بالاتر در پیش‌بینی قیمت‌های سهام، آن را برای کاربردهای سرمایه‌گذاری در بورس تهران مفید می‌سازد. این پژوهش نخستین مطالعه‌ای است که سازوکار توجه را در پیش‌بینی قیمت سهام بورس تهران به کار گرفته و یافته‌ها نشان می‌دهد مدل پیشنهادی می‌تواند به ابزاری کارآمد برای تحلیلگران و معامله‌گران تبدیل شود.

پیش‌بینی بیدرنگ نقص نرم‌افزار با یادگیری چندوجهی مبتنی بر مبدل‌ها

سمیرا کرامت طلایه، محمدرضا ابراهیمی دیشابی، حامد پزشکی، حلیمه مددی

چکیده

در دنیای توسعه نرم‌افزار، که سامانه‌های پیچیده در چرخه‌های «یکپارچه‌سازی مداوم/استقرار مداوم»^۱ تکامل می‌یابند، پیش‌بینی بیدرنگ نقص‌ها به چالشی اساسی بدل شده است. این توانمندی افزون بر کاهش هزینه‌های نگهداری، ریسک‌های عملیاتی را کاهش داده و کیفیت نرم‌افزار را بهبود می‌دهد. رویکردهای سنتی پیش‌بینی نقص عموماً بر روش‌های آماری ساده یا شبکه‌های عصبی پیشرفته و تخصصی مانند CNN و LSTM متکی هستند و با محدودیت‌هایی چون تأخیر پردازشی، ناتوانی در پردازش داده‌های بیدرنگ و تمرکز بیشتر بر داده‌های تک‌وجهی به جای ساختارهای چندوجهی مواجه هستند. در محیط‌های پویای CI/CD، این رویکردها کارایی کافی نداشته و غالباً به صورت برون‌خط عمل می‌کنند و در نتیجه تشخیص زود هنگام نقص‌ها با محدودیت مواجه می‌شود. این پژوهش روشی مبتنی بر یادگیری چندوجهی با معماری مبدل ارائه می‌دهد که پیش‌بینی بیدرنگ نقص‌های نرم‌افزاری را بهبود می‌بخشد. روش پیشنهادی با بهره‌گیری از داده‌های ناهمگن شامل کد منبع، تاریخچه تغییرات، پیام‌های کامیت، معیارهای کد و رخدادهای اجرایی و با استفاده از رمزگذارهای تخصصی و همجوشی تطبیقی مبتنی بر توجه، الگوهای پیچیده را استخراج می‌کند. نتایج آزمایش‌های تجربی روی مخازن کد واقعی نشان‌دهنده کاهش تأخیر زمانی به ۲۳ میلی‌ثانیه می‌باشد که با نیازمندی سیستم بیدرنگ (کاهش تأخیر و کاهش مصرف انرژی) سازگار می‌باشد. همچنین میزان فراخوانی ۰/۷۵ که بیانگر افزایش سه درصدی معیار فراخوانی نسبت به مدل‌های تک‌وجهی و پنج درصدی آن نسبت به مدل‌های شبکه‌های عصبی پیشرفته

1- Continuous Integration/Continuous Deployment (CI/CD)

www.isi.org.ir

برای دریافت
اسلاید سخنرانی‌های انجمن
به وبگاه انجمن
مراجعه کنید.

از انتشارات انجمن انفورماتیک ایران
چاپ دوم

مهارت‌های نرم

راهنمای زندگی تولیدکنندگان نرم‌افزار

نوشته جان سانمز
ترجمه ابراهیم نقیب‌زاده مشایخ



زنگ تفریح

راهنمای معماری سازمانی

یک معمار سازمانی، یک مهندس سخت‌افزار و یک تحلیل‌گر کسب و کار در جزیره دورافتاده‌ای گیر افتاده بودند و تنها یک قوطی کنسرو داشتند، اما در بازکن همراهشان نبود. تحلیل‌گر کسب و کار گفت «بیاید تحلیل سودبران و ارزیابی هزینه-منفعت انجام دهیم تا کارآمدترین روش برای باز کردن این قوطی را پیدا کنیم».

مهندس سخت‌افزار پیشنهاد کرد «نظرتان چیست تا ماشینی از برگ نخل و نارگیل بسازیم که بتواند در قوطی را برایمان باز کند؟» معمار سازمانی نگاهی به آنها انداخت، قوطی را برداشت و به هوا پرتاب کرد. وقتی قوطی به زمین خورد، دهانه‌اش باز شد. مهندس سخت‌افزار و تحلیل‌گر کسب و کار که شگفت‌زده شده بودند، پرسیدند «چطور این کار را کردی؟» معمار سازمانی لبخندی زد و گفت «فقط کافیه خارج از چهارچوب قوطی فکر کنید!»

قانون تحویل سیستم

تخمین خود را دو برابر کنید و آن را با واحد زمانی بعدی جایگزین نمایید. برای مثال، اگر تخمین اولیه شما ۶ هفته است، دو برابر آن می‌شود ۱۲ هفته و با توجه به این که واحد زمانی بعدی ماه است، زمان تحویل سیستم می‌شود ۱۲ ماه!

رابطه زمان و پیچیدگی

مدت زمان اختصاص یافته به هر یک از موارد دستور کار، با مبلغ در نظر گرفته شده برای آن (یعنی پیچیدگی موضوع) نسبت عکس دارد

قانون ۹۰-۹۰

ساخت ۹۰٪ اول هر نرم‌افزار، ۹۰٪ زمان را به خود اختصاص می‌دهد و ۱۰٪ باقیمانده آن نیز ۹۰٪ دیگر زمان را می‌گیرد.

قوانین مورفی

۱- برنامه‌های ساده هیچگاه برای اولین بار کار نمی‌کنند. برنامه‌های پیچیده هم هیچ وقت کار نمی‌کنند.

۲- اشکال‌زدایی، دو برابر سخت‌تر از نوشتن کد از اول است. بنابراین، اگر کد را تا حد امکان هوشمندانه بنویسید، طبق تعریف، به اندازه کافی باهوش نیستید که آن را اشکال‌زدایی کنید.

۳- یک اشتباه رایج که مردم هنگام تلاش برای طراحی یک سیستم کاملاً بی‌نقص مرتکب می‌شوند، دست کم گرفتن نیوغ احمق‌هاست.

قانون سودمندی

سودمندی = لگاریتم فناوری

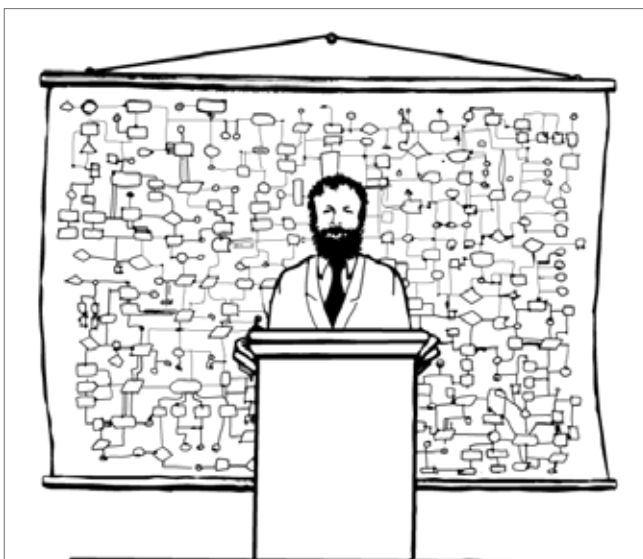
بنابراین، برای دستیابی به بهبود خطی در سودمندی در گذر زمان، لازم است که فناوری رشدی نمایی داشته باشد.

رابطه نرم‌افزار و سخت‌افزار

نرم‌افزار سریع‌تر از سریع‌تر شدن سخت‌افزار کند می‌شود!

دموی نرم‌افزار

احتمال بروز یک خطا در نرم‌افزار، هنگام دموی آن، چهار برابر می‌شود!



اکنون که با مرور کلی سیستم آشنا شدید، آماده‌ایم تا کمی بیشتر وارد جزئیات شویم!

فهرست اعضای حقوقی

انجمن انفورماتیک ایران

اعضای حقوقی فعال در حوزه راه حل های برنامه ریزی منابع سازمان و نرم افزارهای پیشرفته سازمانی

آریانا پرداز آینده



آسان سازان ستاره نوین



ارمغان پارس پرداز



رستاک سامانه ویرا



پرنده های هدایت پذیر از دور



پارس تصمیم



پردازش موازی سامان



پوپک سیستم پارت



بن مایه تحول هوشمند



توسعه یکپارچه ایلیا



توسعه سامانه های مدیریت برید



توسعه فناوری ایده ماندگار



چارگون



سبز داده افزار



سبز افزار پرگاس



سیستم های اطلاعات مدیریت شرق رایا



سورن سامانه مدیریت هوشمند



شرکت داده پردازان پرسیس پویا



گروه نرم افزاری پیوست



گلرنگ سیستم



فرا بوم کسب و کاری نوآوری باز



فرا رو آرمان رایان وفا



فناوری اطلاعات آرادپرداز اسپادانا



فناوری اطلاعات و ارتباطات پاسارگاد آریان (فناپ)



داده ورزی فرادیس البرز



نماتک ایرانیان



مدیریت سیستم های دیجیتال



مهندسی نرم افزاری رایورز



مبنا داده ارتباط شبکه



داده پردازی نیلرام مانا



گروه مشاورین ورنانگر نوین



گام الکترونیک



هورماه رابین خاور



تحول هوشمند ایلیا



تحول هوشمند ایلیا



اعضای حقوقی فعال در حوزه نرم افزارهای کاربردی

آدرین ارتباط حامی



ارقام نگار اندیشه



ابر کاکیا رایانه



ارتباطات جاده ای دیجیتال آسمان



اطمینان فناوری امید



شرکت پارسا نواوران سامان ایرانیان



اطلاع رسانی پیوند داده ها



آریانا پرداز آینده (آریا)



پیک عصر نوین تجارت مجازی

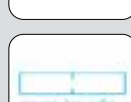


پیشرو تجارت سازان کارا



فهرست اعضای حقوقی

انجمن انفورماتیک ایران

هم اندیشان یگانه تک پرداز		رفاه صنعت آذرخش		بهسازان ملت	
مدیریت صنعت نکو		رایان پرتو نگار		بهین رایان نواندیش داده	
مهندسی راستای دانش		رایانه فیدار سپاهان		پیشگامان فن آور هوداد	
مهندسی راهکار آفرین آدا		رایا محاسب تطبیق		پیشازان سبب طلایی	
مجموع داده ها و سیستم ها (MDS)		راهکار فناوری آرتا		پیشرو نگاه زرین	
فرادید ارتباط نیلگون		سپهر اندیش حساب آسیا		پندار کوشک ایمن	
فناوری سلامت جست		سپها آکو وب		پژند الکترونیک	
ققتوس پردازش سرمایه		سروش رایانه ایرانیان		پردازش ابری نو واژه	
فرآیندهای یاسان		سازه های اطلاعاتی و ارتباطی سامان		پندار پاکان پاندرا	
همراهان سیستم گوهر		ستاره فناوری اطلاعات خاورمیانه		تدبیر گران توسعه انرژی اترک	
همراه الکترونیک آوای پارسی		سپهر		توسعه فناوری اطلاعات جهان افزارنوبین	
نوبن آوازه گران فرا وب		شایگان سیستم		توسعه تجارت هوشمند ماد	
مهندسی رز اندیشه هوشمند		شرکت اطلاع رسانی پیوند داده ها		تصمیم یاران بهینه سپند	
نواوران علوم مکانی		شرکت داده پرداز ماکان سیستم دانشمند		خدمات مهندسی فن آوری های طیف	
ماهان وب گستر آویژه		صنایع فن آوری طراحان بهینه		داده کاوان پیشرو ایده ورائگر	
گروه فناوری اطلاعات آتیه ویستانگر		صالح اندیشان ماندگار		دانبار توسعه پایدار علوم	
		طرح تجارت پردیس نیکان		دیده نگار خلاق	

فهرست اعضای حقوقی

انجمن انفورماتیک ایران

اعضای حقوقی فعال در حوزه اینترنت و پورتال‌ها

ارتباطات مبین نت



داده پردازان ایده شرق



داده پردازای پویای شریف



روند تازه



رایا پروژه



دانش افزار نارون شریف



فرابرد داده‌های ایرانیان



شرکت کلید گستر بینا
(نت بینا)



شبکه سازان باوان وب



شاهراه تجارت و فناوری ابربازار



مهندسی سازه اطلاعات سامان



هاسن ایران



هوشمند پرداز پویا تجارت قدیر



اعضای حقوقی فعال در حوزه شبکه و سخت‌افزار

آکام پردازش پارس



آژند پارت روشنا



توسعه تجارت الکترونیک نگین توسن



شرکت مهندسی آتی سازاب ایرانیان



شبکه الکترونیکی پرداخت کارت
(شاپرک)



شرکت توسعه ارتباطات الکترونیک
تجارت ایرانیان



شرکت صنایع الکترونیک فاران



داده‌پردازان هوشمند آریا توسعه رده



رایانه خدمات امید



زنجیره تبادل هوشمند روشن



شرکت مشاوره رتبه بندی اعتباری ایران



فن‌آوران اطلاعات آسام



فناوری اطلاعات ناواکو



گرایش تازه کیش



گسترش فناوری‌های نوین



صنایع پرسو الکترونیک



شرکت ملی انفورماتیک



مدیریت امن الکترونیکی کاشف



اعضای حقوقی فعال در حوزه بانکی و بانکداری الکترونیکی

آمیتیس همتا



بهسازان ملت



پدیسار انفورماتیک ایران



پردازشگران سامان



پردازش هوشمند هزارستان



پرداخت الکترونیک سداد



پویا



پردازش اطلاعات مالی پارت



پایا انرژی باتاب



پایا تراکنش هزاره سوم



تجارت الکترونیک پارسیان کیش



توسعه فناوری کیان پرداز زاگرس



شرکت راهکارهای هوشمند و یکپارچه آسا



شرکت سکوی کسب و کار الکترونیک



شرکت توسعه سامانه نرم‌افزاری نگین
(توسن)



خدمات انفورماتیک



فهرست اعضای حقوقی

انجمن انفورماتیک ایران

ارتباط داده های فراایده



ایمن ارتباط سبز کاسپین



امین دقیق سامانه



افرا کارا نت



افرا گستر لیان پارس



ابتکار کارو



افرا رابین ارتباط کوشا



شرکت ارتباطات برتر فرادید



ایده آل ارتباطات قرن



اوژن تدبیر پارس



اکسین ایمن نیکراد



پردازش افزار دانا



پیشه و فن آینه سان



پردازش سرو نیوان



پترو تجهیز کیان داریا



پشتیبان زیر ساخت امید



پیشرو گستر افرا مهر



پارس تکنولوژی سداد



بانی رایان پرداز نو



باربن سخت افزار



بین المللی کیکاوس زمان



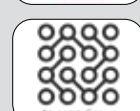
بهاور فناوری ویرا



تسلا الکترونیک پیشتازان



توسعه فناوری و نوآوری جویان



توسعه ارتباطات پیشبرد



توسعه انفورماتیک کارن



تولیدی بازرگانی اشجع باتری



توسعه و تجهیز فدک رایان



توسعه فناوری اطلاعات آپریس



تحلیل نوین داده های اسپادانا



توسعه تجارت الکترونیک توانا



توسعه سخت افزار سوشیان



شرکت تعاونی کارکنان خدمات



ارتباطی رایتل



پندار کوشک ایمن



پیشتازان اندیشه پویا



داتیس آریانا تیم



داده ورزی هوشمند آسا



داده پردازی راسا



داده پردازان نیکو ارتباط



رایمند ارتباطات نقش جهان



رایا پرداز برین



سرو حامی پارس



سام تجهیز خاورمیانه ایرانیان



سرآمدان شبکه و ارتباط نوین



شرکت فنی مهندسی نوآوران



افرا تک هوشمند



شبکه گستر مهر تجارت



فاوا پردازش خلیج فارس



شرکت مهندسی پیما عمران نیرو



شرکت مهندسی آموزندگان سیستم



شرکت شبکه بانان پاژ



فهرست اعضای حقوقی

انجمن انفورماتیک ایران

کمان امن دیاکو



گسترش داده باران آذربایجان



گروه بازرگانی آریانا پارس راژمان



لایف سرویس پارسه



شرکت مهندسی فرداد رایانه
اسپادانا (سهامی خاص)



مهندسی شبکه گستر



مهندسی افق داده ایرانیان



مهندسین مشاور ارتباط گستران شرق



نویان ابر آوران



ویرا رسانه افزار



نوآوران ارتباطات سیمای امید W



نوبین داده پرداز روناش



فن آوری اطلاعات گرین پی



اعضای حقوقی فعال در حوزه مدیریت پروژه

سامانه ورز هزاره



اعضای حقوقی مشتریان بهره‌بردار از
راه‌حل‌های نرم‌افزارهای پیشرفته سازمانی
و بهره‌بردار از فناوری اطلاعات

آجرپایه



سامان کاربن فناوری هیراد



شبکه گستر نسل جدید



لیان شبکه گستر سیراف



فناوران امن آرنیکا



فراگامان سایا ارتباط نوین



فراز اطمینان کیا مهر



فناوری اطلاعات سهلان



فرست کیش



مهندسی ارتباطات دوربرد فارس



مهندسی رایانه پرداز طوفان



مهندسی نرم افزاری گلستان



مفتاح رایانه افزار



هوشمند اندیشان فرایند



شبکه پایدار فناوری اطلاعات



ویرا انرژی مهرماهان



صنایع پرسو الکترونیک



شرکت مهندسی گسترش ارتباطات نو
خاورمیانه



شرکت تجارت سرور ماندگار



شرکت آوای همراه هوشمند
هزاردستان (کافه بازار)



شرکت مهندسی سیستم‌های اطلاعاتی
پیشرو



شبکه اندیش آژمان



شرکت فناوران آتیه گئومات



فناوران رایان کهکشان



رایمند ارتباطات نقش جهان



فناوری ارتباطات هوشمند لیان



تحلیلگران اطلاعات نگاره



توسعه فن آوری ارتباطات و اطلاعات
راهکار مفید پرداز



شرکت رایانه همراه کیان



رایکا شبکه هوشمند



رادمان داده‌پردازی فرنام



خدمات آواژنگ



خدمات رایانه ای بهینه پرداز پویا



سرو رایانه



شرکت نوآوران انفورماتیک



فهرست اعضای حقوقی

انجمن انفورماتیک ایران

فناوری اطلاعات مدبران عصر
فرا اندیشه



مرکز تحقیقات هوش مصنوعی
مبتنی بر سازه پلیمری



داده پردازی خوارزمی



اعضای حقوقی فعال در حوزه مشاوره

توسعه الکترونیک ماهان



داده پردازان ویرا ارتباط



طرح و توسعه الکترونیک افق



شرکت شبکه گستر سینا



نرم افزاری ژابیز پردا



مشاوره مدیریت و مهندسی آرا



اعضای حقوقی فعال در حوزه توسعه وب فارسی

گروه اقتصادی و فناوری اهلیت و
اصلیت ماندگار



شرکت آسیا سرمایه هوشیار



اعضای حقوقی فعال سایر حوزه

ایده های تجارت کاسپین



آنیک الکترونیک کویر یزد



آوا پژوهان هیوا



نوین معتمد اقتصاد ایرانیان (نما)



نوساز محاسب صفاهان



ماندگار اندیشه آوا هوشمند



فرآیند کنترل نقش جهان



فناوری اطلاعات و ارتباطات
اندیشمند آریا



اعضای حقوقی فعال در حوزه آموزش و پژوهش

آینده سازان هوش عدد



آکادمی یاسان



الماس رایان ایرانیان



اندیشه پردازان چوگان



صالح اندیشان ماندگار



نوین طب بنیان راد



پژوهشگاه ارتباطات و فناوری اطلاعات



شرکت توسعه فناوری زعیم



توسعه فناوری ایده ماندگار



صالح اندیشان ماندگار



فناوری اطلاعات بین الملل



ایریسا



ایده پردازش دانش ساخت



ایمن پردازشگران طلوع فارس



پیام ارتباطات نوین عرش



رادمان ارتباط هوشمند افزار



تجارت الکترونیک تدبیر کیش



چکاوک پردازش هوشمند البرز



مجمع همپردازان امروز آسیا



مجمع صنعتی سپاهان باتری



معماران به طرح



مرکز گسترش فناوری اطلاعات (مگفا)



موسسه شهر مدیریت آرال (میراگل)



راه آرمان مهر نیکان



شرکت توسعه خدمات مراکز داده پیشرو



شرکت سروش آفرینان دیبا



شرکت داده نوین خوش سلامت



فرهنگ و توسعه کندو



فهرست اعضای حقوقی

انجمن انفورماتیک ایران

استاد انرژی



توسعه تجارت الکترونیک نوظهور



پیشگامان تحول سبز آریا



پردازش اطلاعات نیک آیین ثمین



کارامیس



فناوری ثروت یارا



نوید طلوع پارسیان



مهرداد آریا



نیکنام تجارت ثمین شرق



آنی نوین هوشمند

